

Propositions As Sessions Taken Apart

Wen Kokke



Doctor of Philosophy
Laboratory for Foundations of Computer Science
School of Informatics
University of Edinburgh

2024

Abstract

The foundations of functional programming are built on the λ -calculus, a powerful model of computation whose canonicity is affirmed by its correspondence with intuitionistic logic. The foundations of concurrent computation are built on less firm ground. There is a wide variety of process calculi, but none enjoy the same canonicity as the λ -calculus, nor an exact correspondence with logic. Since its inception by Girard [1987], Classical Linear Logic has been believed to have some relation to concurrent computation, which has spawned a wealth of research in both logic and programming language theory.

This thesis continues the work towards a foundation for concurrent computation, starting from the propositions-as-sessions correspondence proposed by Wadler [2012]. Drawing on the work of Abramsky [1994], Bellin and Scott [1994], and Caires and Pfenning [2010], among others, Wadler proposed the process calculus CP, which has an exact correspondence with Classical Linear Logic. Drawing on the work of Honda [1993], Honda et al. [1998], and Gay and Vasconcelos [2010], among others, Wadler proposed the session-typed functional language GV, which provides a practical foundation for session-typed concurrency in functional languages. Finally, Wadler connects GV and CP by means of an operational correspondence, and, thereby, connects practice of session-typed concurrency to Classical Linear Logic.

This thesis provides an in-depth study of the propositions-as-sessions correspondence proposed by Wadler, highlighting its strengths and repairing a number of its shortcomings. Due to its adoption of commuting conversions, Wadler's CP is non-confluent, and its semantics are difficult to realise using the abstractions of concurrent computation it purports to model. I address this shortcoming by removing the commuting conversions, and showing that the resulting process calculus is well-behaved and retains its connection to linear logic. Due to its lack of a parallel composition operator, Wadler's CP is ill-suited to an analysis using standard concurrency theory. Wadler's GV lacks an operational semantics and accounts of polymorphism and replication, which were provided by Lindley and Morris [2014; 2015; 2016b]. Unfortunately, Lindley and Morris' GV suffers from similar problems as CP with regard to its treatment of parallel composition, which complicates its metatheory. I address this shortcoming in both systems by adding a parallel composition operator and showing that the resulting calculus is well-behaved and correspond to a hypersequential variant of Classical Linear Logic. Finally, I study Priority CP, a variant of Wadler's propositions-as-sessions correspondence proposed by Dardha and Gay [2018], which increases its expressivity at the cost of its compositionality and correspondence to logic. As with CP, I remove the commuting conversions from Priority CP, and complete the correspondence with Priority GV, which is the analogous variant of GV.

Lay Summary

In recent decades, programming languages have widely adopted *data type checking*, which are a mechanism that helps programmers write safe and secure programs by ruling out data errors. An example of a *data error* would be if the program that manages bank accounts handled the account balance as text rather than as a number. If my balance is £50 and I deposit £20, I would expect my balance to become £70, not the text ‘5020’.

When a programming language is designed with care, it is possible to automatically determine if a program is free from data errors by checking the types of ways in which some data is handled (data types) against each other (e.g. the account balance is handled as a number here but as text there) or against the programmer’s stated intention (e.g. the account balance is handled as text, but the programmer stated that it should always be handled as a number).

With the rise of the internet and the end of Moore’s Law, *concurrent programs* are becoming increasingly important. A concurrent program consists of multiple processes that run at the same time and may share resources or communicate by passing messages. Examples include your browser communicating with a server, an ATM communicating with a bank, and the thousands of graphics processors involved in rendering your favourite game.

Concurrent programs are vulnerable to erroneous behaviours that cannot be prevented by data types. For an example of such an error, let us consider what would happen if the program that manages bank accounts handled each deposit or withdrawal in three steps: (1) read the current balance, (2) calculate the updated balance by adding the amount deposited or subtracting the amount withdrawn, and (3) write the updated account balance. If the account balance for the joint account that I share with my partner is £100, I use one ATM to deposit £60, and my partner uses another ATM to withdraw £50, then we would expect our balance to become £110. If the two interactions happen in sequence, this is what happens. However, each ATM reads and writes the account balance separately, and these reads and writes may be interleaved in any order:

- If my ATM reads the initial balance of £100, then my partner withdraws £50, and then my ATM writes the updated balance of £100 plus the deposited £60, the final balance becomes £160.
- If my partner’s ATM reads the initial balance of £100, then I deposit £60, and then my partner’s ATM writes the updated balance of £100 minus the withdrawn £50, the final balance becomes £50.

This kind of error is known as a *race condition*. Race conditions are difficult to diagnose since, while the program is fundamentally incorrect,

the erroneous behaviour might not happen most of the time. After all, how often do my partner and I use our joint account at exactly the same time? So if I, unaware of my partner's simultaneous withdrawal, call my bank to complain that "I deposited £60, but my balance went down!", an engineer might test the bank's program by depositing £60, see the account balance go up, and conclude that everything is working as intended.

Behavioural type checking is a mechanism which aims to rule out race conditions and other such erroneous behaviours. Whereas data types describe the way that a program must always handle some data (e.g. the account balance must always be handled as a number), behavioural types describe the evolving way a program must interact with its environment over time (e.g. usually, the bank's program should accept any request for a transaction, but while a transaction is ongoing, it should not accept requests for other transactions on the same account).

Most mainstream programming languages do not support behavioural types and, as such, it is not generally possible to automatically determine if a program is free from behavioural errors. In practice, writing correct concurrent programs comes down to the programmer correctly reasoning about how multiple processes might interact and correctly diagnosing and repairing errors, both of which are incredibly difficult tasks.

One reason that behavioural types have not yet seen widespread mainstream adoption is that the mathematical foundation underlying them is several decades younger and less well understood than the mathematical foundation underlying data types.

In this thesis, I investigate several proposed theories that could become the mathematical foundation underlying behavioural types for message-passing communication in concurrent programs. I identify several shortcomings in the proposed theories and, on the basis of my findings, propose a new mathematical foundation for behavioural types. I develop the theory necessary to integrate the proposed behavioural types into existing programming languages and describe a proof of concept implementation.

Acknowledgements

I am grateful to my supervisor, Philip Wadler, and my secondary supervisors, Sam Lindley and J. Garrett Morris, for their support throughout my Ph.D.

I am grateful to my examiners, James Cheney and Peter Thiemann, for accepting to act as my examiners and for their feedback on this thesis.

I am grateful to the proverbial village—to all past members of the ABCD research group, to all past and present members of the Scottish Programming Languages Community, and especially to Robert Atkey, Ornela Dardha, Simon Fowler, and James McKinna.

I am grateful to Fabrizio Montesi and Marco Peressotti for the kindness they showed me upon discovering our parallel developments, and for exposing me to concurrency theory.

I am grateful to Jonathan MacBride and the lovely folks at the Staff Pride Network for making the University of Edinburgh a kinder and safer place.

I am grateful to Sabine Weber for founding and organising the Informatics Open Art Space with me, to all the wonderful folks who joined us every week, and to Tekevwe Kwakpovwe and Marine Demarty for continuing to organise the Open Art Space.

I am grateful to Andrew McLeod and Mihaela Dragomir for organising my first-year trip to the Firbush Outdoor Centre, to the Firbush staff for starting my obsession with windsurfing and sailing, and to the many, many folks who helped me organise and attend nearly a dozen trips to Firbush between 2016 and 2020.

I am grateful to the many, many organisers of CDT Pizza, and especially to Kate McCurdy, Aidan Marnane, Sigrid Passano Hellan, and Siobhán Carroll for organising the 2019 series with me, to the EPSRC Centre for Doctoral Training in Data Science for funding CDT Pizza, and to the CDT in Data Science community for welcoming outsiders with an interest in data science and free pizza.

I am grateful to my mother, Marieke Helsen, and my secondary parents, Odile Jansen and Louis Kolkman, for their support throughout my Ph.D.

I am grateful to my partners, pals, and pets, who have greatly supported me throughout my Ph.D., and especially to April Gonçalves and Dora Gonçalves-Kokke, Grae Rose, Inari Listenmaa, and Zaza Helsen.

The EPSRC Centre for Doctoral Training in Pervasive Parallelism (EP/L01503X/1) supported my studies at the University of Edinburgh.

The Cost Action CA15123 (EUTypes) supported my visit to Fabrizio Montesi and Marco Peressotti at the University of Southern Denmark.

The SIGPLAN Professional Activities Committee (PAC) supported my attendance of the 23rd and 24th International Conference on Functional Programming (ICFP'18 and ICFP'19) and the 46th Symposium on Principles of Programming Languages (POPL'19).

Table of Contents

1	Introduction	11
1.1	Contributions	17
1.2	Conventions	24
1.3	Thesis Structure	25
I	Classical Processes Revisited	27
2	Classical Processes Revisited	29
2.1	Classical Processes	31
2.1.1	Process Structure	34
2.1.2	Link	35
2.1.3	Send and Receive	36
2.1.4	Close and Wait	38
2.1.5	Select and Offer	38
2.1.6	The Absurd Offer	39
2.2	Metatheory	40
2.2.1	Preliminaries	41
2.2.2	Preservation	48
2.2.3	Progress	51
2.2.4	Duality, Dependency, and Deadlock	53
2.2.5	Connection and Right-Branching Form	60
2.2.6	Observational Equivalence	66
2.3	Commuting Conversions Considered Bad	66
2.4	Conclusion	72
2.5	Omitted Proofs	73
II	Taking Linear Logic Apart	77
3	Hypersequent Classical Processes	79
3.1	Hypersequent Classical Processes	85
3.1.1	Process Structure	89
3.1.2	Link	92
3.1.3	Send and Receive	92
3.1.4	Close and Wait	93
3.1.5	Select and Offer	94
3.1.6	The Absurd Offer	94
3.2	Metatheory	95
3.2.1	Preliminaries	95
3.2.2	Preservation	103
3.2.3	Progress	106

3.2.4	Duality, Dependency, and Deadlock	108
3.2.5	Connection and Disentanglement	114
3.2.6	Multiset CP	123
3.2.7	Fission, Fusion, and Disentanglement	126
3.2.8	Label-Transition System and Harmony	134
3.3	Discussion	139
3.3.1	Zap: the Exceptionally Absurd Offer	140
3.3.2	Synchronous Links and Lazy Forwarders	147
3.3.3	Variant Types and Guarded Summation	150
3.3.4	Channel or Endpoint Names?	157
3.3.5	Deep Inference and Display Calculus	159
3.3.6	Hypersequent Calculus	160
3.3.7	Hypersequents, MIX, and MIX ₀	160
3.4	Conclusion	160
3.5	Omitted Proofs	161
4	Hypersequent Good Variation	171
4.1	Legend and Errata	174
4.2	Paper I: Separating Sessions Smoothly	176
4.3	Discussion	230
4.3.1	Cooking GV With Cut	230
4.3.2	Lock Types, Cuts, and Hypersequents	231
4.4	Conclusion	234
III	Prioritise The Best Variation	235
5	Priority Classical Processes & Priority Good Variation	237
5.1	Legend and Errata	240
5.2	Paper II: Prioritise the Best Variation	243
5.3	Discussion	292
5.3.1	Relation to Classical Processes	292
5.3.2	Priority Inference	298
5.4	Conclusion	304
IV	Deadlock-Free Session Types in Linear Haskell	307
6	Implementations	309
6.1	Legend and Errata	310
6.2	Paper III: Deadlock-Free Session Types in Linear Haskell . .	313
6.3	Conclusion	328
A	Glossary	329
A.1	Graph Theory	329
A.2	Multisets	332
	Bibliography	335

Chapter 1

Introduction

The foundations of functional programming are built on the λ -calculus, a powerful Model of Computation whose canonicity is affirmed by its correspondence with intuitionistic logic. The foundations of concurrent computation are built on less firm ground. There is a wide variety of process calculi, but none enjoy the same canonicity as the λ -calculus, nor an exact correspondence with logic.

Since its inception by Girard [1987], Classical Linear Logic (CLL) has been believed to have some relation to concurrent computation, which has spawned a wealth of research in both logic and programming language theory.

This thesis continues the work towards a foundation for concurrent computation, starting from the propositions-as-sessions correspondence proposed by Wadler [2012].

- Drawing on the work of Abramsky [1994], Bellin and Scott [1994], and Caires and Pfenning [2010], among others, Wadler proposed the process calculus Classical Processes (CP), which has an exact correspondence with Classical Linear Logic.
- Drawing on the work of Honda [1993], Honda et al. [1998], and Gay and Vasconcelos [2010], among others, Wadler proposed the session-typed functional language Good Variation (GV), which provides a practical foundation for session-typed concurrency in functional languages.
- Finally, Wadler connects GV and CP by means of an operational correspondence, and, thereby, connects practice of session-typed concurrency to Classical Linear Logic.

The names CP and GV are unstated homages to the authors of the work that inspired them: Classical Processes was based on π DILL by Caires and Pfenning [2010; later Caires et al., 2016]; and Good Variation was based on LAST by Gay and Vasconcelos [2010; the name LAST was given by Lindley and Morris, 2015].

Let us start by discussing Wadler’s propositions-as-sessions correspondence at a glance. The paper contains two separate correspondences:

1. CP demonstrates that the original and most well-known sequent calculus for full CLL, exactly as presented by Girard [1987, pp. 22 and 26-27], can be interpreted as something that syntactically resembles a process calculus. I intentionally avoid calling CP a process calculus—whether or not it *is* is debatable—but it can hardly be denied that CP *looks* like a process calculus.
2. GV, together with the translations from GV to CP [Wadler, 2012, § 3.1; Lindley and Morris, 2015, § 3.2] and from CP to GV [Lindley and Morris, 2015, § 3.1], demonstrates an operational correspondence between CLL and a session-typed concurrent λ -calculus.

Only the first—the correspondence between CLL and CP—is one that I would characterize as a Curry-Howard correspondence:

- It relates a logic and a typed Model of Computation.
- It is complete. Propositions correspond to types, proofs correspond to processes, and cut elimination corresponds to computation.¹
- It is exact. Going from the logic to the Model of Computation, and vice versa, requires no hard work. It is simply a change of notation.

The correspondence is not a *profound coincidence*. CP was *intentionally constructed* to correspond to CLL. Moreover, its inspiration, π DILL, was intentionally constructed to correspond to DILL. A constructed correspondence undermines what is arguably the most important role of Curry-Howard correspondences. It does nothing to reassure us that either system is profound, in the sense of a correspondence between two systems that were independently devised.

One could argue that these correspondences relate the π -calculus to linear logic—in which case they would be profound—but I believe CP and π DILL are too far removed from the π -calculus for this to be true.

Having rejected these correspondences as profound coincidences, let us take a different view: CP *anticipates* a Curry-Howard correspondence. The typed λ -calculus, the most popular typed Model of Computation, corresponds exactly to intuitionistic logic. Linear logic has been said to be relevant to concurrent computation since its inception. So why not cut out the middle-man of independent discovery and directly employ linear logic as a typed Model of Concurrent Computation (MoCC)?

¹The correspondence between CP and CLL is actually incomplete, in the sense that the structural congruence of CP does not correspond to anything in CLL. The equations of the structural congruence are merely valid rewrite rules for proofs. This is easily missed: λ -calculus does not have a structural congruence, and as such the list of correspondences for CP looks as complete as the list for the Curry-Howard correspondence.

To evaluate whether or not CP is successful under this particular view, we must examine what we want from a typed MoCC:

- What is the intended purpose for our typed MoCC?
- Do we want our typed MoCC to be as expressive as the π -calculus?
- What properties do we want our typed MoCC to have?

To me, the purpose of a type system is to rule out erroneous programs. The price of a type system is *necessarily* a reduction in expressivity—for the silly reason that we can no longer express erroneous programs, and for the serious reason that the exact boundary between erroneous and non-erroneous is difficult to capture.

- I do not want to evaluate programs with data-type errors such as `"hello" + true`. The price paid is more difficulty with code reuse [see, e.g. McBride, 2010, Dagand and McBride, 2014, Stump, 2017]
- I do not want to evaluate programs that loop indefinitely, such as `x = x`. The price paid is more difficulty with recursive algorithms due to the undecidability of the halting problem [Davis, 1952, Kleene, 1952, Davis, 1958].
- I do not want to evaluate programs that deadlock. The price paid is more difficulty with certain concurrent communication structures.

It seems unreasonable to expect our typed MoCC to be as expressive as, say, the π -calculus, simply because the π -calculus includes all of the above erroneous behaviors. However, this raises an important question: What behaviors are erroneous?

Let us consider Mazza [2018], who paints a pessimistic picture of the “proofs as processes” program. They argue that neither CLL nor DiLL² [Differential Linear Logic, Ehrhard and Regnier, 2006, Ehrhard, 2018] can satisfactorily encode even an elementary process calculus, as neither can express *confusion*—a kind of non-local non-determinism in which two communications mutually exclude each other. This raises an immediate question: *Do we want confusion?* The question alone elicits an immediate and visceral “No!” from the audience, but once our gut reaction subsides, we should evaluate the question formally.

Mazza [2018] demonstrates that confusion is present in even the most elementary MoCCs, and that confusion is preserved by any encoding that preserves the degree of distribution—i.e. any encoding that is a homomorphism on parallel composition. Therefore, there exists a class of processes for which we must choose (a) to allow confusion, or (b) to accept a reduction in the degree of distribution. Networking algorithms commonly allow for confusion—e.g. *leader election* only requires that

²Note that DiLL refers to Differentiable Linear Logic, whereas DILL refers to Dual Intuitionist Linear Logic. The only difference in writing is the capitalization of the “i”. Presumably, both are pronounced like the herb.

some leader is elected. Therefore, I do not feel comfortable rejecting confusion outright. However, accepting confusion means rejecting strong normalization and the Church-Rosser property. To me, those properties are synonymous with safety, and I cannot reject them outright either.

We may want to allow our programming languages to write programs with local non-determinism or confusion, but we should never allow the programmer to write such a program *by accident*. This seems to require that we capture programs with these properties using some modality. CLL captures non-linearity using the exponentials. DiLL captures failure and local non-determinism using the co-exponentials. I choose not to commit to the wholesale acceptance or rejection of non-determinism and confusion. Rather, my stance is that all is not lost if our typed MoCC does not support confusion. We have simply taken the first step on our journey of trying to claw back the steep price we paid for types.

Dependently-typed λ -calculus recovers a lot of the expressivity of the λ -calculus by allowing output data to depend on input data. For instance, it can express any recursive function whose termination can be proved within the type system of the calculus itself. It stands to reason that a typed process calculus could likewise recover expressivity by allowing the present communication structure of a program to depend on the past communication structure.

I would argue that to build a dependently-typed process calculus, we need the solid foundation of a simply-typed process calculus. I believe—and intend to convince you in this thesis—that the simply-typed fragment of CP is not yet the solid foundation we are looking for.

- The correspondence between CP and process calculus is lacking, because CP does not easily admit a behavioural theory, e.g. a labelled-transition system and a behavioural equivalence.
- The correspondence between CP and CLL is lacking, because the sequent calculus for CLL is not the canonical proof theory. Proof nets are the canonical representation of CLL proofs! This is an issue with CP as a typed MoCC, because the sequent calculus has a lower degree of distribution than the proof nets.
- CLL itself is lacking as a foundation for a typed MoCC. It does not describe non-determinism, unlike DiLL [Ehrhard, 2018]. It describes parallelism as independence, and its dual as dependence, but cannot describe sequential dependence, unlike Pomset Logic [Retoré, 1997] and BV [Guglielmi, 2007].

The first assignment Philip Wadler gave me as my Ph.D. supervisor was to add support for non-determinism to CP. I did, albeit poorly [Kokke, 2017]. However, while doing so, I discovered more ways in which CP is lacking as a process calculus. Since CP does not have a standalone process construct for parallel composition, the extension I

introduced has *three different* process constructs that contain a parallel composition (two inherited from CP, and a third that permits non-determinism). Parallel composition is the most crucial connective in a process calculus. Not having it as a standalone process construct in CP is untenable. Consequently, I pivoted to restructuring CP to include parallel composition as a standalone process construct and tighten its correspondence with the π -calculus, while retaining its tight correspondence with CLL. This work is the principal focus of this thesis and is presented in Chapter 3 and Chapter 4. At times, I will allude to the second and third points, and arguably some of the work on priorities in Chapter 5 applies to the third point. However, I consider them out of scope for this thesis, as I believe that each could correspond to several theses in their own right.

The result should, at least at its core, be a conservative extension of CP. There are three main reasons for wanting this.

1. It retains some correspondence with Classical Linear Logic, even if the correspondence is no longer as obvious as in Wadler’s case.
2. It ensures that any previous work on CP, such as Wadler’s work on polymorphism and unrestricted usage and Lindley and Morris’ work on fixed points, carries over without too much trouble.
3. It eases the future integration of the features offered by DiLL, Pomset Logic, and BV, since these logics are all extensions of CLL.

As I mentioned, my first assignment was to add support for non-determinism to CP, which I did, though in a manner I have always found unsatisfactory [Kokke, 2017]. Later, Qian et al. [2021] found the exact structure I was looking for. In hindsight, the answer was there all along, in logic! DiLL extends CLL with co-exponentials, which make cut elimination non-deterministic [Ehrhard, 2005, 2018]. Qian, Kavvos, and Birkedal’s work, which extends CP with DiLL’s co-exponentials, uses the work on hypersequents presented in Chapter 3.

What role does GV play in the propositions-as-sessions correspondence? CP demonstrates a correspondence between a process calculus and CLL—a correspondence so exact that it requires no proof—but CP is unsatisfactory as a foundation for a programming language for one simple reason: CP does not have functions. In CLL, everything is defined in terms of duality. This matches well with session types and channel-based communication, where two processes act in dual ways, but does not match well with functions. CLL’s implication is *classical*, and does not lend itself to an interpretation as a function type. A significant portion of programming language theory is built on the foundations of functions and the λ -calculus. Hence, adopting CP as the foundation for concurrent programming means parting ways with a significant body of work.

In essence, GV extends CP with functions. Extending CP with a function type requires extending CLL with an intuitionistic implication, which is

more complicated than simply adding a connective, as CLL’s sequents are fundamentally classical. There are some approaches to mixing classical and intuitionistic connectives in the literature on logic, such as Display Logic [Belnap, 1982]. However, none of these approaches have the widespread acceptance of CLL, nor are they easily amenable to a term calculus. Instead, Wadler [2012] approaches GV from the perspective of concurrent λ -calculus, and adapts Gay and Vasconcelos’ LAST [Gay and Vasconcelos, 2010] to fit a correspondence with CP. From a theoretical viewpoint, GV embeds an axiomatisation of CP into the linear λ -calculus using constants. From a practical viewpoint, GV resembles the way in which concurrency is exposed to the user in real-world programming languages, such as the POSIX Threads API for the C programming language. GV offers the user an API—a collection of constant functions which create threads and channels, send messages, etc—but concrete threads, parallel composition, and name restriction are not representable in the static language. The user cannot write a program that is the parallel composition of two processes, only a program which creates that configuration. This makes it very easy to implement GV’s safe concurrency primitives as a library in existing programming languages, on top of the language’s unsafe concurrency primitives, and doubly so if the languages already supports some form of linear types [see, e.g. Lindley and Morris, 2016a, Kokke, 2019, Kokke and Dardha, 2021b].

To put it plainly, CP is a theoretical tool for studying foundational well-behaved concurrent systems, but GV is what you actually implement.

What role does the correspondence between CP and GV play? Wadler [2014, p. 385] writes that the correspondence formalises, for the first time, “a tight connection between a standard presentation of session types and linear logic”, which formally confirms the previously assumed connection. However, this leaves us with a question. Now that the connection is formally confirmed, what is the purpose of continuing to maintain the correspondence as we make changes to CP and GV? I am motivated to do so for two reasons.

Firstly, an exact correspondence, especially one that preserves parallel composition, reassures us that GV is correct, in that the communication primitives of GV capture the same communication structures as CP. Unfortunately, constructing an exact correspondence requires some amount of tedium, as one must show that CP’s communication channels can correctly emulate GV’s functions and data structures. This is well-established and an area of GV significantly less likely to contain mistakes—or, at least, mistakes that are easily caught by a translation into process calculus. Consequently, any part of GV that does not interact with concurrency should be omitted from the correspondence.³ To further

³Wadler’s GV only has the concurrency primitives, product types—which are used to type the receive primitive—and the unit type—which is used to type the primitive that

simplify the correspondence, the translation from GV to CP is commonly factored into two translations: a translation from GV into fine-grain GV, which removes higher-order control flow; and a translation from fine-grain GV into CP.⁴

Secondly, the translations between CP and GV are helpful in guiding you to the correct typing and implementation of new concurrency primitives. For instance, Wadler [2014, p. 409] considers several alternative designs, but remarks that “these designs are difficult to translate into CP, which suggests they may suffer from deadlock.” Indeed, the suggested alternatives *do* suffer from deadlock. Hence, by following the translation, Wadler correctly chose to abandon alternative concurrency primitives, even though they appear more principled.

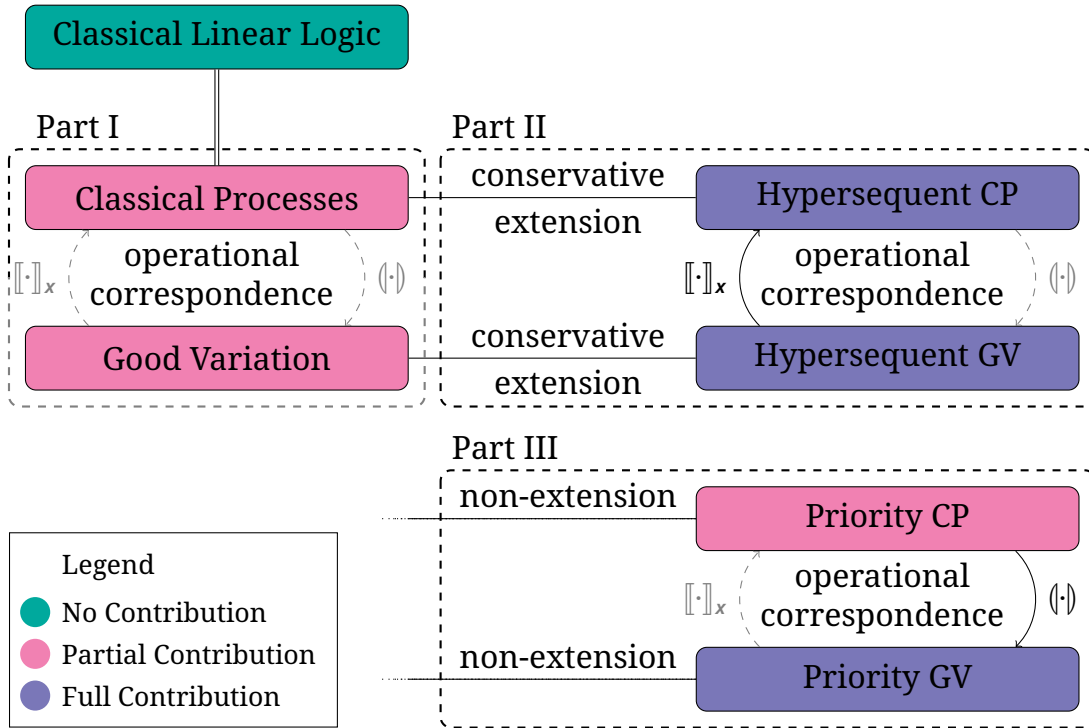
1.1 Contributions

The contributions of this thesis are centred around Wadler’s propositions-as-sessions correspondence, and are divided into four parts. Part I is concerned with Wadler’s propositions-as-sessions correspondence. It discusses CP and its metatheory. Part II is concerned with the hypersequential variants. It discusses Hypersequent CP (HCP), Hypersequent GV (HGV), their metatheory and operational correspondence, as well as their relation to CP and GV. Part III is concerned with the priority-based variants. It discusses Priority CP (PCP), Priority GV (PGV), their metatheory and operational correspondence, as well as their relation to CP and GV. Part IV is concerned with implementations of Good Variation and Priority GV in Linear Haskell.

The following visualisation gives an overview of the landscape of the formal systems discussed in this thesis, and the level of my contributions.

closes a channel. Lindley and Morris [2014] add polymorphism and replicated sessions to GV, which correspond to CP’s polymorphism and exponentials. Lindley and Morris [2015] add sum types to GV, which are used to replace the concurrency primitives for choice by simply sending and receiving a value of a sum type.

⁴The clearest presentation of this decomposed translation can be found in Fowler et al. [2023], which will be presented in Chapter 4, but its origins can be traced back all the way to Lindley and Morris [2014], where fine-grain GV is referred to as HGV π . (Note that the “H” in HGV π stands for “Harmonious”, not “Hypersequent” as in Chapter 4.).



The sole *green* box signifies that I do not claim to have made any serious contribution to Classical Linear Logic. The *pink* boxes signify that, while Classical Processes, Good Variation, and Priority CP were not originally developed by me, I have since contributed to these theories in one way or another. The *periwinkle* boxes signify that Hypersequent CP, Hypersequent GV, and Priority GV were developed by me and colleagues.

The arrows represent the translations and operational correspondences between the various systems. The arrows are labelled by the conventional names of these translations. By convention, any translation from a variant of GV to a variant of CP is denoted by double square brackets with an endpoint name as a subscript, i.e. $\llbracket \cdot \rrbracket_x$, and any translation from a variant of CP to a variant of GV is denoted by double parentheses, i.e. $\langle \cdot \rangle$. The solid arrows—the translation from Hypersequent GV to Hypersequent CP and the translation from Priority CP to Priority GV—and were developed by me and colleagues. The translations corresponding to the dashed and greyed-out arrows are not discussed in this thesis.

The solid lines represent the extension and non-extension results for the various systems with respect to CP and GV, and were developed by me and colleagues.

The dashed boxes—labelled Part I, Part II, and Part III—roughly group the systems by the part of this thesis in which they are discussed.

In the remainder of this section, I give a detailed account of the contributions made in this thesis.

Part I: Classical Processes Revisited

Part I consists of Chapter 2, which revisits Classical Processes. The principal focus is to present, in detail, a variant of CP without commuting conversions, which break confluence in Wadler’s CP, and show that this variant has adequate canonical forms.

Chapter 2 contains the following contributions by me:

- I drop the commuting conversions from the reduction semantics, which cause Wadler’s CP to be non-confluent, and prove that *progress* (Proposition 2.32) continues to hold, albeit with a different canonical form (Definition 2.30).

The reduction semantics without commuting conversions first appeared in Kokke [2017].

- I formalise the relation between CP with and without the commuting conversions. I prove that (1) any process in canonical form can reduce to a process in Wadler’s canonical form using only commuting conversions (Proposition 2.57); and (2) any reduction with commuting conversions is equivalent to a reduction without commuting conversions followed by a reduction using only commuting conversions (Proposition 2.58).

To the best of my knowledge, this is the first publication in which these properties are explicitly stated and proven for CP, though they are hinted at by Lindley and Morris [2016b].

- I formalise the notion of *dependency graph* (Definition 2.35) and *deadlock* (Definition 2.38) for CP, prove that CP processes are deadlock-free (Proposition 2.39), and prove that my canonical forms are adequate, by showing that processes in canonical form are blocked on free endpoints (Corollary 2.47).

To the best of my knowledge, this is the first publication which explicitly defines dependency and deadlock for CP, though it is related to Lindley and Morris’ definitions of dependency and deadlock for GV, and to PCP’s priority constraints [Dardha and Gay, 2018]. The proof of deadlock freedom is adapted from a similar proof of deadlock freedom for GV by Lindley and Morris [2015]. My characterisation of “blocked on free endpoints” is stronger than Lindley and Morris’ characterisation, which corresponds to my Proposition 2.34.

- I formalise the notion of *connection graph* for CP (Definition 2.49), prove that CP’s connection graphs are trees (Proposition 2.50), prove that every CP process can be rewritten into *right-branching form* (Proposition 2.51), and prove that connection graphs are canonical representatives of processes up to a restricted structural congruence (Proposition 2.53).

I have informally referred to the connection graph as the *process structure* or *communication structure* [Kokke and Dardha, 2021a,b]. The definition and corresponding theorems are adapted from work by Simon Fowler for Hypersequent GV [Fowler et al., 2023, Definition 3.9], who formulate the *abstract process structure* of the free names in a process in terms of its typing environment and a set of co-names.

Good Variation is introduced in Chapter 4, together with the proof that HGV is a conservative extension of GV, and the variant of GV with cuts is described in § 4.3.1 and § 4.3.2.

I believe that my changes to CP (in Chapter 2), are more important than my changes to GV (in § 4.3.1). Unlike GV, CP has an exact correspondence with CLL. Furthermore, my changes affect its reduction semantics in a fundamental way, and these changes are used in HCP and PCP. My changes to GV, on the other hand, are immediately superseded by the introduction of hypersequents in HGV in Chapter 4.

If you require a detailed discussion of GV without hypersequents, I recommend Fowler’s Ph.D. thesis, *Typed Concurrent Functional Programming with Channels, Actors, and Sessions* [Fowler, 2019, Chapter 3].

Part II: Taking Linear Logic Apart

Part II consists of Chapter 3 and Chapter 4. Chapter 3 introduces Hypersequent CP, which is a session-typed process calculus based on CP with a tighter correspondence to the π -calculus. Chapter 4 introduces Hypersequent GV, which is the corresponding session-typed concurrent λ -calculus.

Chapter 3 contains the following contributions by me:

- I introduce Hypersequent CP with its typing rules and reduction semantics, and prove *preservation* (Proposition 3.30) and *progress* (Proposition 3.35).

The main innovation—HCP’s type system—was developed independently by myself and by Fabrizio Montesi & Marco Peressotti⁵. The reduction semantics were primarily developed by me, and the label-transition semantics were primarily developed by Fabrizio Montesi & Marco Peressotti.

⁵Our initial developments had minor errors. My initial development permitted hyper-environments to occur in all logical rules, which breaks progress. Progress can be repaired for the multiplicative rules, by using delayed actions [as in Kokke et al., 2019a], but not for the additive rules. These errors made it into the first publication [Kokke et al., 2019b]. Fabrizio Montesi & Marco Peressotti’s initial development [Montesi and Peressotti, 2018] typed the terminated process using the empty environment, as opposed to the empty hyper-environment, which admits Mix_0 .

The similarities between my early work on Hypersequent CP and Hypersequent Calculus [Avron, 1991] were noted by Simon Castellan at an ABCD meeting.⁶

- I adapt the notions of *dependency graph* (Definition 3.36) and *deadlock* (Definition 3.39) for HCP, prove that HCP processes are deadlock-free (Corollary 3.42), and prove that my canonical forms are adequate (Corollary 3.50).
- I adapt the notion of *connection graph* for HCP (Definition 3.51), and prove that HCP's connection graphs are forests (Proposition 3.52), and that every HCP process can be rewritten into *right-branching forest form* (Proposition 3.57).
- I formalise disentanglement from Hypersequent CP processes to multisets of CP processes, and prove that disentanglement preserves typing, structural congruence, and reduction.

Disentanglement was first described by Kokke et al. [2019b] as a rewrite relation that preserves provability, but the presentation in this thesis is novel. To the best of my knowledge, this is the first publication which shows that disentanglement preserves structural congruence and reduction.

- I introduce a label-transition semantics for HCP and prove *harmony* between the reduction and label-transition semantics (Proposition 3.93).

The label-transition semantics is a minor variation of the label-transition semantics by Montesi and Peressotti [2018]. To the best of my knowledge, this is the first publication that proves harmony for a reduction and label-transition semantics for HCP.

- I introduce a variant of HCP which uses an alternative *exceptional* semantics for the additive units, with both reduction and label-transition semantics, and extend the proofs of preservation, progress, and harmony.
- I introduce a variant of HCP which uses an alternative *synchronous* semantics for link, based on identity expansion, which, I claim, continues to work in the presence of polymorphism, and discuss the consequent changes in canonical forms, progress, and adequacy.
- I introduce a variant of HCP which uses a combination of variant types and focusing to further tighten the correspondence to the π -calculus by permitting prefixing and (guarded) summation as

⁶My talk at the ABCD meeting on Monday, December 18, 2017, was titled *Taking Apart Classical Processes*. Oddly, it appears to refer to the calculus as Hypersequent CP, which implies that I used the word “hypersequent” *before* I was aware of Avron’s work on Hypersequent Calculus.

syntactic constructs, extend the proof of preservation, and sketch a proof of operational correspondence between HCP and the variant.

Chapter 4 consists primarily of the paper *Separating Sessions Smoothly* by Fowler et al. [2023], written in collaboration with Simon Fowler, Ornela Dardha, Sam Lindley, and J. Garrett Morris.

- We introduce Hypersequent GV with its typing rules and reduction semantics, and prove *preservation* (Theorem I.3.3), the *tree-structure* of connections in configurations (Theorem I.3.14), *global progress* (Theorem I.3.20), the *diamond property* (Theorem I.3.21), and *termination* (Theorem I.3.22).
- We define a translation from GV to HGV (Theorem I.4.3)
- We define a translation from HGV to GV (Corollary I.4.7).
- We prove an operational correspondence between HGV and HCP. We define fine-grain call-by-value HGV (HGV*), define translations from HGV to HGV* and from HGV* to HCP, and prove that the latter translation preserves types (Lemma I.5.9) and is a sound and complete operational correspondence (Theorem I.5.11).

To the best of my knowledge, this is the first publication of a sound and complete correspondence between any variant of CP and GV.

- We define a translation from HCP to HGV.

I co-developed Hypersequent GV and was primarily responsible for the translations and operational correspondence between HGV and HCP.

The remainder of the chapter contains the following contributions by me:

- I introduce a variant of GV that uses cuts rather than lock typing. Consequently, its structural congruence preserves types, which significantly simplifies its metatheory, and I argue that under lock types, the parallel composition construct is equivalent to a cut.

Part III: Prioritise the Best Variation

Part III consists of Chapter 5, which revisits Priority CP, a variant of CP with priorities whose typing rules permit benign cyclic process configurations, and introduces Priority GV, which is the corresponding session-typed concurrent λ -calculus.

Chapter 5 consists primarily of the paper *Prioritise The Best Variation* by Kokke and Dardha [2021a], written in collaboration with Ornela Dardha.

- We drop the commuting conversions from the reduction semantics of Priority CP, which cause Dardha and Gay's CP to be non-confluent, and prove that *progress* (Theorem II.4.4) continues to hold, albeit with a different canonical form.

These changes are similar to the changes made to CP in Part I.

- We introduce Priority GV with its typing rules and reduction semantics, and prove *preservation* (referred to as *subject reduction*, Theorem II.3.5) and *global progress* (Theorem II.3.14).
- We define a translation from PCP to PGV, and prove that the translation preserves types (Theorem II.4.6) and is a sound and complete operational correspondence (Theorems II.4.7 and II.4.10).

I co-developed Priority GV and was primarily responsible for the initial draft of its theory and metatheory. I adapted the changes to the reduction semantics of Priority CP from my previous work on CP.

The remainder of the chapter contains the following contributions by me:

- I prove that PCP is not an extension of CP (Counterexample 5.1).
- I define priority inference for PCP, and prove that priority inference is sound and complete with respect to typing for PCP (Proposition 5.5 and Proposition 5.7).

Part IV: Deadlock-Free Session Types in Linear Haskell

Part IV consists of Chapter 6, which describes a library in Linear Haskell which implements session-typed channels based on GV and Priority GV.

Chapter 6 consists primarily of the paper *Deadlock-Free Session Types in Linear Haskell* by Kokke and Dardha [2021b], written in collaboration with Ornela Dardha.

- We implement session-typed channels in Linear Haskell and argue that a restricted interface to those channels corresponds to GV's channels and therefore enjoys GV's safety guarantees.
- We define a refined variant of Priority GV's type system, which gives exact rather than approximate priority bounds.
- We define the monadic reflection of Priority GV's priority-based type system into a graded linear monad, where the grading is pairs of lower and upper priority bounds.
- We implement priority-based session-typed channels in Linear Haskell and argue that those channels correspond to PGV's channels and therefore enjoys PGV's safety guarantees.
- We compare our Haskell library to existing Haskell libraries for session-typed channels, extending the comparison by Orchard and Yoshida [2017].

I implemented the Haskell library, developed the refined version of Priority GV's type system and its monadic reflection, wrote the initial

draft of the paper, and co-authored the comparison with existing Haskell libraries for session-typed channels.

1.2 Conventions

In this section, I introduce several conventions that I intend use throughout this thesis.

Barendregt’s Convention

My proofs use Barendregt’s Convention [1985]. If a term occurs in a proof, then all bound variables in that term are chosen to be different from each other and all the free variables. This applies equally to variables in λ -calculus terms and endpoint names in process calculus processes. In each case, the proofs can be made more formal by permitting and applying α -conversion where necessary, or by using a nameless term syntax such as deBruijn indices.

Pattern Matching

I use the phrase “is of the form” to imply pattern matching, where any unbound meta variable on the right hand side of that phrase is implicitly existentially quantified, and the entire phrase is converted into an equality, e.g. the phrase “ P is of the form $Q \parallel R$ ” should be interpreted as “there exist some Q and R such that $P = Q \parallel R$ ”.

Syntax Highlighting

Except where noted, each chapter is focused primarily on one system.

The terms of that system are printed in **red**, its types are printed in **blue**, and its annotation—where applicable—are printed in **yellow**, and all three are rendered in a sans-serif font.

To save on accessible colour combinations, the terms, types, and annotations of *any other system* are printed in **pink**, **green**, and **periwinkle**, respectively, and all three are rendered in an italicised font with serif. The relations of other system, such as typing and reduction, are marked by a subscript.

Let us discuss CP’s syntax highlighting as an example. In Chapter 2, which focuses on CP, its processes and types of are printed in **red** and **blue**, respectively, and are rendered in a sans-serif font. However, in Chapter 3, which focuses on HCP, CP’s processes and types are rendered in **pink** and **green**, respectively, and both are rendered in an italicised font with serif. Furthermore, its relations, such as its typing relation, are marked by a subscript C, even when it was unmarked in Chapter 2.

As a matter of personal preference, the metavariables for those syntactic categories that are collections of terms, types, and annotations, such as sets of names or typing environments, are typeset in the same color as the terms, types, and annotations of that system, respectively. However, the constructors of these collections are not typeset in that color, e.g. $\Gamma, x : A \vdash P$ and not $\Gamma, x : A \vdash P$. (The comma and colon are printed in blue in the second statement.) Likewise, functions that act on syntactic categories are not typeset in the corresponding colours, e.g. $\bar{A}$ and not $\bar{A}$.

Omitted Proof Sections

Some chapters end with a section titled “Omitted Proofs”. These sections contain proofs that were omitted from the main text of the chapter. Any proposition whose proof is omitted should provide a summary of the proof that references the relevant omitted proofs section. For lemmas, the proof is simply omitted, and only appears in the relevant omitted proofs section.

Embedded Papers and Pink Pages

This thesis contains three embedded papers. The text in these papers was not composed exclusively by myself, and does not always respect the conventions set out in this section. In an effort to signpost the beginning and end of these papers, they are preceded and followed by one pink page, in the following colour:



The pink page preceding each paper summarises its publication history and my contributions, and the pink page following each paper is empty.

To ensure an easy transition into the text of these papers, the section prior to each paper provide a legend—which summarises the differences in its conventions, notations, and definitions—and an erratum—which clarifies any errors in the paper.

1.3 Thesis Structure

This thesis proceeds as follows.

- In Part I, I present a variant of Classical Processes without commuting conversions, which break confluence in Wadler’s CP, and show that this variant is free from deadlock and has adequate canonical forms.

- In Part II, I present Hypersequent CP, which is a session-typed process calculus based on CP with a tighter correspondence to the π -calculus, and Hypersequent GV, which is the corresponding session-typed concurrent λ -calculus.
- In Part III, I present a variant of Priority CP without commuting conversions and introduce Priority GV which is the corresponding session-typed concurrent λ -calculus.
- In Part IV, I describe an implementation of session-typed channels based on GV and Priority GV in Linear Haskell.

Part I

Classical Processes Revisited

Chapter 2

Classical Processes Revisited

This chapter presents Classical Processes (CP), a session-typed process calculus introduced by Wadler [2012] based on Classical Linear Logic [Girard, 1987, CLL].

Wadler went to great lengths to ensure an *exact* correspondence between CP and CLL:

1. If you erase all that’s written in **red** from the typing rules of CP, you get *exactly* the inference rules of CLL, as given by Girard [1987].
2. The operational semantics of CP resembles that of a process calculus (§ 2.1), but if you examine the proof of progress [Wadler, 2012, Proposition 2], which implements the concrete reduction strategy for CP, the manner in which the reduction rules are used resembles a proof of cut elimination for CLL in the style of Gentzen, as described for classical logic by Girard et al. [1989, § 13.2].¹

Wadler chose to make serious concessions to CP as a process calculus in order to achieve its exact correspondence with CLL:

1. He combines name restriction and parallel composition into a single term constructor, corresponding to the logical *Cut* rule.

As a consequence, the structural congruence for CP looks quite unlike that of any process calculus, and many of the concepts used in concurrency theory, such as labelled transition systems, bisimilarity, and observational equivalents, cannot easily be applied to CP. For example, Atkey [2017] defines an observational equivalence for CP but, in order to do so, must introduce a second,

¹It is easy to mistake Wadler’s claim for a correspondence between CP’s reduction and the proof of cut elimination for the sequent calculus for CLL, as given by Girard [1987]. The citations “Girard [1987]” and “Girard et al. [1989]” are similar; and, given the context, it is reasonable to assume that Wadler is referencing a proof by Girard [1987]—except that Girard [1987] gives no such proof! He proves cut elimination for the proof nets of CLL, proves that the sequent calculus is sound and complete with respect to the proof nets, and (implicitly) obtains cut elimination for the sequent calculus as a corollary.

separate term language that once again separates name restriction and parallel composition.

2. He introduces additional reduction rules, named *commuting conversions*, into CP's operational semantics.

The commuting conversions are not canonical for process calculi, and they are difficult to justify if we want CP to model parallel or distributed computation, since they require that if some process is blocked, waiting for input on some channel, then any other process can also become blocked, waiting for input on that same channel, whether they have access to that channel or not, and *without any communication*. (I discuss the commuting conversions in detail in § 2.3.)

The version of CP presented in this section retains the first part of the correspondence, which requires that we keep name restriction and parallel composition combined in a single term construct. There are several approaches to separating the two, which I present in Chapter 3, Chapter 4, and Chapter 5, but each weakens the strict correspondence between the typing rules and the inference rules of CLL. Therefore, I believe that adopting any of these for CP would compromise its status as a canonical reference point.

The version of CP presented in this section weakens the second part of the correspondence by dropping the commuting conversions from the operational semantics. This results in a reduction strategy much closer to that of a process calculus and repairs several defects in Wadler's CP. I believe that adopting this change does not compromise the status of CP as canonical reference point, which I discuss in detail in § 2.3 (see Proposition 2.57 and Proposition 2.58). The reduction strategy presented in this section still resembles a proof of cut elimination, just not in Gentzen's style.

In this chapter, CP's processes are printed in **red**, its types are printed in **blue**, and both are rendered in a sans-serif font. To save on accessible colour combinations, the terms and types of *any other system* are printed in **pink** and **green**, respectively, both are rendered in an italicised font with serif, and any relations, such as typing and reduction, are marked by a subscript.

This chapter proceeds as follows:

- In § 2.1, I introduce CP.
- In § 2.2, I introduce the metatheory for CP.

Notably, I prove *preservation* (Proposition 2.27) and *progress* (Proposition 2.32), that its processes are deadlock-free (Proposition 2.39), that its canonical forms are adequate (Corollary 2.47), and that its connection graphs are trees (Proposition 2.50).

- In § 2.3, I discuss the relation between CP with and CP without commuting conversions.
- Finally, § 2.5 contains all omitted proofs.

2.1 Classical Processes

In this section, I introduce Classical Processes (CP), a session-typed process calculus that was introduced by Wadler [2012] and based on π DILL by Caires and Pfenning [2010]. CP's process calculus resembles the π -calculus [Milner et al., 1992b], and its type system is Classical Linear Logic [Girard, 1987, CLL].

The fundamental notions of programs and computation in CP—as in most process calculi—are processes and message-passing communication. Processes communicate by passing messages over channels. Unlike the π -calculus, CP's channels are *binary*, which means that communication always takes place between two processes. Each channel has exactly two endpoints. Each endpoint is held by exactly one process. Names refer to *channel endpoints* rather than the channels themselves. CP has no notion of multiparty communication, though there are extensions that address this [e.g. Carbone et al., 2016].

In the π -calculus, channel names fulfil two roles. Channel names are used as communication channels and as labels, compared with the conditional operator. In CP, the two roles are separate. Channel endpoints and labels are in different syntactic sorts, and are communicated by different send and receive operators. To simplify the presentation, CP is commonly restricted to the labels 'left' and 'right', written as **inl** and **inr**. This is no less general, as any *finite* set of labels can be encoded as a sequence of binary choices. For any practical implementation, this restriction is easily lifted.

Processes (ranged over by **P**, **Q**, **R**) are defined by the following grammar:

P, Q, R	::=	$x \leftrightarrow y$	link
		$(\nu x \bar{x})(P \parallel Q)$	cut
		$x[y].(P \parallel Q)$	send
		$x(y).P$	receive
		$x[] . 0$	close
		$x().P$	wait
		$x \triangleleft \text{inl}.P$	select left
		$x \triangleleft \text{inr}.P$	select right
		$x \triangleright \{\text{inl}: P; \text{inr}: Q\}$	choice
		$x \not\leq N$	absurd offer

The names **x**, **y**, **z**, and **w** range over the endpoints of communication channels—'channel endpoints' or 'endpoints' for short. The names **$\bar{x}$** , **$\bar{y}$** ,

$\bar{z}$, and $\bar{w}$ as well as a , b , and c also range over endpoints, though they are only used in specific circumstances, which I discuss shortly. The names N and M range over sets of endpoints.

An endpoint is *bound* in the following cases:

- In $(\nu x\bar{x})(P \parallel Q)$, x and $\bar{x}$ are bound in P and Q , respectively.
- In $x[y].(P \parallel Q)$, y is bound in P , but not in Q .
- In $x(y).P$, y is bound in P .

An endpoint is *free* if it is not bound. Notably, for $x \notin N$, x and all names in N are free. I write $\text{fn}(P)$ to denote the set of free endpoints in P . By convention, the names a , b , and c are used as a shorthand to imply to the reader that the endpoint is free.

Two endpoints are *dual* if they are bound by the same name restriction, e.g. in $(\nu x\bar{x})(P \parallel Q)$, x is bound in P , $\bar{x}$ is bound in Q , and x and $\bar{x}$ are dual. By convention, the names $\bar{x}$, $\bar{y}$, $\bar{z}$, and $\bar{w}$ are used as a shorthand to imply duality to the reader, e.g. I use x and $\bar{x}$ when they are dual endpoints of the same channel.

In CP, actions are not a well-defined syntactic sort, unlike in the π -calculus. Nonetheless, I informally write “action” in reference to the bit before the dot, e.g. the action for $x[y].(P \parallel Q)$ is $x[y]$. For the troublemakers without a dot, $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ and $x \notin N$, I write $x \triangleright \text{inl}$, $x \triangleright \text{inr}$, and $x \notin$, respectively. (The syntax $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ is the primary obstacle to separating actions out into their own syntactic sort. I discuss a variant that decomposes the offer construct in this manner in § 3.3.3.)

Types (ranged over by A , B) are the formula of CLL, as defined by the following grammar:

$$A, B ::= A \otimes B \mid \mathbf{1} \mid A \oplus B \mid \mathbf{0} \\ \mid A \wp B \mid \perp \mid A \& B \mid \top$$

Duality plays an important role in CP as in CLL. Viewed from the perspective of a logic, it corresponds to negation. Viewed from the perspective of a process calculus, it guarantees that processes act on dual endpoints of the same channel in dual ways, e.g. one process sends when the other receives. As in CLL, duality is not defined as a type constructor, but as a function on types:

$$\begin{array}{lll} \overline{A \otimes B} \triangleq \overline{A} \wp \overline{B} & \overline{\mathbf{1}} \triangleq \perp \\ \overline{A \wp B} \triangleq \overline{A} \otimes \overline{B} & \overline{\perp} \triangleq \mathbf{1} \\ \overline{A \oplus B} \triangleq \overline{A} \& \overline{B} & \overline{\mathbf{0}} \triangleq \top \\ \overline{A \& B} \triangleq \overline{A} \oplus \overline{B} & \overline{\top} \triangleq \mathbf{0} \end{array}$$

As we will see, dual endpoints have dual types. The notation for duality (A , $\bar{A}$) and the naming convention for dual endpoints (x , $\bar{x}$) were chosen to emphasize this. Duality is involutive.

$$\begin{array}{c}
\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \text{T-LINK} \qquad \frac{P \vdash \Gamma, x : A \quad Q \vdash \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(P \parallel Q) \vdash \Gamma, \Delta} \text{T-CUT} \\
\\
\frac{P \vdash \Gamma, y : A \quad Q \vdash \Delta, x : B}{x[y].(P \parallel Q) \vdash \Gamma, \Delta, x : A \otimes B} \text{T-SEND} \qquad \frac{P \vdash \Gamma, y : A, x : B}{x(y).P \vdash \Gamma, x : A \wp B} \text{T-RECV} \\
\\
\frac{}{x[] . 0 \vdash x : \mathbf{1}} \text{T-CLOSE} \qquad \frac{P \vdash \Gamma}{x().P \vdash \Gamma, x : \perp} \text{T-WAIT} \\
\\
\frac{P \vdash \Gamma, x : A}{x \triangleleft \text{inl}.P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_1 \qquad \frac{P \vdash \Gamma, x : B}{x \triangleleft \text{inl}.P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_2 \\
\\
\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B}{x \triangleright \{\text{inl}: P; \text{inr}: Q\} \vdash \Gamma, x : A \& B} \text{T-OFFER} \qquad \frac{N = \text{fn}(\Gamma)}{x \not\vdash N \vdash \Gamma, x : \top} \text{T-ABSURD}
\end{array}$$

Figure 2.1: Typing Rules for Classical Processes

Lemma 2.1. $\bar{\bar{A}} = A$

Typing environments (ranged over by Γ, Δ) are sets of type assignments, as defined by the following grammar:

$$\Gamma, \Delta ::= \cdot \mid \Gamma, x : A$$

The set of free endpoint names in a typing environment, written $\text{fn}(\Gamma)$, is defined by recursion on the typing environment, i.e. $\text{fn}(\cdot) \triangleq \emptyset$ and $\text{fn}(\Gamma, x : A) \triangleq \text{fn}(\Gamma) \cup \{x\}$. The extension $\Gamma, x : A$ is only defined when x is not free in Γ , i.e. $x \notin \text{fn}(\Gamma)$.

I write Γ, Δ for the concatenation of typing environments Γ and Δ . The concatenation Γ, Δ is only defined when the names in Γ and Δ are unique, i.e. $\text{fn}(\Gamma) \cap \text{fn}(\Delta) = \emptyset$.

The typing judgment $P \vdash \Gamma$ means that P is well-typed if, for each type assignment $x : A$ in Γ , exactly one process in P uses the endpoint x according to the session type A .

Definition 2.2 (Typing). A process P is well-typed under typing environment Γ if there exists a derivation with conclusion $P \vdash \Gamma$ that uses the typing rules in Figure 2.1.

The names $\mathcal{E}$ and $\mathcal{F}$ range over *evaluation contexts*. Evaluation contexts are one-hole process contexts that consist only of cuts.

Definition 2.3 (Evaluation Context). Evaluation contexts are one-hole process contexts, as defined by the following grammar:

$$\mathcal{E}, \mathcal{F} ::= \square \mid (vx\bar{x})(\mathcal{E} \parallel Q) \mid (vx\bar{x})(Q \parallel \mathcal{E})$$

$$\begin{array}{lll}
x \leftrightarrow y & \equiv y \leftrightarrow x & \text{SC-LINKCOMM} \\
(vx\bar{x})(P \parallel Q) & \equiv (v\bar{x}x)(Q \parallel P) & \text{SC-CUTCOMM} \\
(vx\bar{x})((vy\bar{y})(P \parallel Q) \parallel R) & \equiv (vy\bar{y})((vx\bar{x})(P \parallel R) \parallel Q) & \text{SC-CUTASSOC} \\
& \text{where } x \notin Q \text{ and } y \notin R
\end{array}$$

Figure 2.2: Structural Congruence for Classical Processes

Plugging is defined by replacing the one hole with a process:

$$\begin{array}{l}
\Box \quad [P] \triangleq P \\
(vx\bar{x})(\mathcal{E} \parallel Q)[P] \triangleq (vx\bar{x})(\mathcal{E}[P] \parallel Q) \\
(vx\bar{x})(Q \parallel \mathcal{E})[P] \triangleq (vx\bar{x})(Q \parallel \mathcal{E}[P])
\end{array}$$

I write $\text{fn}(\mathcal{E})$ for the free endpoints in $\mathcal{E}$.

I write $\text{bn}(\mathcal{E})$ for the endpoints bound by $\mathcal{E}$.

I write $\mathcal{E} \vdash \Gamma' \rightarrow \Gamma$ to mean that the evaluation context $\mathcal{E}$ is well-typed under input typing context Γ' and output typing context Γ .

$$\frac{}{\Box \vdash \Gamma \rightarrow \Gamma} \quad \frac{\mathcal{E} \vdash \Gamma' \rightarrow \Gamma, x : A \quad Q \vdash \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(\mathcal{E} \parallel Q) \vdash \Gamma' \rightarrow \Gamma, \Delta} \quad \frac{Q \vdash \Gamma, x : A \quad \mathcal{E} \vdash \Delta' \rightarrow \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(Q \parallel \mathcal{E}) \vdash \Delta' \rightarrow \Gamma, \Delta}$$

Processes are considered equivalent up to structural congruence.

Definition 2.4 (Structural Congruence). Structural congruence, written $P \equiv Q$, is the congruence closure over processes which satisfies the rules in Figure 2.2.

The semantics of CP processes is given by *reduction*. Reduction is closed over evaluation contexts, and structural congruence is embedded in reduction by allowing pre- and post-composition using E-CONG, written as $\equiv \rightarrow$, $\rightarrow \equiv$, or $\equiv \rightarrow \equiv$.

Definition 2.5 (Reduction). Reduction, written $P \rightarrow Q$, is the smallest relation on processes defined by the rules in Figure 2.3.

In the remainder of the section, I discuss each process construct together with its typing rule and operational semantics, either by itself—e.g. *link*—or together with its dual—e.g. *send* and *receive*.

2.1.1 Process Structure

The process $(vx\bar{x})(P \parallel Q)$ denotes a *cut*, which creates a communication channel with two endpoints, x and $\bar{x}$. For communication safety, the two endpoints must have dual types. Hence, $x : A$ and $\bar{x} : \bar{A}$. For deadlock freedom, the two endpoints must be used in different processes, and those processes cannot share any other channels. Hence, x is bound in

$$\begin{array}{lll}
(\nu x \bar{x})(x \leftrightarrow w \parallel P) & \longrightarrow P\{w/\bar{x}\} & \text{E-LINK} \\
(\nu x \bar{x})(x[y]. (P \parallel Q) \parallel \bar{x}(\bar{y}). R) & \longrightarrow (\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel R)) & \text{E-SEND} \\
(\nu x \bar{x})(x[\bar{\cdot}]. 0 \parallel \bar{x}(). Q) & \longrightarrow Q & \text{E-CLOSE} \\
(\nu x \bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) & \longrightarrow (\nu x \bar{x})(P \parallel Q) & \text{E-SELECT}_1 \\
(\nu x \bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) & \longrightarrow (\nu x \bar{x})(P \parallel R) & \text{E-SELECT}_2
\end{array}$$

$$\begin{array}{c}
\text{E-EQUIV} \\
\frac{P \equiv P' \quad P' \longrightarrow Q' \quad Q' \equiv Q}{P \longrightarrow Q}
\end{array}
\qquad
\begin{array}{c}
\text{E-CONG} \\
\frac{P \longrightarrow P'}{\mathcal{E}[P] \longrightarrow \mathcal{E}[P']}
\end{array}$$

Figure 2.3: Reduction for Classical Processes

P and $\bar{x}$ is bound in Q , and the typing environment is split between P and Q .

$$\frac{P \vdash \Gamma, x : A \quad Q \vdash \Delta, \bar{x} : \bar{A}}{(\nu x \bar{x})(P \parallel Q) \vdash \Gamma, \Delta} \text{T-CUT}$$

The cut combines two process constructs from the π -calculus: name restriction $(\nu x \bar{x})$ and parallel composition $P \parallel Q$. The two do not detach. In CP, $(\nu x \bar{x})P$ and $P \parallel Q$ are not syntactically well-formed processes.

Cuts are *commutative* and *quasi-associative*:

$$\begin{array}{lll}
(\nu x \bar{x})(P \parallel Q) & \equiv & (\nu \bar{x} x)(Q \parallel P) & \text{SC-CUTCOMM} \\
(\nu x \bar{x})((\nu y \bar{y})(P \parallel Q) \parallel R) & \equiv & (\nu y \bar{y})((\nu x \bar{x})(P \parallel R) \parallel Q) & \text{SC-CUTASSOC} \\
& & \text{where } x \notin Q \text{ and } y \notin R
\end{array}$$

The SC-CUTASSOC rule is not quite an associativity rule. Hence, *quasi-associative*. Wadler [2012] refers to the rule as “(Assoc)” in reference to the change of bracketing. In fact, the rule combines the associativity rule of parallel composition with scope extrusion, which commutes parallel composition and name restriction, as we will see in § 3.2.7 (see Proposition 3.76).

There is no specific reduction rule for cut. Rather, all other reduction rules resolve dual communication actions *under a cut*, i.e. over an existing channel. E-CONG allows reduction under cuts.

$$\begin{array}{c}
\text{E-CONG} \\
\frac{P \longrightarrow P'}{\mathcal{E}[P] \longrightarrow \mathcal{E}[P']}
\end{array}$$

2.1.2 Link

The process $x \leftrightarrow y$ denotes a *link*. It forwards any messages received on x to y , and vice versa. For communication safety, the two endpoints must have dual types. Hence, $x : A$ and $y : \bar{A}$.

$$\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \text{T-LINK}$$

Links are *commutative*. If two channels are connected by a link, the order in which they are connected is irrelevant. This property is captured by the following equivalence:

$$x \leftrightarrow y \equiv y \leftrightarrow x \quad \text{SC-LINKCOMM}$$

CP's semantics for link does not explicitly forward messages, but treats links as suspended α -renaming. When a link $x \leftrightarrow y$ reduces, it renames all occurrences of the dual of x to y , or all occurrences of the dual of y to x . In essence, this updates all the processes connected to the one side of the link to point directly at the other side, circumventing the link.

$$(\nu x \bar{x})(x \leftrightarrow y \parallel P) \longrightarrow P\{y/\bar{x}\} \quad \text{E-LINK}$$

The renaming targets a bound name. Hence, there cannot be any other occurrences of that name, and the link can be removed. The rule E-LINK gives link an asynchronous semantics, as P is not required to be ready on $\bar{x}$, whereas all other actions are synchronous. (I discuss synchronous semantics for link in § 3.3.2.)

2.1.3 Send and Receive

The send and receive actions are dual:

- The process $x[y].(P \parallel Q)$ denotes a *send* action. It creates a fresh channel, names one endpoint of that channel y , sends the other endpoint over x , then continues as P and Q in parallel, where y is bound in P and x is bound in Q .
- The process $x(y).P$ denotes a *receive* action. It receives an endpoint over x , names it y , then continues as P .

The typing rules for send and receive are as follows:

$$\frac{P \vdash \Gamma, y : A \quad Q \vdash \Delta, x : B}{x[y].(P \parallel Q) \vdash \Gamma, \Delta, x : A \otimes B} \text{T-SEND} \quad \frac{P \vdash \Gamma, y : A, x : B}{x(y).P \vdash \Gamma, x : A \wp B} \text{T-RECV}$$

The behaviour of send and receive is given by the following rule:

$$(\nu x \bar{x})(x[y].(P \parallel Q) \parallel \bar{x}(\bar{y}).R) \longrightarrow (\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel R)) \quad \text{E-SEND}$$

After a send action, the endpoints y and x are passed to different parallel processes, which ensures that they are handled independently. After a receive action, the two endpoints are kept by the same process. This may seem restrictive—and it is—but it is paramount to type preservation and deadlock freedom.

The send action is a *bound* send. It cannot send just any old endpoint that the process happens to have laying around. It only sends *fresh* endpoints. As part of the reduction, the send action *creates* a fresh channel with two

fresh endpoints. It sends one of those endpoints, and keeps the other. You can see this in E-SEND. The left-hand side has one name restriction, $(\nu x\bar{x})$, but the right-hand side has two, creating a fresh channel with $(\nu y\bar{y})$.

You can think of the bound send action as a suspended cut. For deadlock freedom, the two processes connected by a cut cannot share any other channel. (The relation between $\otimes$ and cut is widely known in logic, e.g. Girard [1987] treats cut as a special case of $\otimes$.)

I could extend CP with an *unbound* send action. Lindley and Morris [2015] discuss such an extension [section 3.1, under “A Simpler Send”]. They propose the following syntax, typing, and reduction rules:

$$P, Q, R ::= \dots \mid x\langle y \rangle. P$$

T-USEND

$$\frac{P \vdash \Gamma, y : A, x : B}{x\langle y \rangle. P \vdash \Gamma, x : A \otimes B, y : \bar{A}}$$

E-USEND

$$(\nu x\bar{x})(x\langle y \rangle. P \parallel \bar{x}(z). Q) \longrightarrow (\nu x\bar{x})(P \parallel Q\{y/z\})$$

The unbound send $x\langle y \rangle. P$ sends the *free* endpoint y over x , then continues as P . In the π -calculus, unbound send is strictly more expressive than bound send [see Boreale, 1996, Sangiorgi, 1996]. In CP, unbound send can be defined in terms of bound send, but the converse does not hold. Unbound send can be defined as follows:

$$x\langle y \rangle. P \triangleq x[z]. (z \leftrightarrow y \parallel P)$$

The definition has the same reduction behavior as E-USEND, albeit with an extra step for the link reduction.

$$\begin{aligned} (\nu x\bar{x})(x\langle y \rangle. P \parallel \bar{x}(z). Q) &\longrightarrow (\nu x\bar{x})(P \parallel Q\{y/z\}) \\ (\nu x\bar{x})(x\langle y \rangle. P \parallel \bar{x}(z). Q) &\triangleq (\nu x\bar{x})(x[w]. (w \leftrightarrow y \parallel P) \parallel \bar{x}(z). Q) \\ &\longrightarrow (\nu x\bar{x})(P \parallel (\nu w\bar{w})(w \leftrightarrow y \parallel Q\{\bar{w}/z\})) \\ &\longrightarrow (\nu x\bar{x})(P \parallel Q\{y/z\}) \end{aligned}$$

The converse does not hold. One might try and define bound send in terms of unbound send as follows:

$$x[y]. (P \parallel Q) \triangleq (\nu y\bar{y})(P \parallel x\langle \bar{y} \rangle. Q)$$

However, the two do not behave quite the same. In the former, reductions in both P and Q are blocked on the send action. In the latter, only reductions in Q are blocked on the send action, while reductions in P can happen before the send action.

In Wadler’s CP, the adoption of unbound send breaks the proof of progress, as it introduces cuts that cannot easily be eliminated by means of a commuting conversion. The cut in $(\nu x\bar{x})(a\langle x \rangle. P \parallel x(). Q)$ cannot be eliminated by a commuting conversion, as the unbound send action $a\langle x \rangle$

cannot commute past it, which breaks the simple left-to-right evaluation strategy proposed by Wadler [2012]. In CP, as presented in this chapter, the adoption of unbound send would not pose any problem, as my canonical form accounts for and justifies ineliminable top-level cuts.

By its very nature, bound send can only send endpoints, which means it does not easily generalise to sending other data. Hence, unbound send is important in the context of GV, where programs can send arbitrary data.

In the interest of an exact correspondence with CLL, and a close correspondence to Wadler’s CP, the version of CP presented in this chapter uses bound send.

2.1.4 Close and Wait

The close and wait actions are dual:

- The process $x[] . 0$ denotes a *close* action.
It sends a ping over x , then terminates.
- The process $x(). P$ denotes a *wait* action.
It receives a ping over x , then continues as P .

(I say ‘ping’ to imply the interaction between close and wait is merely a synchronisation and does not transmit any information.)

The typing rules for close and wait are as follows:

$$\frac{}{x[] . 0 \vdash x : \mathbf{1}} \text{ T-CLOSE} \qquad \frac{P \vdash \Gamma}{x(). P \vdash \Gamma, x : \perp} \text{ T-WAIT}$$

The behaviour of these two actions is given by the following rule:

$$(vx\bar{x})(x[] . 0 \parallel \bar{x}(). Q) \longrightarrow Q \text{ E-CLOSE}$$

After a close action, the process must terminate, but after a wait action, the process continues. As with send and receive, this seems restrictive—*because it is restrictive*—but is paramount to type preservation.

The close and wait actions close the channel, which removes the cut. Any two processes connected by a cut cannot share any other channel. If both processes were allowed to continue, that would leave them unconnected. This would break an important property of CP—that all processes are connected—and relaxing it would be logically equivalent to admitting the Mix rule (see § 3.3.7) and would break the exact correspondence with CLL. (Such a variant of CP is explored by Atkey et al. [2016].)

2.1.5 Select and Offer

The select and offer actions are dual:

- The process $x \triangleleft \text{inl}. P$ denotes a *left selection* action. It sends the label inl over x , then continues as P .
- The process $x \triangleleft \text{inr}. P$ denotes a *right selection* action. It sends the label inr over x , then continues as P .
- The process $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ denotes a *choice* action. It receives a label over x , and then continues as either P or Q , depending on which label was received.

The typing rules for select and offer are as follows:

$$\frac{P \vdash \Gamma, x : A}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_1 \quad \frac{P \vdash \Gamma, x : B}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_2$$

$$\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B}{x \triangleright \{\text{inl}: P; \text{inr}: Q\} \vdash \Gamma, x : A \& B} \text{T-OFFER}$$

The behaviour of these actions is given by the following rules:

$$\begin{aligned} (\nu x \bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) &\rightarrow (\nu x \bar{x})(P \parallel Q) \quad \text{E-SELECT}_1 \\ (\nu x \bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) &\rightarrow (\nu x \bar{x})(P \parallel R) \quad \text{E-SELECT}_2 \end{aligned}$$

The syntax for the select and offer actions was adapted from Dardha and Gay [2018], rather than from Wadler [2012]. It is easily generalized from binary choice to variant types by removing the restriction that labels must always be drawn from $\{\text{inl}, \text{inr}\}$. (I discuss variant types in § 3.3.3.)

2.1.6 The Absurd Offer

The process $x \not\triangleright N$ denotes the *absurd offer*. It waits to receive a choice between *zero* alternatives. Such a choice cannot be made, which means that there is no corresponding select action, and no corresponding reduction rule. In essence, an absurd offer is inert. The absurd offer is the sole process that is allowed to leave endpoints unused, and the set of those unused endpoints is denoted by N .

$$\frac{N = \text{fn}(\Gamma)}{x \not\triangleright N \vdash \Gamma, x : \top} \text{T-ABSURD}$$

The ‘inert semantics’ for absurd is unsatisfying—something has clearly gone wrong, and instead of cleaning up, we leave the garbage strewn around. However, it is not *incorrect*. By accident, it satisfies all the expected properties of normalisation for CP. (The accident is the fact that CP processes are fully connected. Hence, even if $x \triangleright \{\}$ did not have immediate license to kill some process P , it would have obtained it eventually.) In variants where processes are not always fully connected, the inert semantics break progress (see, e.g. Chapter 3). (I present an *exceptional* semantics for the absurd offer in § 3.3.1².)

²I use ‘exceptional’ to mean ‘reminiscent of exceptions’, but I also happen to believe they are very good semantics.

The syntax for the absurd offer deviates from the naming scheme for the binary select and offer actions. If I were to follow the naming scheme, I would write ‘ $x \triangleright \{\}$ ’. However, it is important for the absurd offer to record which endpoints it discards. For one, recording those names allows us to define linearity by counting names (as in § 2.2.1). More importantly, to give an exceptional semantics to the absurd offer, those names are needed to be able to erase types. For an intuition, let us look at the absurd offer in Wadler’s CP. It does not record the discarded endpoints and it has the following reduction rule:

$$\frac{a \triangleright \{\} \vdash \Gamma, a : T, x : A \quad P \vdash \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(a \triangleright \{\} \parallel P) \vdash \Gamma, \Delta, a : T} \longrightarrow \frac{}{a \triangleright \{\} \vdash \Gamma, \Delta, a : T}$$

If the absurd offer discarded an endpoint x , it may kill the process P that owns the dual endpoint $\bar{x}$. The reduction rule appears innocuous, but a problem surfaces when we erase types and consider the equivalent reduction rule on *processes*:

$$(vx\bar{x})(a \triangleright \{\} \parallel P) \longrightarrow a \triangleright \{\}$$

It is no longer possible to see whether P owns any of the endpoints dual to those discarded, which means that $a \triangleright \{\}$ has license to kill *any process*³. Surprisingly, by itself this does not pose a *huge* problem for CP. Ultimately, all processes in CP are connected, so even if $a \triangleright \{\}$ were to kill a process it was not immediately connected to, it would have gotten there eventually. However, if the language were extended with parallel but unconnected processes—as in, e.g. HCP—this leads to immediate problems, where some $a \triangleright \{\}$ kills a process it was in no way connected to. It also leads to issues with scope extrusion in SC-CUTASSOC:

$$(vx\bar{x})((vy\bar{y})(P \parallel a \triangleright \{\}) \parallel b \triangleright \{\}) \equiv (vy\bar{y})((vx\bar{x})(P \parallel b \triangleright \{\}) \parallel a \triangleright \{\})$$

If $a \triangleright \{\}$ is responsible for killing x on the left-hand side of the structural congruence, then on the right-hand side it moves out of the scope of $(vx\bar{x})$.

My syntax, $a \not\triangleright N$, is an effort to avoid clutter: I omit the absurd continuation and add a lightning bolt, which, to me, implies that something has gone wrong⁴.

2.2 Metatheory

In this section, I introduce the metatheory for Classical Processes. This section proceeds as follows:

³Wadler was unaware of the issue with the absurd offer and type erasure.

⁴The lightning bolt is an homage to Simon Fowler’s Exceptional GV, where it marks “zapper threads” [Fowler, 2019, Figure 9.4], and to a poster titled “zap is why we can’t have nice things” that I presented at my CDT’s 2016 Industrial Engagement Event.

- In § 2.2.1, I give several preliminary definitions that are used throughout the discussion of the metatheory.
- In § 2.2.2, I prove preservation (Proposition 2.27). In essence, the proof is a reproduction of Wadler’s proof of preservation, but without the commuting conversions and with small changes to account for the my changes to CP’s syntax.
- In § 2.2.3, I define canonical form (Definition 2.30) and prove progress (Proposition 2.32). The proof cannot be reproduced from Wadler [2012], as the reduction rules I give in § 2.1 are a proper subset of those given by Wadler [2012]. I could adapt the proof by Lindley and Morris [2015], but it relies on the property that CP processes can be rewritten to right-branching form (see Proposition 2.51), which does not hold for the variants of CP that I discuss in later chapters. The proof for progress I give in this section generalises Lindley and Morris’ appeal to right-branching form using process contexts, which allows me to reuse the same proof structure in later chapters. This proof of progress first appeared in my M.Sc. thesis [Kokke, 2017].
- In § 2.2.4, I define dependency graphs for CP processes (Definition 2.35). I prove that CP is deadlock-free, as its dependency graphs are always acyclic (Proposition 2.39), and I prove that my definition of canonical form is adequate (Corollary 2.47) and stronger than Wadler’s definition.
- In § 2.2.5, I define connection graphs for CP processes (Definition 2.49) and prove that CP’s connection graphs are always trees (Proposition 2.50). The validity of right-branching form for CP follows as a corollary.
- In § 2.2.6, I sketch the basis of a behavioural theory for CP.

2.2.1 Preliminaries

Configuration Contexts

Processes in CP can be rewritten to right-branching form:

$$(\nu x_1 \bar{x}_1)(P_1 \parallel \dots (\nu x_n \bar{x}_n)(P_n \parallel P_{n+1}) \dots)$$

Right-branching form is a convenient form to work with. It neatly captures the prefix of top-level cuts and the connected processes, and it relates the endpoints and processes in a predictable manner—each x_i is free in P_i , and each $\bar{x}_i$ is free in some P_j where $j > i$.

Right-branching form has its downsides. Proving that processes can be rewritten to right-branching form is far from trivial (Proposition 2.51). More importantly, right-branching form is a quirk of CP’s rigid process

structure. In variants of CP with more expressive process structure, processes cannot be rewritten to right-branching form.

I generalise right-branching form using *configuration contexts*. Configuration contexts are multi-hole process contexts that consist only of cuts. For instance, for the process above, an example configuration context is

$$(\nu x_1 \bar{x}_1)(\square_1 \parallel \dots (\nu x_n \bar{x}_n)(\square_n \parallel \square_{n+1}) \dots)$$

The expressions $\square_1, \dots, \square_{n+1}$ are the $n+1$ holes, numbered from left to right for convenience. The name $\mathcal{C}^k$ ranges over configuration contexts, where k denotes the number of holes. Plugging, written $\mathcal{C}^k[P_1, \dots, P_k]$, replaces the holes, from left to right, with the processes $P_1, \dots, P_k$, such that each $\square_i$ is replaced with the process P_i (for $1 \leq i \leq k$). If $\mathcal{C}^{n+1}$ refers to the example configuration context shown above and $P_1, \dots, P_{n+1}$ to the processes of the example process shown above, then $\mathcal{C}^{n+1}[P_1, \dots, P_{n+1}]$ is exactly equal to the example process. The use of right-branching form requires that the prefix of cuts is a right-branching tree, i.e. a list. Configuration contexts permit arbitrary trees.

Definition 2.6 (Configuration Context). Configuration contexts are n -hole process contexts, as defined by the following grammar:

$$\mathcal{C}, \mathcal{D} ::= \square \mid (\nu x \bar{x})(\mathcal{C} \parallel \mathcal{D})$$

If there is risk of ambiguity, we explicitly write the number of holes in a configuration context with a superscript, e.g. as $\mathcal{C}^n$.

Plugging is defined by replacing the n holes with n processes, left to right:

$$\begin{aligned} \square & [P] \triangleq P \\ (\nu x \bar{x})(\mathcal{C}^n \parallel \mathcal{D}^k) [P_1, \dots, P_n, P_{n+1}, \dots, P_{n+k}] & \triangleq (\nu x \bar{x})(\mathcal{C}^n[P_1, \dots, P_n] \parallel \mathcal{D}^k[P_{n+1}, \dots, P_{n+k}]) \end{aligned}$$

I write $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n]$ for the evaluation context focused on the i 'th hole in $\mathcal{C}$ such that $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n][P_i] = \mathcal{C}[P_1, \dots, P_i, \dots, P_n]$.

I write $\text{dn}(\mathcal{C})$ for the unordered pairs of dual endpoints bound by $\mathcal{C}$.

$$\begin{aligned} \text{dn}(\square) & \triangleq \emptyset \\ \text{dn}((\nu x \bar{x})(\mathcal{C} \parallel \mathcal{D})) & \triangleq \{\{x, \bar{x}\}\} \cup \text{dn}(\mathcal{C}) \cup \text{dn}(\mathcal{D}) \end{aligned}$$

I write $\text{bn}(\mathcal{C})$ for the endpoints bound by $\mathcal{C}$, i.e. $\text{bn}(\mathcal{C}) \triangleq \bigcup \text{dn}(\mathcal{C})$.

I write $\mathcal{C}^n \vdash \Gamma_1 \mid \dots \mid \Gamma_n \rightarrow \Gamma$ to mean that the configuration context $\mathcal{C}^n$ is well-typed under input typing contexts $\Gamma_1, \dots, \Gamma_n$ and output typing context Γ .

$$\frac{\square \vdash \Gamma \rightarrow \Gamma}{\mathcal{C}^n \vdash \Gamma_1 \mid \dots \mid \Gamma_n \rightarrow \Gamma, x : A \quad \mathcal{D}^k \vdash \Delta_1 \mid \dots \mid \Delta_k \rightarrow \Delta, \bar{x} : \bar{A}} \quad (\nu x \bar{x})(\mathcal{C}^n \parallel \mathcal{D}^k) \vdash \Gamma_1 \mid \dots \mid \Gamma_n \mid \Delta_1 \mid \dots \mid \Delta_k \rightarrow \Gamma, \Delta$$

Shallow Structural Congruence

Reduction is closed under evaluation contexts, not under arbitrary process contexts, and only acts on the topmost actions. As such, reduction only needs a structural congruence that is similarly closed under evaluation contexts—a *shallow* structural congruence. Therefore, it is useful to distinguish different kinds of structural congruence, based on which portions of the process they are permitted to rewrite.

Definition 2.7 (Shallow Structural Congruence). Shallow structural congruence, written $\overset{s}{\equiv}$, is the smallest symmetric relation over processes that satisfies the rules in Figure 2.2 (p. 34) and is closed under evaluation contexts, as per the following rule:

$$\frac{\text{SC-CONG} \quad P \overset{s}{\equiv} P'}{\mathcal{E}[P] \overset{s}{\equiv} \mathcal{E}[P']}$$

Most rules of the structural congruence target name restriction and parallel composition. The odd one out is SC-LINKCOMM, which rewrites a link action. It will be useful to single out the portions of a structural congruence that rewrite links.

Definition 2.8 (Link-Preserving Structural Congruence). Link-preserving structural congruence, written $\overset{l}{\equiv}$, is the congruence closure over processes that satisfies the rules in Figure 2.2 except for SC-LINKCOMM.

Finally, it will be useful to have variants which combine these restrictions. In practice, I only need link-preserving shallow structural congruence and deep structural congruence.

Definition 2.9 (Link-Preserving Shallow Structural Congruence). Link-preserving shallow structural congruence, written $\overset{ls}{\equiv}$, is the intersection of link-preserving and shallow structural congruence.

Definition 2.10 (Deep Structural Congruence). Deep structural congruence, written $\overset{d}{\equiv}$, is the equivalence closure of the complement of $\overset{ls}{\equiv}$ with respect to $\equiv$.

Any structural congruence can be decomposed into its link-preserving shallow and deep structural components.

Lemma 2.11. If $P \equiv Q$, then there exists some R such that $P \overset{ls}{\equiv} R$ and $R \overset{d}{\equiv} Q$.

Ready Processes and Threads

A process is ready if it is ready to perform some communication action, i.e. if it is a link or it is prefixed by an action. The formal definition of *ready* processes is a bit verbose, since actions are not a syntactic sort in

CP, so “prefixed by an action” is not well-defined. To work around this, I enumerate the process constructors that contain an action.

Definition 2.12 (Ready). A process P is ready to act on x , written $\text{ready}(P, x)$, if it is of one of the forms:

$$\begin{array}{llll} x \leftrightarrow y & x[y].(P \parallel Q) & x[].0 & x \triangleleft \text{inl}.P \quad x \triangleright \{\text{inl}.P; \text{inr}.Q\} \\ y \leftrightarrow x & x(y).P & x().P & x \triangleleft \text{inr}.P \quad x \not\triangleleft N \end{array}$$

A process is ready if it is ready to act on some endpoint.

In particular, links $x \leftrightarrow y$ are considered ready to act on both x and y , and absurd $x \not\triangleleft N$ is *not* considered ready to act on the channels $y \in N$.

A process can be decomposed into a prefix of its cuts, and a series of threads connected by those cuts. Such a prefix is the *maximum configuration context*, in the sense that no further cuts can be added.

Definition 2.13 (Maximum Configuration Context). The maximum configuration context $\mathcal{C}^n$ of a process P is the configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 1$) and (for $1 \leq i \leq n$) each P_i is ready. The processes P_i are the threads of P . Every process has a unique maximum configuration context.

Likewise, evaluation contexts are maximal if no further cuts can be added. Informally, maximal evaluation contexts are paths to the threads contained within some process, so each maximum configuration context $\mathcal{C}^n$ gives us n distinct maximal evaluation contexts.

Definition 2.14 (Maximal Evaluation Context). A maximal evaluation context $\mathcal{E}$ of a process P is an evaluation context such that $P = \mathcal{E}[Q]$ and Q is ready. If $\mathcal{C}$ is the maximum configuration context of P , then $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n]$ is a maximal evaluation context of P .

Finally, I refer to top-level ready processes as *threads*. A significant portion of CP’s metatheory deals with threads. Let the metavariable T range over threads.

Definition 2.15 (Thread). A process P is a thread of Q if there exists some evaluation context $\mathcal{E}$ of Q such that $Q = \mathcal{E}[P]$ and P is ready. I say that Q contains the thread P to mean that P is a thread of Q . I say P is a thread when the process Q of which P is a thread can be inferred from context. Let T range over threads.

The use of the thread metavariable allows us to succinctly decompose any process P into its maximum configuration context and its threads, by stating “ P is of the form $\mathcal{C}[T_1, \dots, T_n]$ ”, as the notation implies that each thread T_i is a ready process and therefore that $\mathcal{C}$ is the maximum configuration context.

In CP, every thread is ready. Therefore, it is tempting to do away with the separate definition of threads and let T range over ready processes.

However, the two definitions diverge for GV, since GV's threads have internal reduction behaviour and therefore not all threads are ready. To keep my terminology consistent, I keep the definitions of ready processes and threads separate.

Process Contexts

For the occasional convenience, I define full n -hole process contexts, which are arbitrary processes with any number of holes. Process contexts may contain any process construct, and may contain holes in any position, including nested under actions. As such, process contexts generalise evaluation and configuration contexts.

Definition 2.16 (Process Context). *Process contexts are defined by the grammar for processes, extended with the hole constructor, written $\square$. A process context may have any number of holes.*

The names $P[\cdot]$, $Q[\cdot]$, and $R[\cdot]$ range over process contexts, where the trailing $[\cdot]$ is intended to help distinguish between processes and process contexts, and denotes the position of the arguments for plugging. I write $P^n[\cdot]$ to denote that the process context $P[\cdot]$ has n holes.

Plugging, written $P[P_1, \dots, P_n]$, is defined by replacing the n holes in the process context $P[\cdot]$ with the processes $P_1, \dots, P_n$ in order from left to right.

Linearity

CP has a *linear* type system. It ensures that resources are always used exactly once, and never copied or dropped. Given the correspondence between CP and linear logic, that seems almost painfully obvious. However, due to CP's reuse of endpoint names, it may appear that resources are used multiple times. (For instance, the process $x().x[] . 0$ appears to use the endpoint x twice.)

The true measure of a linear calculus lies in its execution. Any execution should use all linear resources exactly once. Unfortunately, such a formulation of linearity is a bit tedious.

For my purposes, it suffices to define linearity by counting the uses of free endpoint names—taking the sum across any parallel composition, and the union across an offer. This formulation reveals the implicit rebinding of names in CP's actions. For instance, the action $x(y).P$ uses the name x , but then binds a *fresh* x for the remainder of the session.

Definition 2.17 (Free Name Count). *The multiset of free endpoints in P , written $\text{fn}(P)$, is a multiset (see Definition A.3) with support set $\text{fn}(P)$ and*

multiplicity function $\mu_{\text{fn}(P)}$.

$$\begin{array}{ll}
\text{fn}(x \leftrightarrow y) & \triangleq \{x, y\} \\
\text{fn}((\nu x \bar{x})(P \parallel Q)) & \triangleq (\text{fn}(P) \setminus \{x\}) + (\text{fn}(Q) \setminus \{\bar{x}\}) \\
\text{fn}(x[y].(P \parallel Q)) & \triangleq \{x\} + (\text{fn}(P) \setminus \{y\}) + (\text{fn}(Q) \setminus \{x\}) \\
\text{fn}(x[] . 0) & \triangleq \{x\} \\
\text{fn}(x(y). P) & \triangleq \{x\} + (\text{fn}(P) \setminus \{x, y\}) \\
\text{fn}(x(). P) & \triangleq \{x\} + \text{fn}(P) \\
\text{fn}(x \triangleleft \text{inl}. P) & \triangleq \{x\} + (\text{fn}(P) \setminus \{x\}) \\
\text{fn}(x \triangleleft \text{inr}. P) & \triangleq \{x\} + (\text{fn}(P) \setminus \{x\}) \\
\text{fn}(x \triangleright \{\text{inl}: P; \text{inr}: Q\}) & \triangleq \{x\} + ((\text{fn}(P) \setminus \{x\}) \cup (\text{fn}(Q) \setminus \{x\})) \\
\text{fn}(x \not\leq N) & \triangleq \{x\} + \{w \mid w \in N\}
\end{array}$$

Note that the operation $\mathcal{X} \setminus X$ removes all occurrences of the elements in the set X from the multiset $\mathcal{X}$ (see Definition A.3).

The only unusual case is the case for the absurd offer $x \not\leq N$, which counts all names in N as used. This seems wrong, since the absurd offer never uses the names in N . However, it is correct in the tedious sense. All executions of $x \not\leq N$ use all the names in N , trivially, since $x \not\leq N$ is never executed.

For actions, fn removes all previous occurrences of the relevant name, which hides any potential non-linear usage of that name, e.g. $\text{fn}(x(). P) \triangleq \{x\} + (\text{fn}(P) \setminus \{x\})$ hides the count of x in P . Hence, fn is *shallow*, in the sense that it only accurately counts the top-most parallel usages of the name. This suffices for my purposes.

Linearity states that, for well-typed processes, each endpoint in the typing environment is used exactly once in the process, and vice versa.

Proposition 2.18 (Linearity). *If $P \vdash \Gamma$, then:*

- $x \in^k \text{fn}(P) \implies k = 1$
- $x \in^1 \text{fn}(P) \iff x \in \text{fn}(\Gamma)$

Proof. By induction on the derivation of $P \vdash \Gamma$.

□

For any well-typed process $\mathcal{E}[P]$, each endpoint bound by $\mathcal{E}$ is used exactly once in the process P , and vice versa.

Corollary 2.19. *If $\mathcal{E}[P] \vdash \Gamma$, then:*

- $x \in^k \text{fn}(P) \implies k = 1$
- $x \in^1 \text{fn}(P) \iff x \in \text{bn}(\mathcal{E})$
- $\text{bn}(\mathcal{E}) \subseteq \text{fn}(P)$

For any well-typed configuration $\mathcal{E}^n[P_1, \dots, P_n]$, the processes $P_1, \dots, P_n$ must collectively use all the endpoints bound by $\mathcal{E}^n$ exactly once.

Corollary 2.20. *If $\mathcal{C}^n[P_1, \dots, P_n] \vdash \Gamma$, then:*

- $x \in^k \bigcup_{1 \leq i \leq n} \underline{\text{fn}}(P_i) \implies k = 1$
- $x \in^1 \bigcup_{1 \leq i \leq n} \underline{\text{fn}}(P_i) \iff x \in \text{bn}(\mathcal{C})$
- $\text{bn}(\mathcal{C}) \subseteq \bigcup_{1 \leq i \leq n} \underline{\text{fn}}(P_i)$

Separation

Separation relates configuration contexts and evaluation contexts—it ‘zooms in’, from viewing a process as a series of connected processes, to viewing two specific processes and the cut connecting them. Separation also captures an essential property of CP’s type system: dual endpoints must be in distinct processes, separated by a cut.

Lemma 2.21 (Separation). *If $P \vdash \Gamma$, $\mathcal{C}^n$ is a configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 2$), and there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$ such that $x \in \text{fn}(P_i)$ and $\bar{x} \in \text{fn}(P_j)$ (for some $1 \leq i, j \leq n$), there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that either*

1. $P = \mathcal{E}[(vx\bar{x})(\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j])]$, or
2. $P = \mathcal{E}[(v\bar{x}x)(\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i])]$.

Proof. By induction on the structure of the configuration context $\mathcal{C}$.

There are two cases:

- P_i and P_j are on different sides of the cut.

There are two cases, corresponding to the two equations:

- P_i is to the left and P_j is to the right.

$$P = (vx\bar{x})(\mathcal{C}_i[P_1, \dots, P_i, \dots, P_k] \parallel \mathcal{C}_j[P_{k+1}, \dots, P_j, \dots, P_n]).$$

By $P \vdash \Gamma$ and T-CUT, the bound endpoints must be $\{x, \bar{x}\}$.

$$\begin{aligned} \mathcal{C} &= \square \\ \text{Let } \mathcal{F}_i &= \mathcal{C}_i[P_1, \dots, \square_i, \dots, P_k] \\ \mathcal{F}_j &= \mathcal{C}_j[P_{k+1}, \dots, \square_j, \dots, P_n] \end{aligned}$$

The result follows.

- P_i is to the right and P_j is to the left.

As above.

- P_i and P_j are on the same side of the cut.

There are two cases:

- P_i and P_j are both to the left.

$$P = (vz\bar{z})(\mathcal{D}[P_1, \dots, P_k] \parallel R) \text{ and } 1 \leq i, j \leq k.$$

By $P \vdash \Gamma$ and T-CUT, the bound endpoints $\{z, \bar{z}\}$ cannot be $\{x, \bar{x}\}$.

By induction, one of the following holds:

- * $\mathcal{D}[P_1, \dots, P_k] = \mathcal{E}'[(\nu x \bar{x})(\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j])]$; or
- * $\mathcal{D}[P_1, \dots, P_k] = \mathcal{E}'[(\nu \bar{x} x)(\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i])]$.

Let $\mathcal{E} = (\nu z \bar{z})(\mathcal{E}' \parallel R)$.

The result follows.

- P_i and P_j are both to the right.

As above.

□

The separation lemma is rather precise, and gives us one of two equalities. However, both cases are equivalent up to structural congruence.

Corollary 2.22 (Separation). *If $P \vdash \Gamma$, $\mathcal{C}^n$ is a configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 2$), and there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$ such that $x \in \text{fn}(P_i)$ and $\bar{x} \in \text{fn}(P_j)$ (for some $1 \leq i, j \leq n$), there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that $P \equiv \mathcal{E}[(\nu x \bar{x})(\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j])]$.*

Proof. By Lemma 2.21 and SC-CUTCOMM.

□

Evaluation contexts commute with cuts.

Lemma 2.23. *If $\text{fn}(P) \cap \text{bn}(\mathcal{E}) = \emptyset$, then $\mathcal{E}[(\nu x \bar{x})(P \parallel Q)] \equiv (\nu x \bar{x})(P \parallel \mathcal{E}[Q])$.*

Proof. By induction on the structure of the evaluation context $\mathcal{E}$.

□

2.2.2 Preservation

Structural congruences preserve typing. If some process P is well-typed and is equivalent to some process Q under structural congruence, then Q is well-typed under the same typing environment.

Lemma 2.24. *If $P \equiv Q$, then $P \vdash \Gamma$ if and only if $Q \vdash \Gamma$.*

Proof. By induction on the derivation of the equivalence $P \equiv Q$.

The case for reflexivity follows immediately. The cases for symmetry and transitivity follow immediately by induction. The case for congruence closure follows by induction and the injectivity of the type derivation rules. The cases for applications of SC-LINKCOMM, SC-CUTCOMM, and SC-CUTASSOC are as follows, presented as equivalences on type derivations:

- Case SC-LINKCOMM:

$$\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \equiv \frac{\frac{y \leftrightarrow x \vdash y : \bar{A}, x : \bar{\bar{A}}}{y \leftrightarrow x \vdash y : \bar{A}, x : A}}{} \text{Lemma 2.1}$$

- Case SC-CUTCOMM:

$$\frac{\frac{P \vdash \Gamma_1, x : A \quad Q \vdash \Gamma_2, \bar{x} : \bar{A}}{(vx\bar{x})(P \parallel Q) \vdash \Gamma_1, \Gamma_2}}{\text{III}} \frac{P \vdash \Gamma_1, x : \bar{\bar{A}}}{P \vdash \Gamma_1, x : A} \text{Lemma 2.1}$$

$$\frac{Q \vdash \Gamma_2, \bar{x} : \bar{A} \quad \frac{P \vdash \Gamma_1, x : \bar{\bar{A}}}{P \vdash \Gamma_1, x : A}}{(v\bar{x}x)(Q \parallel P) \vdash \Gamma_1, \Gamma_2}$$

- Case SC-CUTASSOC:

$$\frac{\frac{P \vdash \Gamma_1, x : A, y : B \quad Q \vdash \Gamma_2, \bar{y} : \bar{B}}{(vy\bar{y})(P \parallel Q) \vdash \Gamma_1, \Gamma_2, x : A} \quad R \vdash \Gamma_3, \bar{x} : \bar{A}}{(vx\bar{x})((vy\bar{y})(P \parallel Q) \parallel R) \vdash \Gamma_1, \Gamma_2, \Gamma_3}$$

$$\frac{\frac{P \vdash \Gamma_1, x : A, y : B \quad R \vdash \Gamma_3, \bar{x} : \bar{A}}{(vx\bar{x})(P \parallel R) \vdash \Gamma_1, \Gamma_3, y : B} \quad Q \vdash \Gamma_2, \bar{y} : \bar{B}}{(vy\bar{y})((vx\bar{x})(P \parallel R) \parallel Q) \vdash \Gamma_1, \Gamma_2, \Gamma_3}$$

The above derivation for the left-hand side (symmetrically, right-hand side) is the only derivation, as $x \notin Q$ (symmetrically, $y \notin R$).

□

Renaming preserves typing. If a process is well-typed, then renaming any free endpoint does not affect its typing.

Lemma 2.25. *If $P \vdash \Gamma, x : A$, then $P\{w/x\} \vdash \Gamma, w : A$.*

Proof. The result follows by induction.

□

Plugging with any form of process context preserves typing.

Lemma 2.26. *If $P^n[\cdot] \vdash \Gamma_1 \mid \dots \mid \Gamma_n \rightarrow \Gamma$ and $Q_i \vdash \Gamma_i$ (for $1 \leq i \leq n$), then $P^n[Q_1, \dots, Q_n] \vdash \Gamma$.*

Proof. By induction on the derivation of $P^n[\cdot] \vdash \Gamma_1 \mid \dots \mid \Gamma_n \rightarrow \Gamma$.

□

Reduction preserves typing. If a process P is well-typed and reduces to some other process Q , then Q is well-typed under the same typing environment.

Proposition 2.27 (Preservation). *If $P \vdash \Gamma$ and $P \longrightarrow Q$, then $Q \vdash \Gamma$.*

Proof. By induction on the derivation of the reduction $P \longrightarrow Q$.

The cases for E-CONG follow by induction and Lemma 2.26. The case for E-EQUIV follows by induction and Lemma 2.24. The cases for the rules E-LINK, E-SEND, E-CLOSE, E-SELECT₁, and E-SELECT₂ are as follows, presented as reductions on type derivations:

- Case E-LINK:

$$\frac{x \leftrightarrow w \vdash x : A, w : \bar{A} \quad P \vdash \Gamma, \bar{x} : \bar{A}}{(vx\bar{x})(x \leftrightarrow w \parallel P) \vdash \Gamma, w : \bar{A}} \quad \downarrow$$

$$\frac{P \vdash \Gamma, \bar{x} : \bar{A}}{P\{w/\bar{x}\} \vdash \Gamma, w : \bar{A}} \text{ Lemma 2.25}$$

- Case E-SEND:

$$\frac{\frac{P \vdash \Gamma_1, y : A \quad Q \vdash \Gamma_2, x : B}{x[y].(P \parallel Q) \vdash \Gamma_1, \Gamma_2, x : A \otimes B} \quad \frac{R \vdash \Gamma_3, \bar{y} : \bar{A}, \bar{x} : \bar{B}}{\bar{x}(\bar{y}).R \vdash \Gamma_3, \bar{x} : \bar{A} \wp \bar{B}}}{(vx\bar{x})(x[y].(P \parallel Q) \parallel \bar{x}(\bar{y}).R) \vdash \Gamma_1, \Gamma_2, \Gamma_3} \quad \downarrow$$

$$\frac{P \vdash \Gamma_1, y : A \quad \frac{Q \vdash \Gamma_2, x : B \quad R \vdash \Gamma_3, \bar{y} : \bar{A}, \bar{x} : \bar{B}}{(vx\bar{x})(Q \parallel R) \vdash \Gamma_2, \Gamma_3, \bar{y} : \bar{A}}}{(vy\bar{y})(P \parallel (vx\bar{x})(Q \parallel R)) \vdash \Gamma_1, \Gamma_2, \Gamma_3}$$

- Case E-CLOSE:

$$\frac{\frac{x[] . 0 \vdash x : \mathbf{1}}{(vx\bar{x})(x[] . 0 \parallel \bar{x}().Q) \vdash \Gamma_2} \quad \frac{Q \vdash \Gamma_2}{\bar{x}().Q \vdash \Gamma_2, \bar{x} : \perp}}{\longrightarrow Q \vdash \Gamma_2}$$

- Case E-SELECT₁:

$$\frac{\frac{P \vdash \Gamma_1, x : A}{x \triangleleft \text{inl}. P \vdash \Gamma_1, x : A \oplus B} \quad \frac{Q \vdash \Gamma_2, \bar{x} : \bar{A} \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{\bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_2, \bar{x} : \bar{A} \& \bar{B}}}{(vx\bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) \vdash \Gamma_1, \Gamma_2} \quad \downarrow$$

$$\frac{P \vdash \Gamma_1, x : A \quad Q \vdash \Gamma_2, \bar{x} : \bar{A}}{(vx\bar{x})(P \parallel Q) \vdash \Gamma_1, \Gamma_2}$$

- Case E-SELECT₂:

$$\begin{array}{c}
 \frac{P \vdash \Gamma_1, x : B}{x \triangleleft \text{inr}. P \vdash \Gamma_1, x : A \oplus B} \quad \frac{Q \vdash \Gamma_2, \bar{x} : \bar{A} \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{\bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_2, \bar{x} : \bar{A} \& \bar{B}} \\
 \hline
 (vx\bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) \vdash \Gamma_1, \Gamma_2 \\
 \downarrow \\
 \frac{P \vdash \Gamma_1, x : B \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{(vx\bar{x})(P \parallel R) \vdash \Gamma_1, \Gamma_2}
 \end{array}$$

□

2.2.3 Progress

What should the canonical forms be for processes in CP? The usual starting point would be to ask what the canonical forms are for *closed* processes. For instance, in the λ -calculus, a closed process of the function type must be a lambda, and in the π -calculus, a closed process must be equivalent to the terminated process. Alas, that question is meaningless for CP, which *has no closed processes*. Any closed process would correspond to a proof of the empty sequent, which is not provable in CLL. A CP process is never *done*. There is no terminated process. There are only processes that are temporarily stuck, blocked on a free endpoint—one endpoint of a channel whose other endpoint has not yet been connected. Unfortunately, it is not sufficient to simply require that the process contains an action on a free endpoint. For instance, the process

$$(vx\bar{x})(a(). P \parallel (vy\bar{y})(y[]. 0 \parallel \bar{y}(). Q))$$

has an action that is blocked on an external channel, but it can still reduce. Hence, the definition of canonical form should also capture the fact that no further reduction rules can be applied. Unsurprisingly, that condition is in and of itself sufficient, so let us formalise it as a starting point.

The simplest and most useless definition would be to say that P is in canonical form if and only if $P \not\rightarrow$. However, it would be more satisfying if we could characterise the conditions under which no reduction rules apply.

First, note that E-LINK acts differently from most reduction rules. Most reduction rules, i.e. E-SEND, E-CLOSE, E-SELECT₁, and E-SELECT₂, are synchronous interactions between two processes and are blocked unless both processes are ready to act on dual endpoints. However, link is *asynchronous*. Any ready link can evaluate to a renaming in the process it communicates with, regardless of whether that process is ready to act on the relevant channel.

To capture this difference, I divide reduction into α -reduction, which captures link evaluating as α -renaming, and β -reduction, which captures all other reduction. In essence, α -reduction captures asynchronous reduction, and β -reduction captures synchronous reduction.

Definition 2.28 (α -Reduction). A process P α -reduces to Q , written $P \rightarrow_\alpha Q$, if there exists a reduction $P \rightarrow Q$ which only uses the rules *E-LINK*, *E-CONG*, and *E-EQUIV*.

Definition 2.29 (β -Reduction). A process P β -reduces to Q , written $P \rightarrow_\beta Q$, if there exists a reduction $P \rightarrow Q$ that does not use the rule *E-LINK*.

Any reduction is either an α -reduction or a β -reduction.

I could say that P is in canonical form if and only if $P \not\rightarrow_\alpha$ and $P \not\rightarrow_\beta$. However, I have already given a tighter characterisation above. A process P is in canonical form if and only if

1. P contains no link thread ready to act on a bound endpoint; and
2. P contains no two threads ready to act on dual endpoints.

If (1) does not hold, *E-LINK* applies, and if (2) does not hold, at least one of the other reduction rules applies... *and that is progress in a nutshell!*

Definition 2.30 (Canonical Form). A process P is in canonical form, written $\text{canonical}(P)$, if P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 1$) and (for $1 \leq i, j \leq n$)

1. no T_i is a link thread ready to act on an endpoint $x \in \text{bn}(\mathcal{C})$; and
2. no T_i and T_j are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$.

If condition (1) holds, $P \not\rightarrow_\alpha$. If condition (2) holds, $P \not\rightarrow_\beta$.

If a process contains two threads ready to act on dual endpoints, then it can reduce. The following lemma abstracts over the reduction rules for CP's many dual actions.

Lemma 2.31 (Reduction). If $(\nu x \bar{x})(P \parallel Q) \vdash \top$, and P and Q are ready to act on x and $\bar{x}$, respectively, there exists some R such that $(\nu x \bar{x})(P \parallel Q) \rightarrow R$.

Proof. By inversion on the derivation of $(\nu x \bar{x})(P \parallel Q) \vdash \top$. There are eight cases, which correspond exactly to *E-SEND*, *E-CLOSE*, *E-SELECT*₁, *E-SELECT*₂, and the variants under *E-EQUIV* with *SC-CUTCOMM*.

□

Progress states that any process is either in canonical form or can reduce. In essence, the proof shows that conditions (1) and (2) of the definition of canonical form correspond to the absence of α - and β -reduction.

Proposition 2.32 (Progress). If $P \vdash \top$, then either P is in canonical form, or there exists some Q such that $P \rightarrow Q$.

Proof. P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 1$).

If P is in canonical form, the result follows.

Otherwise, $n \geq 2$, and there are two cases:

Condition (1) does not hold. Some T_i (for $1 \leq i \leq n$) is a link thread ready to act on an endpoint $x \in \text{bn}(\mathcal{C})$. By definition, there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$. By Corollary 2.22 and Proposition 2.18, there exists some P_j (for $1 \leq j \leq n$ and $i \neq j$) such that $\bar{x} \in \text{fn}(P_j)$.

$$\begin{aligned}
 P &\stackrel{s}{\equiv} \mathcal{E}[(vx\bar{x})(\mathcal{F}_1[x \leftrightarrow y] \parallel \mathcal{F}_2[P_j])] && \langle \text{by Corollary 2.22 and SC-LINKCOMM} \rangle \\
 &\stackrel{s}{\equiv} \mathcal{E}[\mathcal{F}_1[(vx\bar{x})(x \leftrightarrow y \parallel \mathcal{F}_2[P_j])] && \langle \text{by Lemma 2.23} \rangle \\
 &\stackrel{s}{\equiv} \mathcal{E}[\mathcal{F}_1[\mathcal{F}_2[(vx\bar{x})(x \leftrightarrow y \parallel P_j)]]] && \langle \text{by Lemma 2.23} \rangle \\
 &\longrightarrow \mathcal{E}[\mathcal{F}_1[\mathcal{F}_2[P_j\{y/\bar{x}\}]]] && \langle \text{by E-LINK and E-CONG} \rangle
 \end{aligned}$$

Condition (2) does not hold. Some threads T_i and T_j (for $1 \leq i, j \leq n$) are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$.

$$\begin{aligned}
 P &\stackrel{s}{\equiv} \mathcal{E}[(vx\bar{x})(\mathcal{F}_i[T_i] \parallel \mathcal{F}_j[T_j])] && \langle \text{by Corollary 2.22} \rangle \\
 &\stackrel{s}{\equiv} \mathcal{E}[\mathcal{F}_i[(vx\bar{x})(T_i \parallel \mathcal{F}_j[T_j])] && \langle \text{by Lemma 2.23} \rangle \\
 &\stackrel{s}{\equiv} \mathcal{E}[\mathcal{F}_i[\mathcal{F}_j[(vx\bar{x})(T_i \parallel T_j)]]] && \langle \text{by Lemma 2.23} \rangle \\
 &\longrightarrow \mathcal{E}[\mathcal{F}_i[\mathcal{F}_j[R]]] && \langle \text{by Lemma 2.31 and E-CONG} \rangle
 \end{aligned}$$

□

2.2.4 Duality, Dependency, and Deadlock

The definition of canonical form (Definition 2.30) requires justification: it defines “canonical forms” as “processes that do not reduce”, which is an easy way to get in trouble by admitting processes as “canonical” when they are, in fact, stuck for undesirable reasons such as deadlock.

It is straightforward to argue that processes in canonical form *must* contain a process ready to act on a free endpoint, which puts Definition 2.30 on par with Wadler’s definition of canonical form (see § 2.3) and Lindley and Morris’ definition of blocking. The argument relies on the simple combinatorics of configuration contexts. Any configuration context with n holes must contain exactly $n - 1$ cuts, and therefore must create exactly $n - 1$ channels.

Lemma 2.33. *For any $\mathcal{C}^n$, $|\text{dn}(\mathcal{C}^n)| = n - 1$ and $|\text{bn}(\mathcal{C}^n)| = 2(n - 1)$.*

Proof. By induction on the structure of the configuration context $\mathcal{C}$.

□

The fact that *some* process must be ready to act on a free endpoint follows, since you cannot have n processes ready to act on $n - 1$ channels without having at least two of them ready to act on the same channel.

Proposition 2.34 (Canonical Form). *Any well-typed process P in canonical form contains a thread that is ready to act on a free endpoint.*

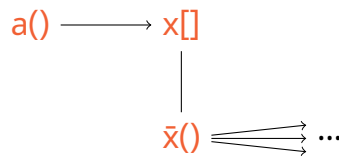
Proof. By contradiction. We have P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 1$). Assume no thread T_i is ready to act on a free endpoint. Then all n threads $T_1, \dots, T_n$ must be ready to act on bound endpoints. By linearity, all n threads $T_1, \dots, T_n$ must act on distinct bound endpoints. By Lemma 2.33, $\mathcal{C}^n$ binds $2(n - 1)$ endpoints in $n - 1$ dual pairs, so any subset of n bound endpoints must contain at least one dual pair. Therefore, at least two threads must be ready to act on dual endpoints, which contradicts the premise that P is in canonical form. $\square$

Unfortunately, that characterisation is *inadequate*, since it does not guarantee that the process cannot reduce. An adequate definition should to match the intuition that *all communication* in processes that cannot reduce is blocked on free endpoints. To formalise this notion, we must ensure that every ready action is blocked on an action on (1) a free endpoint or (2) a bound endpoint whose dual depends on some ready action that is blocked. For instance, in the process

$$(vx\bar{x})(a().x[].0 \parallel \bar{x}().P)$$

the action $a()$ is blocked on the free endpoint a , but the action $\bar{x}()$ is blocked because its dual action, $x[]$, depends on the action $a()$, which is blocked on the free endpoint a .

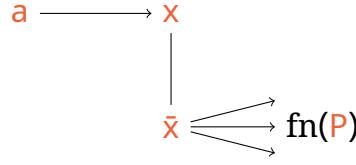
I formalise the notion of one action *depending* on another as the *shallow dependency graph* of a process. The dependency graph is a mixed graph, where undirected edges represent connected endpoints, either by link or by cut, and directed edges—or arcs—represent sequential dependencies. For instance, the shallow dependency graph for the above process is



where the arrows out of $\bar{x}()$ connect to the actions in P . The graph is *shallow* because it only tracks dependencies up to the *first action*, e.g. any dependencies within P are not tracked.

To define the dependency graph, we need some way to uniquely refer to actions. Unfortunately, the actions themselves are not unique—consider, e.g. $x \triangleleft \text{inl}. x \triangleleft \text{inl}. x[]. 0$. For shallow dependencies, it suffices to use endpoint names to refer to the first action on that endpoint. The

dependency graph of the above process then becomes



For the *deep* dependency graph, we could augment names with indices tracking the usage, e.g. letting $(x \triangleleft \text{inl}, 1)$ and $(x \triangleleft \text{inl}, 0)$ refer to the first and second occurrence of $x \triangleleft \text{inl}$ in $x \triangleleft \text{inl}. x \triangleleft \text{inl}. x[] . 0$. Alternatively, we could assign fresh vertices to each action (as is done in Priority CP, see Chapter 5). Fortunately, the shallow variants suffice for my purposes in this chapter.

The dependency graph is a mixed graph. I informally revisit the relevant definitions. For a detailed discussion, see § A.1.

- A *mixed graph* G has a set of vertices (denoted V_G , ranged over by u, v), a set of edges (denoted E_G), a set of arcs (denoted A_G). Edges are unordered pairs denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$. The set of edges may not contain loops uu . Arcs are ordered pairs denoted by juxtaposition overset with an arrow to indicate the direction, i.e. $\vec{uv} \triangleq (u, v)$. The set of arcs may not contain loops $\vec{uu}$.
- For any graph G with vertices $u, v \in V_G$, u is *adjacent* to v when there exists some edge $uv \in E_G$ or some arc $\vec{uv} \in A_G$.
- A *walk* w is a sequence of pairwise adjacent vertices.
- A *path* p is a walk with no repeated vertices, except possibly the first and last.
- A *cycle* c is a path that begins and ends at the same vertex.
- A walk is *essentially directed* when it contains at least one arc.
- A graph is *essentially acyclic* if and only if it contains no essentially directed cycles.
- The *undirected reachability* relation (denoted by $\sim_G$) is the equivalence closure over E_G .
- The *essentially directed reachability* relation (denoted by $<_G$) is the transitive closure over A_G quotiented by $\sim_G$.

The vertices of the dependency graph are endpoint names, which are a proxy for the first action on that endpoint. The edges represent channels, created by either links or cuts. The arcs represent dependencies, created by prefixing. For instance, in $a(). b[] . 0$, the process $b[] . 0$ is prefixed with the action $a()$. Hence, the action on b depends on the action on a .

Definition 2.35 (Dependency graph). *The shallow dependency graph of a process P , written $\text{Dep}(P)$, is a mixed graph (see § A.1) defined by the structure of the process P . The process P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for*

some $n \geq 1$). The shallow dependency graph $\text{Dep}(\mathbf{P})$ is defined as:

$$\begin{aligned} V_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \text{fn}(\mathbf{T}_i) \\ E_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \{\overrightarrow{xy} \mid \mathbf{T}_i = \mathbf{x} \leftrightarrow \mathbf{y}\} \cup \{\mathbf{x}\bar{x} \mid \{\mathbf{x}, \bar{x}\} \in \text{dn}(\mathcal{C})\} \\ A_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \{\overrightarrow{xy} \mid \mathbf{x}, \mathbf{y} \in \text{fn}(\mathbf{T}_i), \text{ready}(\mathbf{T}_i, \mathbf{x}) \wedge \neg \text{ready}(\mathbf{T}_i, \mathbf{y})\} \end{aligned}$$

By Lemma 2.21, $E_{\text{Dep}(\mathbf{P})}$ and $A_{\text{Dep}(\mathbf{P})}$ contain no loops. If G is (the subgraph of) some dependency graph, I write $\text{fn}(G)$ for its vertices, i.e. $\text{fn}(G) = V_G$.

The dependency graph gives us *duality* on actions, which is undirected reachability in the dependency graph.

Definition 2.36 (Duality). An endpoint $\mathbf{x}$ is dual to some endpoint $\mathbf{y}$ in $\mathbf{P}$, written $\mathbf{x} \sim_{\mathbf{P}} \mathbf{y}$, if and only if there exists an undirected path from $\mathbf{x}$ to $\mathbf{y}$ in $\text{Dep}(\mathbf{P})$.

If $\mathbf{x} \sim_{\mathbf{P}} \mathbf{y}$, the corresponding path in $\text{Dep}(\mathbf{P})$ may be arbitrarily long, as undirected edges arise from both cuts and links. Consider the process

$$(\nu \mathbf{x}\bar{\mathbf{x}})(\mathbf{a} \leftrightarrow \mathbf{x} \parallel \bar{\mathbf{x}} \leftrightarrow \mathbf{b})$$

The duality $\mathbf{a} \sim \mathbf{b}$ is witnessed by the path $(\mathbf{a}\mathbf{x}, \mathbf{x}\bar{\mathbf{x}}, \bar{\mathbf{x}}\mathbf{b})$. The structure generated by cuts and links does not fork, i.e. each component of the undirected subgraph of the dependency graph is a path.

The dependency graph also gives us *dependency* on actions, which is the converse of essentially directed reachability in the dependency graph.

Definition 2.37 (Dependency). An endpoint $\mathbf{x}$ depends on some endpoint $\mathbf{y}$, written $\mathbf{x} >_{\mathbf{P}} \mathbf{y}$, if and only if there exists an essentially directed path from $\mathbf{y}$ to $\mathbf{x}$ in $\text{Dep}(\mathbf{P})$.

One quirk of using endpoints as a proxy for actions is that the duality and dependency appear to “leak” restricted names, i.e. $\sim_{\mathbf{P}}$ and $>_{\mathbf{P}}$ are not relations over $\text{fn}(\mathbf{P})$, but relations over $\text{fn}(\mathbf{P}) \cup \text{bn}(\mathcal{C})$, where $\mathcal{C}$ is the maximum configuration context of $\mathbf{P}$. However, as discussed, these relations should be viewed as relations on the first actions on those endpoints, not the endpoints themselves.

A process is in deadlock if the dependency relation is not antisymmetric, or, equivalently, if there is an essentially directed cycle in the dependency graph that contains at least one arc.

Definition 2.38 (Deadlock). A process $\mathbf{P}$ is in deadlock, written $\text{deadlock}(\mathbf{P})$, if $\text{Dep}(\mathbf{P})$ contains an essential cycle.

Every well-typed CP process is deadlock-free.

Proposition 2.39. If $\mathbf{P} \vdash \Gamma$, then $\neg \text{deadlock}(\mathbf{P})$.

Proof. By induction on the derivation of $\mathbf{P} \vdash \Gamma$.

- Case T-LINK.

The process P is of the form $x \leftrightarrow y$.

$\text{Dep}(x \leftrightarrow y)$ is essentially acyclic as it has no arcs.

- Case T-CUT.

The process P is of the form $(\nu x \bar{x})(P_1 \parallel P_2)$. By induction $\text{Dep}(P_1)$ and $\text{Dep}(P_2)$ are essentially acyclic.

By T-CUT, $\text{fn}(P_1) \cap \text{fn}(P_2) = \emptyset$. Under the Barendregt convention, all bound names in P_1 and P_2 are distinct. Hence, $\text{Dep}(P_1)$ and $\text{Dep}(P_2)$ are disjoint.

$\text{Dep}(P)$ is essentially acyclic by Lemma A.2.

- Case T-SEND, T-RECV, T-CLOSE, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, or T-ABSURD.

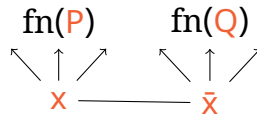
The process P is ready, i.e. $\text{ready}(P, x)$ for some endpoint x .

$\text{Dep}(P)$ is essentially acyclic as it has no edges and only arcs out of x .

□

The definition of deadlock is *shallow*. If a process is “free from deadlock”, that means the process is not in *immediate deadlock*. However, it does not mean that the process can never become deadlocked. Fortunately, the latter follows by type preservation. Since well-typed processes are free from immediate deadlock, and reduction preserves types, no well-typed process can ever reduce to a deadlocked process.

Care should be taken to only use Definition 2.38 for well-typed processes, as it does not imply session fidelity, the property that dual endpoints are used in dual ways. Hence, there are ill-typed processes that are morally in deadlock, but whose dependency graphs are essentially acyclic. For instance, the ill-typed process $(\nu x \bar{x})(x(). P \parallel \bar{x}(). Q)$ is *morally* in deadlock, but its dependency graph is essentially acyclic:



A *blocking action* is an action that blocks a process from making progress. For instance, in the process

$$(\nu x \bar{x})(a(). x[], 0 \parallel \bar{x}(). P)$$

the action $a()$ is *blocking*. However, not every ready action is *blocking*. The action $\bar{x}()$ is ready, but not blocking. Rather, it is *blocked*: its dual $x[]$ depends on $a()$, so it cannot reduce until $a()$ does.

As with dependency, I approximate blocking actions by their endpoints. *Blocking endpoint* are the maxima of the dependency relation, or, equivalently, the sources of the dependency graph. The blocking set of a process is the set of all sources of its dependency graph. Every ready action in a process is blocked on one of the endpoints in the blocking set.

Definition 2.40 (Blocking Set). *The blocking set of endpoints of a process P , written $\text{blocking}(P)$, is the set of sources of $\text{Dep}(P)$, i.e. $\{x \in V_{\text{Dep}(P)} \mid \nexists y.x \succ_P y\}$.*

The blocking set is closed under duality.

Lemma 2.41. *If $P \vdash \Gamma$ and $x \sim_P y$, then $x \in \text{blocking}(P) \implies y \in \text{blocking}(P)$.*

Each endpoint in the blocking set corresponds to a ready action.

Lemma 2.42. *If $P \vdash \Gamma$ and $x \in \text{blocking}(P)$, then P is of the form $\mathcal{C}[T]$ such that $\text{ready}(T, x)$.*

Due to the dualities generated by links, the blocking set may contain more endpoints than necessary. For instance, the blocking set of the process

$$(\forall x \bar{x})(x \leftrightarrow a \parallel \bar{x}().P)$$

is $\{x, \bar{x}, a\}$. An action in P is blocked on all of these endpoints. However, I want to be able to say that every action is blocked on a free name, and, in this case, the set $\{a\}$ suffices. A process is blocked on a set of endpoints if any action is blocked on at least one endpoint in that set.

Definition 2.43 (Blocked). *A process P is blocked on a set of endpoints X , written $\text{blocked}(P, X)$, if closing X under duality yields the blocking set $\text{blocking}(P)$, i.e. if $x \in X$ and $x \sim_P y$, then $y \in \text{blocked}(P)$.*

Any process is blocked on its blocking set.

Lemma 2.44. *If $P \vdash \Gamma$, then $\text{blocked}(P, \text{blocking}(P))$.*

If a process is blocked on some set of endpoints, it is blocked on the set formed by replacing any endpoint in that set with its dual.

Lemma 2.45. *If $P \vdash \Gamma$ and $x \sim_P y$, then $\text{blocked}(P, X) \implies \text{blocked}(P, X\{y/x\})$.*

If a process cannot β -reduce, then it is blocked on some set of free names.

Proposition 2.46. *If $P \vdash \Gamma$, then $P \not\rightarrow_\beta \iff \exists A \subseteq \text{fn}(P). \text{blocked}(P, A)$.*

Proof. Let P be of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 1$).

There are two cases:

- Case $(\implies)$.

By contradiction. Assume $x \in \text{blocking}(P)$ and $\nexists a \in \text{fn}(P). x \sim_P a$.

There are two cases:

- If $x \in \text{fn}(P)$, then $x \sim_P x$.
- If $x \in \text{bn}(\mathcal{C})$, then there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$.

There are two cases:

- * If $\bar{x} \in \text{blocking}(P)$, then there exist threads T_i and T_j that are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C})$. By Lemma 2.31, P is not β -free.
- * If $\bar{x} \notin \text{blocking}(P)$, then there exists some y such that $\bar{x} >_P y$. By definition, $x\bar{x} \in E_{\text{Dep}(P)}$. Therefore, $x >_P y$ and $x \notin \text{blocking}(P)$.

- Case $(\Leftarrow)$.

By contradiction. Assume $P \rightarrow_\beta$.

By inversion, there exist some T_i and T_j (for $1 \leq i, j \leq n$) that are ready to act on dual endpoints x and $\bar{x}$.

By definition, x and $\bar{x}$ only have outgoing arcs in $\text{Dep}(P)$, and the only edge connected to either is $x\bar{x}$. Therefore, $\{x, \bar{x}\} \subseteq \text{blocking}(P)$ and $\nexists a \in \text{fn}(P). x \sim_P a \vee \bar{x} \sim_P a$.

□

Corollary 2.47. *If $P \vdash \top$, then $\text{canonical}(P) \implies \text{blocking}(P) \subseteq \text{fn}(P)$.*

Proof. By Proposition 2.46, there exists a set $A \subseteq \text{fn}(P)$ such that $\text{blocked}(P, A)$. By definition, $A \subseteq \text{blocking}(P)$. It remains to show that $\text{blocking}(P) \subseteq A$.

By contradiction. Assume $x \in \text{blocking}(P)$ and $x \notin A$.

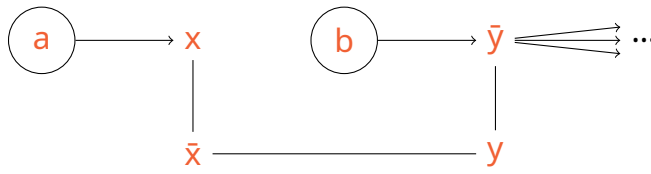
By definition, there exists some $a \in A$ such that $x \sim_P a$. The duality $x \sim_P a$ corresponds to some undirected path $p_{xa} = (x, \dots, a)$ in $\text{Dep}(P)$, which must contain at least one edge that connects some *bound* name y to some *free* name b , say, $y\bar{b}$. By definition, any edge in $\text{Dep}(P)$ generated by a cut connects two *bound* names. Therefore, $y\bar{b}$ must be generated by a link. By Lemma 2.42, $P = \mathcal{C}[y \leftrightarrow b]$. Hence, $P \rightarrow_a$.

□

Unfortunately, “blocked on free endpoints” does not characterise canonical forms, as $\text{blocking}(P) \subseteq \text{fn}(P) \not\Rightarrow P \not\rightarrow_a$. For instance, the process

$$(\nu x\bar{x})(a().x[].0 \parallel (\nu y\bar{y})(\bar{x} \leftrightarrow y \parallel b().\bar{y}().P))$$

can α -reduce. Its dependency graph, with blocking endpoints circled, is



In conclusion, my definition of canonical form (Definition 2.30) is *adequate*: any process in canonical form is blocked on some set of free endpoints. Due to the behaviour of links, “blocked on free endpoints” is insufficient to characterise canonical forms. As this would be a desirable property to have, I consider an alternative semantics for the link construct when discussing HCP in § 3.3. I have not adopted any of these alternatives as standard to maintain backwards compatibility with the work based on Wadler’s CP.

2.2.5 Connection and Right-Branching Form

In this section, I formalise the notion of the *connection graph* of a process, and show that any CP process can be rewritten into right-branching form.

Definition 2.48 (Right-branching Form). *A process P is in right-branching form if P is ready or P is of the form $(\nu x_1 \bar{x}_1)(T_1 \parallel \dots (\nu x_n \bar{x}_n)(T_n \parallel T_{n+1}) \dots)$ for some $n \geq 1$.*

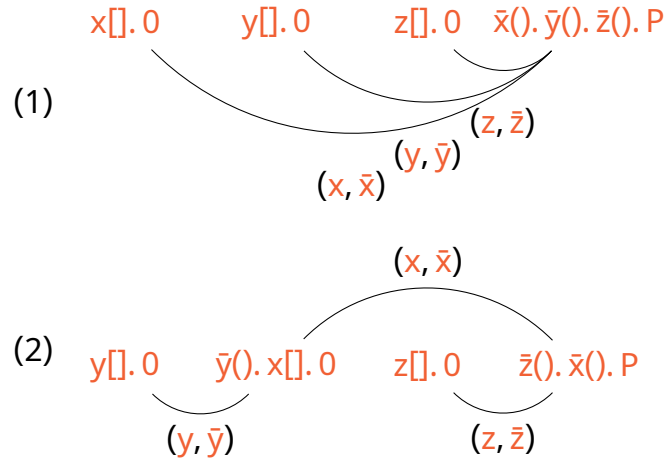
As discussed earlier, right-branching form is often used in the literature, but my metatheory avoids it, since it does not generalise to variants of CP that relax its rigid connection structure. So why am I talking about it? Of course, it is nice to justify a piece of reasoning that is frequently used. However, as it turns out, the theory I develop for converting processes to right-branching form will be important in the discussion of HCP.

Let us start by examining the connection graphs and right-branching forms of two simple example processes:

- (1) $(\nu x \bar{x})(x[], 0 \parallel (\nu y \bar{y})(y[], 0 \parallel (\nu z \bar{z})(z[], 0 \parallel \bar{x}().\bar{y}().\bar{z}().P)))$
- (2) $(\nu x \bar{x})(\nu y \bar{y})(y[], 0 \parallel \bar{y}().x[], 0) \parallel (\nu z \bar{z})(z[], 0 \parallel \bar{z}().\bar{x}().P)$

The connection graph of a process is the graph formed of all ready sub-processes and the channels that connect them. In the case of the example processes, each one consists of four ready sub-processes which

are connected by three channels:



Process (1) is already in right-branching form, though we could reorder the first three processes. Process (2) is not yet in right-branching form, but we can use the connection graph to convert it to right-branching form. The procedure picks a leaf from the connection graph, moves the corresponding cut and thread to the topmost, leftmost position, removes the leaf, and continues until all of the graph is empty. The process below is one of numerous different right-branching forms:

$$(\nu y \bar{y})(y[] . 0 \parallel (\nu x \bar{x})(\bar{y}().x[] . 0 \parallel (\nu z \bar{z})(z[] . 0 \parallel \bar{z}().\bar{x}().P)))$$

Our notion of connection graph is *shallow*, as opposed to *deep*, as we only account for the connections up to the maximum configuration context. However, the shallow definition suffices for our purposes and, as we shall see in Chapter 3, can be used to reason about the deep connection graph of a process by reasoning about the collection of shallow connection graphs of all sub-processes.

The connection graph is a undirected edge-labelled graph. I informally revisit the relevant definitions. For a detailed discussion, see § A.1.

- A *undirected edge-labelled graph* G has a set of vertices (denoted V_G , ranged over by u, v), a set of edges (denoted E_G), a set of edge labels (denoted $\mathcal{L}_G$), and an edge-labeling function (denoted ℓ_G) that assigns labels to edges. Edges are unordered pairs denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$. The set of edges may not contain loops uu . (It suffices to define V_G and ℓ_G , since $E_G \triangleq \text{dom}(\ell_G)$ and $\mathcal{L}_G \triangleq \text{cod}(\ell_G)$.)
- Two vertices $u, v \in V_G$ are *adjacent* when there exists an edge $uv \in E_G$.
- A *walk* w is a sequence of pairwise adjacent vertices.
- A *path* p is a walk with no repeated vertices, except possibly the first and last.
- A *cycle* c is a path that begins and ends at the same vertex.
- A graph is *acyclic* when it does not contain a cycle.

- A graph is *connected* when there is a path between any two vertices.
- A *tree* T is a graph that is connected and acyclic.

Definition 2.49 (Connection graph). *The shallow connection graph of a well-typed process P , written $\text{Con}(P)$, is an undirected edge-labelled graph (see § A.1) where the vertices are threads, the edges are the channels that connect those threads, and the edges are labelled by unordered pairs of their endpoints. The process P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 1$). The shallow connection graph $\text{Con}(P)$ is defined as:*

$$\begin{aligned} V_{\text{Con}(P)} &\triangleq \{T_1, \dots, T_n\} \\ \ell_{\text{Con}(P)} &\triangleq \{T_i T_j \mapsto (x, \bar{x}) \mid T_i, T_j \in V_{\text{Con}(P)} \wedge \{x, \bar{x}\} \in \text{dn}(\mathcal{C}) \wedge x \in T_i \wedge \bar{x} \in T_j\} \end{aligned}$$

By Lemma 2.21, $E_{\text{Con}(P)}$ contains no loops.

By Proposition 2.18, $\ell_{\text{Con}(P)}$ is a function.

If G is a subgraph of some connection graph, I write $\text{fn}(G)$ for the free names in the vertices of G , i.e. $\text{fn}(G) \triangleq \bigcup_{P \in V_G} \text{fn}(P)$.

In CP, the connection graph of a process is always a tree.

Proposition 2.50 (Connection Tree). *If P is well-typed, $\text{Con}(P)$ is a tree.*

Proof. By induction on the structure of $\mathcal{C}$ and inversion on the structure of P and the derivation of $P \vdash \Gamma$.

There are two cases:

- Case $\mathcal{C}$ is of the form $(vx\bar{x})(\mathcal{C}_1 \parallel \mathcal{C}_2)$.

By inversion, the typing derivation is of the form

$$\frac{Q \vdash \Gamma, x : A \quad R \vdash \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(Q \parallel R) \vdash \Gamma, \Delta} \text{T-CUT}$$

The process P is of the form $(vx\bar{x})(Q \parallel R)$.

By T-CUT, $\text{fn}(Q)$ and $\text{fn}(R)$ as well as $\text{Con}(Q)$ and $\text{Con}(R)$ are disjoint.

By induction, $\text{Con}(Q)$ and $\text{Con}(R)$ are trees.

The vertices of $\text{Con}(P)$ are exactly the union of those of $\text{Con}(Q)$ and $\text{Con}(R)$. Furthermore, the unordered pair $\{x, \bar{x}\}$ is the only element of $\text{dn}(\mathcal{C})$ that is not present in $\text{dn}(\mathcal{C}_1)$ or $\text{dn}(\mathcal{C}_2)$. By Proposition 2.18, there is exactly one $T_i \in V_{\text{Con}(Q)}$ and one $T_j \in V_{\text{Con}(R)}$ such that $x \in \text{fn}(T_i)$ and $\bar{x} \in \text{fn}(T_j)$. Therefore,

$$\begin{aligned} V_{\text{Con}(P)} &= V_{\text{Con}(Q)} \cup V_{\text{Con}(R)} \\ \ell_{\text{Con}(P)} &= \{T_i T_j \mapsto (x, \bar{x})\} \cup \ell_{\text{Con}(Q)} \cup \ell_{\text{Con}(R)} \end{aligned}$$

The result follows, as $\text{Con}(P)$ is formed by connecting trees $\text{Con}(Q)$ and $\text{Con}(R)$ with the single edge $T_i T_j$, and connecting two trees with a single edge always yields another tree.

- Case $\mathcal{C}$ is of the form $\square$.

The result follows, as $\text{Con}(\mathbf{P})$ is the singleton graph, which is a tree. $\square$

The tree structure of the connection graph can be used to rewrite any process to right branching form.

Proposition 2.51 (Right-branching Form). *For any well-typed process $\mathbf{P}$, there exists a process $\mathbf{Q}$ such that $\mathbf{P} \equiv \mathbf{Q}$ and $\mathbf{Q}$ is in right-branching form.*

Proof. By Proposition 2.50, $\text{Con}(\mathbf{P})$ is a tree. Let $\mathbf{T}_i$ be any leaf of $\text{Con}(\mathbf{P})$ (for $1 \leq i \leq |V_{\text{Con}(\mathbf{P})}|$). As $\mathbf{T}_i$ is a leaf, exactly one free endpoint in $\mathbf{T}_i$ is bound in $\mathbf{P}$. Let us name that endpoint $\mathbf{x}$. There exist some $\mathcal{E}$, $\mathcal{F}$, and $\mathbf{P}'$ such that:

$$\begin{aligned} \mathbf{P} &\equiv \mathcal{E}[(\mathbf{v}\mathbf{x}\bar{\mathbf{x}})(\mathcal{F}[\mathbf{T}_i] \parallel \mathbf{P}')] && \langle \text{by Corollary 2.22} \rangle \\ &\equiv \mathcal{E}[\mathcal{F}[(\mathbf{v}\mathbf{x}\bar{\mathbf{x}})(\mathbf{T}_i \parallel \mathbf{P}')]] && \langle \text{by Lemma 2.23 and SC-CUTCOMM} \rangle \\ &\equiv (\mathbf{v}\mathbf{x}\bar{\mathbf{x}})(\mathbf{T}_i \parallel \mathcal{E}[\mathcal{F}[\mathbf{P}']]) && \langle \text{by Lemma 2.23} \rangle \end{aligned}$$

By Lemma 2.24, $\mathcal{E}[\mathcal{F}[\mathbf{P}']]$ is well-typed. By induction on the process $\mathcal{E}[\mathcal{F}[\mathbf{P}']]$, there exists some $\mathbf{Q}'$ such that $\mathcal{E}[\mathcal{F}[\mathbf{P}']] \equiv \mathbf{Q}'$ and $\mathbf{Q}'$ is in right-branching form. Let $\mathbf{Q}$ be $(\mathbf{v}\mathbf{x}\bar{\mathbf{x}})(\mathbf{T}_i \parallel \mathbf{Q}')$. The result follows, as $\mathbf{P} \equiv \mathbf{Q}$ and $\mathbf{Q}$ is in right-branching form. $\square$

The tree structure of the connection graph implies deadlock freedom.

Proposition 2.52. *If $\text{Con}(\mathbf{P})$ is a tree, then $\neg \text{deadlock}(\mathbf{P})$.*

Sketch. Let us give the following definitions:

- $\mathbf{P}$ is of the form $\mathcal{C}^n[\mathbf{T}_1, \dots, \mathbf{T}_n]$ (for some $n \geq 1$).
- T is the connection graph $\text{Con}(\mathbf{P})$.
- G is the dependency graph $\text{Dep}(\mathbf{P})$.
- E_G^c are the cut-edges in G , i.e. $\{\mathbf{x}\bar{\mathbf{x}} \mid \{\mathbf{x}, \bar{\mathbf{x}}\} \in \text{dn}(\mathcal{C})\}$.
- E_G^l are the link-edges in G , i.e. $\bigcup_{1 \leq i \leq n} \{\mathbf{x}\mathbf{y} \mid \mathbf{T}_i = \mathbf{x} \leftrightarrow \mathbf{y}\}$.
- G/A_G is the quotient of G by its arcs A_G .

The result follows from the following facts:

1. Any essentially directed path p in G must alternate cut-edges with arcs or link-edges. (Cut-edges connect endpoints in different threads. Link-edges and arcs connect endpoints in the same thread. By definition, each ready link only generates a single link-edge, and every other thread only generates arcs out of the same vertex.)
2. There exists a surjective graph homomorphism from G/A_G to T , which preserves cut-edges and contracts the vertices connected by each link-edge, defined as:

$$f \triangleq \{\mathbf{X} \mapsto \mathbf{T}_i \mid 1 \leq i \leq n \wedge \mathbf{X} \subseteq \text{fn}(\mathbf{T}_i)\}$$

Assume P is in deadlock, i.e. there exists an essential cycle c in G .

There are three cases:

- If c contains no cut-edges, then, by fact (1), c must be a loop.
By definition, G contains no loops.
- If c contains one cut-edge, then, by fact (1), T must contain a loop.
By definition, T contains no loops.
- Otherwise, by fact (2), T contains a cycle.
This contradicts our premise.

□

Connection graphs are not merely a tool for converting processes to right-branching form and proving deadlock freedom. They are a full-fledged alternative representation for processes, with the interesting property that they represent the maximum configuration context of a process without any spurious ambiguity. Processes with equivalent maximum configuration contexts have *equal* connection graphs. Moreover, the correspondence works both ways! Processes with equal connection graphs have equivalent maximum configuration contexts.

Proposition 2.53. *If P is well-typed, then $P \stackrel{\text{ls}}{=} Q \iff \text{Con}(P) = \text{Con}(Q)$.*

Proof. There are two cases:

- Case $(\Rightarrow)$.

By induction on the proof of the structural congruence $P \stackrel{\text{ls}}{=} Q$. The cases for reflexivity, transitivity, symmetry, and SC-CONG follow by induction and those same properties of equality. The cases for SC-CUTCOMM and SC-CUTASSOC follow immediately.

- Case $(\Leftarrow)$.

Let $\text{Proc}(T)$ be the set of processes in right-branching form obtained from the connection tree T of a well-typed process by Proposition 2.51. (This is a set because Proposition 2.51 defines a non-deterministic procedure.)

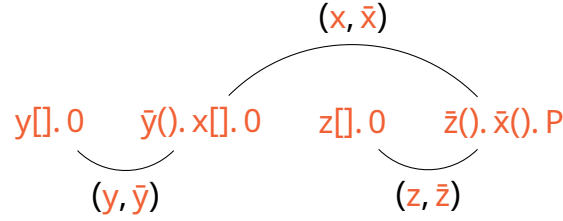
Pick any $R \in \text{Proc}(\text{Con}(P))$. As $\text{Con}(P) = \text{Con}(Q)$, $\text{Proc}(\text{Con}(P)) = \text{Proc}(\text{Con}(Q))$. Hence, $R \in \text{Proc}(\text{Con}(Q))$. By definition, $P \stackrel{\text{ls}}{=} R$ and $Q \stackrel{\text{ls}}{=} R$. Hence, $P \stackrel{\text{ls}}{=} Q$.

□

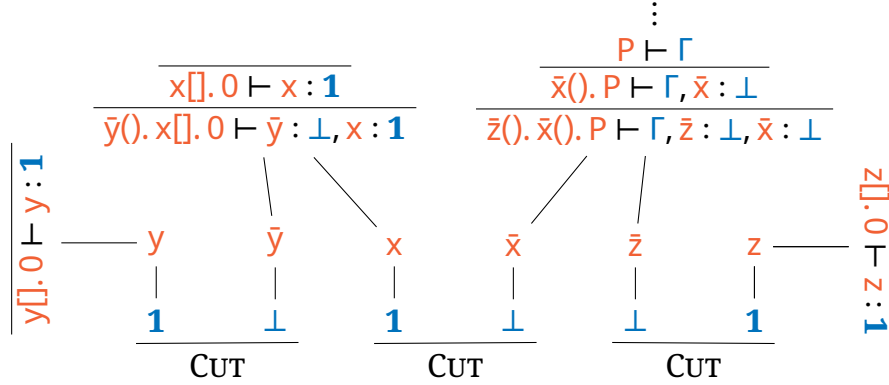
In conclusion, connection graphs are a unique representation for maximum configuration contexts that exactly capture link-preserving shallow structural congruence. That sounds awfully familiar. The canonical representation for Classical Linear Logic is *proof nets* [see Girard, 1987, p. 28]. Proof nets are a graphical representation of proofs

described as “classical natural deduction” which equate proofs up to various commutations. Connection graphs are *exactly* shallow proof nets that equate proofs up to the commutations SC-CUTCOMM and SC-CUTASSOC!

The translation from connection graphs to shallow proof nets is simple. Consider the connection graph for example (2) as presented above:



To create the corresponding shallow proof net, convert each process to its corresponding sequent calculus proof, convert each *edge* of the connection graph to a cut node, and connect each port of the cut node to the corresponding proposition in the sequent calculus proofs. (The vertical dots are the sequent calculus proof corresponding to the process P .)



The correspondence between connection graphs is well-known, e.g. Honda and Laurent [2010] discuss this relation between structural congruence and proof nets in the context of polarised linear logic. For CP, this correspondence opens up several interesting avenues for research. For instance, we could leverage the correspondence to construct a type-checking algorithm for CP’s connection graphs, rather than its process terms, based on the various correctness criteria for proof nets, such as Girard’s long trip criterion [Girard, 1987, p. 30], the Danos-Regnier criterion [Danos and Regnier, 1989], or Melliès’ ribbon criterion [Melliès, 2004]. Moreover, it would be interesting to examine the proof structures that correspond to stronger versions of structural congruence, such as the shallow or deep congruences defined in this chapter, or strong bisimulation with delayed actions, as defined, e.g. by Kokke et al. [2019a, DHCP], which I suspect has a much tighter correspondence to proof nets.

2.2.6 Observational Equivalence

In this section, I provide the basis for a behavioural theory for CP, namely, definitions for an observability predicate and barbed congruence. The purpose of this section is merely to show that these definitions can be given without the additional machinery introduced by Atkey [2017]. Therefore, I will not delve too deeply into the details. Access to an observability predicate also somewhat eases the discussion of commuting conversions in § 2.3.

The definition of observability predicates is non-standard, as it is not defined in terms of a label-transition system, but it is well-known, and follows Sangiorgi and Walker [2003, p. 56, Exercise 2.1.3]. The definition of the barbed congruence is standard, and follows Kokke et al. [2019a].

Definition 2.54 (Observability predicates). *The observability predicate $\downarrow_x P$ holds if and only if P is of the form $\mathcal{E}[Q]$ such that $\text{ready}(Q, x)$ and $\{y \in \text{fn}(Q) \mid \text{ready}(Q, y)\} \cap \text{bn}(\mathcal{E}) = \emptyset$.*

Definition 2.55 (Barbed congruence). *Barbed congruence, written $\approx$, is the largest symmetric relation on well-typed processes that is:*

1. *type preserving (i.e. if $P \approx Q$ and $P \vdash \Gamma$ then $Q \vdash \Gamma$)*
2. *barb preserving (i.e. if $P \approx Q$ and $\downarrow_x P$ then $\downarrow_x Q$)*
3. *context-closed (i.e. if $P \approx Q$ then $R[P] \approx R[Q]$ for any well-typed $R[\cdot]$)*

For a coarser barbed congruence, which distinguishes, e.g. $x \triangleleft \text{inl}.P$ and $x \triangleleft \text{inr}.P$, index the observability predicates with coarser labels [Yoshida et al., 2002, Atkey, 2017, Kokke et al., 2019a]. However, such considerations are outside of the scope of this section.

2.3 Commuting Conversions Considered Bad

Wadler’s CP has additional reduction rules, known as *commuting conversions* [Wadler, 2014, § 3.6]. The reason for adding these rules is that they allow the reduction strategy for Wadler’s CP to correspond step-by-step to a standard proof of cut elimination in Gentzen’s style [see Girard et al., 1989, § 13.2].

The commuting conversions are summarised in Figure 2.4. For clarity, I separate the rules into two different reduction relations [following Lindley and Morris, 2015]. I write $\longrightarrow$ for cut reductions and $\longrightarrow_{cc}$ for commuting conversions.⁵ Wadler’s reduction relation is defined as as

⁵My names for the reduction relations differ from those of Lindley and Morris [2015]. I use the right arrow ($\longrightarrow$) for CP’s the reduction relation without commuting conversions, which I consider the canonical reduction relation and is consistent with later chapters, and the right arrow labelled with a ‘W’ for Wadler ($\longrightarrow_w$) for Wadler’s reduction. Lindley and Morris [2015] use the right arrow ($\longrightarrow$) for Wadler’s reduction, which they consider the canonical reduction relation, and use the right arrow

$$\begin{array}{llll}
(\nu \bar{z}\bar{z})(x[y]. (P \parallel Q) \parallel R) & \longrightarrow_{cc} & x[y]. (P \parallel (\nu \bar{z}\bar{z})(Q \parallel R)) & \text{if } z \notin P \quad \text{CC-SEND}_1 \\
(\nu \bar{z}\bar{z})(x[y]. (P \parallel Q) \parallel R) & \longrightarrow_{cc} & x[y]. ((\nu \bar{z}\bar{z})(P \parallel R) \parallel Q) & \text{if } z \notin Q \quad \text{CC-SEND}_2 \\
(\nu \bar{z}\bar{z})(x(y). P \parallel Q) & \longrightarrow_{cc} & x(y). (\nu \bar{z}\bar{z})(P \parallel Q) & \text{CC-RECV} \\
(\nu \bar{z}\bar{z})(x(). P \parallel Q) & \longrightarrow_{cc} & x(). (\nu \bar{z}\bar{z})(P \parallel Q) & \text{CC-WAIT} \\
(\nu \bar{z}\bar{z})(x \triangleleft \text{inl}. P \parallel Q) & \longrightarrow_{cc} & x \triangleleft \text{inl}. (\nu \bar{z}\bar{z})(P \parallel Q) & \text{CC-SELECT}_1 \\
(\nu \bar{z}\bar{z})(x \triangleleft \text{inr}. P \parallel Q) & \longrightarrow_{cc} & x \triangleleft \text{inr}. (\nu \bar{z}\bar{z})(P \parallel Q) & \text{CC-SELECT}_2 \\
(\nu \bar{z}\bar{z})(x \triangleright \{\text{inl}: P; \text{inr}: Q\} \parallel R) & \longrightarrow_{cc} & & \\
& & x \triangleright \{\text{inl}: (\nu \bar{z}\bar{z})(P \parallel R); \text{inr}: (\nu \bar{z}\bar{z})(Q \parallel R)\} & \text{CC-OFFER}
\end{array}$$

$$\begin{array}{c}
\text{CC-ABSURD} \\
\frac{z \in N, \bar{z} \in \text{fn}(P) \quad N' = (N \setminus \{z\}) \cup (\text{fn}(P) \setminus \{\bar{z}\})}{(\nu \bar{z}\bar{z})(x \not\triangleleft N \parallel P) \longrightarrow_{cc} x \not\triangleleft N'}
\end{array}$$

$$\begin{array}{c}
\text{CC-SC} \\
\frac{P \equiv P' \quad P' \longrightarrow_{cc} Q' \quad Q' \equiv Q}{P \longrightarrow_{cc} Q}
\end{array}
\qquad
\begin{array}{c}
\text{CC-CONG} \\
\frac{P \longrightarrow_{cc} P'}{\mathcal{C}[P] \longrightarrow_{cc} \mathcal{C}[P']}
\end{array}$$

Figure 2.4: Commuting Conversions for Classical Processes

their union, written $\longrightarrow_w$:

$$\longrightarrow_w \triangleq \longrightarrow \cup \longrightarrow_{cc}$$

A consequence of the commuting conversions is that Wadler's CP has an appealingly simple canonical form: 'canonical form' simply means 'not a cut'. I call this *weak cut-free form*, by analogy to weak-head normal form.

Definition 2.56 (Weak Cut-Free Form). *A process is in weak cut-free form if it is not a cut, i.e. it is not of the form $(\nu x\bar{x})(P \parallel Q)$ for any $x, \bar{x}, P$, and Q .*

Equivalently, a process is in weak cut-free form if it is ready.

The commuting conversions are required to reduce processes to weak cut-free form. For instance, the process

$$(\nu x\bar{x})(a(). \bar{x}(). 0 \parallel x(). P)$$

is not in weak cut-free form, but cannot reduce with any cut reduction.

Lindley and Morris [2015, Page 574] claim that it suffices to reduce using only cut reductions, and use commuting conversions only to reduce from canonical form to weak cut-free form. I formalize and prove their claim as a pair of theorems.

labelled with a 'C' for *principal cut* ($\longrightarrow_c$) for the reduction relation without commuting conversions.

First, I relate Lindley and Morris' claim to my definition of canonical form by showing that any process in canonical form reduces to a process in weak cut-free form using only commuting conversions.

Proposition 2.57. *For any process P in canonical form, there is a reduction using only commuting conversions $P \rightarrow_{cc}^* P'$ such that P' is in weak cut-free form.*

Proof. There are two cases:

1. If P is ready, it is in weak cut-free form.
2. Otherwise, we have $P \equiv (\nu x_1 \bar{x}_1)(P_1 \parallel \dots (\nu x_n \bar{x}_n)(P_n \parallel P_{n+1}) \dots)$ such that (1) each P_i is ready; (2) no P_i is a link; and (3) no P_i and P_j are ready to act on dual endpoints of the same channel.

At least one process P_i is ready to act on a free endpoint: There are n pairs of dual endpoints $(x_1, \bar{x}_1, \dots, x_n, \bar{x}_n)$ and $n + 1$ processes $(P_1, \dots, P_{n+1})$. By (1), each process P_i is ready. If all $n + 1$ processes $P_1, \dots, P_{n+1}$ are ready to act on endpoints from the bound channels $x_1, \bar{x}_1, \dots, x_n, \bar{x}_n$, at least two processes must be ready to act on dual endpoints, which contradicts (3).

By (2), the process ready to act on a free endpoint cannot be a link. Therefore, we can use the commuting conversion to propagate that action to the top, and get the reduction $P \rightarrow_{cc}^* P'$ such that P' is in weak cut-free form.

□

Secondly, I show that it is *sufficient* to reduce using only cut reductions followed by only commuting conversions, by showing that any reduction in Wadler's CP can be rewritten to this form.

Proposition 2.58. *For any reduction $P \rightarrow_w^* R$, there is a reduction $P \rightarrow^* \rightarrow_{cc}^* R$.*

Proof. By iteratively commuting all commuting conversion past each reduction using Lemma 2.59. (For the full proof, see § 2.5.)

□

Lemma 2.59. *If $P \rightarrow_{cc}^* \rightarrow_{cc}^* R$, then $P \rightarrow \rightarrow_{cc}^* R$.*

Proof. By iteratively permuting each commuting conversion past the reduction using Lemma 2.60. (For the full proof, see § 2.5.)

□

Lemma 2.60. *If $P \rightarrow_{cc} R$, then either $P \rightarrow \rightarrow_{cc} R$ or $P \rightarrow R$.*

Proof. I will give the high-level intuition behind the proof, and give the formal proof for several interesting cases in § 2.5.

Recall that commuting conversions commute an action past a cut, e.g. to

$$(\text{vz}\bar{z})(\underline{x(y).P \parallel Q}) \longrightarrow_{cc} \underline{x(y).(vz\bar{z})(P \parallel Q)}$$

Recall that cut reductions eliminate actions, e.g.

$$(\text{vx}\bar{x})(\underline{x[y].(P \parallel Q)} \parallel \underline{\bar{x}(\bar{y}).R}) \longrightarrow (\text{vy}\bar{y})(P \parallel (\text{vx}\bar{x})(Q \parallel R))$$

It follows that for any reduction $P \longrightarrow_{cc} \longrightarrow R$ there are two options. The commuting conversion and cut reduction either act on *the same action* or on *different actions*:

- If they act on *the same action*, the commuting conversion is being used to move the action into place for the cut reduction. We can do this using the structural congruence instead.
- Otherwise, the two reduction steps act on *different actions* in different subprocesses. We can postpone the commuting conversion.

Unfortunately, we cannot easily formalize this high-level intuition, as actions are not a standalone syntactic sort in CP. Therefore, the formal proof proceeds by induction on the cut reduction $\longrightarrow R$ followed by inversion on the commuting conversion $P \longrightarrow_{cc}$.

(For the full proof, see § 2.5.)

□

What is the relation between weak cut-free form and canonical form? Can we split some reduction $P \xrightarrow{w} R$ into $P \xrightarrow{*} Q$ and $Q \xrightarrow{cc} R$ using Proposition 2.58, then forget about the commuting conversions, and expect the process Q to be in canonical form? The short answer is “no”.

Any process in weak cut-free form must be an action on a free endpoint, since there can be no bound endpoints at the top-level of a process. The commuting conversions can move an action on a free endpoint past any number of cuts. Therefore, a process can reduce to weak cut-free form using only the commuting conversions if it has any action on a free name under only cuts. As I discussed in § 2.2.4, that is inadequate as a definition for canonical form, e.g. the process

$$(\text{vx}\bar{x})(a().P \parallel (\text{vy}\bar{y})(y[], 0 \parallel \bar{y}().Q))$$

has an action on a free endpoint under a single cut but can still reduce.

Lemma 2.60 reveals an issue with commuting conversions: *redundancy*. Pairs of dual actions must be adjacent in the process configuration

before any reduction rule applies. This is a problem of syntax. The structural congruence dictates that process configurations can be permuted. However, in Wadler's CP, there are at least two ways to move dual actions into adjacent positions. We can either use the structural congruence or the commuting conversions. Worse, by combining both, we can pick any number of unrelated processes for the dual actions to commute over.

Example 2.61 (Redundancy). *The following process has two distinct reductions that eliminate the dual actions $x[]$ and $\bar{x}()$. One uses structural congruence, and the other uses commuting conversions.*

$$\begin{array}{ccc}
 & (vx\bar{x})(x[], 0 \parallel (vy\bar{y})(\bar{x}(). Q \parallel R)) & \\
 \equiv & \searrow_{cc} & \\
 (vx\bar{x})(x[], 0 \parallel \bar{x}(). (vy\bar{y})(Q \parallel R)) & & (vy\bar{y})((vx\bar{x})(x[], 0 \parallel \bar{x}(). Q) \parallel R) \\
 \searrow & & \swarrow \\
 & (vy\bar{y})(Q \parallel R) &
 \end{array}$$

The reduction strategy described by Wadler [2014, Proposition 2] does not meaningfully use the structural congruence; its only use is to derive the symmetric versions of the reduction rules, such as, e.g. the symmetric version of E-LINK:

$$\frac{\frac{(v\bar{x}x)(P \parallel x \leftrightarrow w) \equiv (vx\bar{x})(x \leftrightarrow w \parallel P) \quad (vx\bar{x})(x \leftrightarrow w \parallel P) \longrightarrow P\{w/\bar{x}\}}{(v\bar{x}x)(P \parallel x \leftrightarrow w) \longrightarrow P\{w/\bar{x}\}}$$

There is another issue with commuting conversions. To illustrate the issue, let us have a look at CC-WAIT:

$$(vz\bar{z})(x(). P \parallel Q) \longrightarrow_{cc} x(). (vz\bar{z})(P \parallel Q)$$

On the left-hand side of the reduction rule, Q is a top-level process. Therefore, any action on a free endpoint in Q can be observed, and any reduction $Q \longrightarrow_w Q'$ can happen. On the right-hand side of the reduction, both of these things are blocked by the action $x()$.

Example 2.62 (Blocked Observable). *The reduction with CC-WAIT on the action $a()$ blocks the observation $\downarrow_b$*

$$\begin{array}{c}
 (vx\bar{x})(a(). P \parallel b(). Q) \text{ has observables } \downarrow_a, \downarrow_b \\
 \downarrow \\
 a(). (vx\bar{x})(P \parallel b(). Q) \text{ has observables } \downarrow_a
 \end{array}$$

Example 2.63 (Blocked Reduction). *The reduction with CC-WAIT on the action $a()$ blocks the reduction with E-CLOSE on the actions $y[]$ and $y()$*

$$\begin{array}{c}
 (\nu x\bar{x})(a(). P \parallel (\nu y\bar{y})(y[], 0 \parallel \bar{y}(). Q)) \\
 \downarrow \qquad \qquad \qquad \downarrow \\
 (\nu x\bar{x})(a(). P \parallel Q) \quad a(). (\nu x\bar{x})(P \parallel (\nu y\bar{y})(y[], 0 \parallel \bar{y}(). Q)) \\
 \downarrow \qquad \qquad \qquad \downarrow \\
 a(). (\nu x\bar{x})(P \parallel Q)
 \end{array}$$

Example 2.63 reveals that Wadler’s CP is non-confluent! The processes at the bottom of the tree, $a(). (\nu x\bar{x})(P \parallel Q)$ and $a(). (\nu x\bar{x})(P \parallel (\nu y\bar{y})(y[], 0 \parallel \bar{y}(). Q))$, are distinct weak cut-free forms of $(\nu x\bar{x})(a(). P \parallel (\nu y\bar{y})(y[], 0 \parallel \bar{y}(). Q))$.

Canonical form is a stronger property than weak cut-free form. For some reductions $P \rightarrow_w^* R$ where R is weak cut-free, when we apply Proposition 2.58 to decompose the reduction into $P \rightarrow^* Q$ and $Q \rightarrow_{cc}^* R$, we find the process Q is not in canonical form: further cut reductions were initially possible, but became blocked by the commuting conversions. Formally, weak cut-free form corresponds to a variant of canonical form, called *weak canonical form*, which swaps the universal “all processes are ready” for an existential “some process is ready”.

Definition 2.64 (Weak Canonical Form). *A process P is in weak canonical form if P is of the form $\mathcal{E}[Q]$ such that Q is ready to act on a free endpoint that is not bound by $\mathcal{E}$.*

Equivalently, a process P is in weak canonical form if there is a reduction $P \rightarrow_{cc}^ P'$ such that P' is in weak cut-free form.*

As a consequence of the blocking behaviour of commuting conversions, the parallel composition of Wadler’s CP is not truly parallel composition. To illustrate this, let us each imagine an implementation of Wadler’s CP in our favourite programming language, and let us recall CC-WAIT:

$$(\nu z\bar{z})(x(). P \parallel Q) \rightarrow_{cc} x(). (\nu z\bar{z})(P \parallel Q)$$

Imagine that we implement each ready action P_i in a process as some program t_i running in its own thread. The program t_1 implements $x(). P$: it is blocked, waiting to receive a ping on the channel x . The programs $t_2, \dots, t_N$ implement Q : they are a collection of programs connected by channels, each program running in its own thread. There may be some communication between the programs $t_2, \dots, t_N$, but due to the typing rule for cut, only one of them is connected to t_1 .

The commuting conversion CC-WAIT requires that, because t_1 is blocked on x , *all* of the programs $t_2, \dots, t_N$ can become blocked on x . An implementation would require *some* communication between t_1 and

$t_2, \dots, t_N$, but this *cannot be* communication on shared channels, as t_1 only shares a channel with one of $t_2, \dots, t_N$. Therefore, when we implement the cut from Wadler's CP, we cannot simply use two parallel threads with a shared channel. We *must* add some kind of global scheduler. Consequently, implementations of Wadler's CP *must be* less parallel than we would expect from the syntax.

Formally, simulations of Wadler's CP in the π -calculus *must* decrease the degree of distribution, i.e. it is not possible to translate parallel composition to parallel composition.

Claim 2.65. *There exists no translation from the processes of Wadler's CP to processes of the π -calculus, written $\llbracket \cdot \rrbracket_\pi$, such that:*

- *The translation preserves the degree of distribution, i.e. is a homomorphism on parallel composition, which, due to CP's glued syntax, requires that the following equalities hold*

$$\begin{aligned} \llbracket (\nu x \bar{x})(P \parallel Q) \rrbracket_\pi &= (\nu x)(\llbracket P \rrbracket_\pi \parallel \llbracket Q\{x/\bar{x}\} \rrbracket_\pi) \\ \llbracket x[y].(P \parallel Q) \rrbracket_\pi &= x[y].(\llbracket P \rrbracket_\pi \parallel \llbracket Q \rrbracket_\pi) \end{aligned}$$

(Where the terms on the right-hand side denote π -calculus processes, i.e. (νx) is channel name restriction and $x[y]$ is bound output.)

- *The translation is a simulation, i.e. $P \rightarrow_w Q$ implies $\llbracket P \rrbracket_\pi \rightarrow_{\pi^*} \llbracket Q \rrbracket_\pi$. (Where $\rightarrow_{\pi^*}$ denotes the reflexive-transitive closure of the reduction relation for the π -calculus.)*

In summary, the commuting conversions in Wadler's CP are redundant, break confluence, and break the interpretation of parallel composition as parallel composition.

The version of CP presented in this chapter is confluent, but this is difficult to prove for the reduction semantics, as it requires showing that the structural congruence preserves the observable actions. Ultimately, a proof of congruence can be constructed by its operational correspondence with HCP (Chapter 3), which is the subject of the next chapter, and the proof of congruence for HCP by Montesi and Peressotti [2021].

2.4 Conclusion

In this chapter, I revisited Classical Processes and its metatheory. I dropped the commuting conversions, which cause Wadler's CP to be non-confluent, from the reduction semantics. I proved *preservation* (Proposition 2.27), and proved that *progress* (Proposition 2.32) continues to hold, albeit with a different canonical form (Definition 2.30). I proved that CP's processes are deadlock-free (Proposition 2.39), and that the new canonical forms are adequate (Corollary 2.47). I proved

that CP's connection graphs are trees (Proposition 2.50), and that any process can be rewritten to right-branching form (Proposition 2.51). Finally, I discussed the relation between CP with and CP without commuting conversions. I proved that any process in canonical form can reduce to a process in Wadler's canonical form using only commuting conversions (Proposition 2.57); and that any reduction that uses commuting conversions is equivalent to a reduction that does not use commuting conversions followed by a reduction using only commuting conversions (Proposition 2.58).

2.5 Omitted Proofs

Proposition 2.58. *For any reduction $P \rightarrow_w^* R$, there is a reduction $P \rightarrow^* \rightarrow_{cc}^* R$.*

Proof. By induction on the length of the reduction $P \rightarrow_w^* R$.

There are three cases:

1. The length is 0.

The result follows by reflexivity.

2. The length is $N + 1$ and the reduction is of the form $P \rightarrow_w^* R' \rightarrow R$.

By induction, we have $P \rightarrow^* \rightarrow_{cc}^* R'$. Hence, $P \rightarrow^* \rightarrow_{cc}^* R' \rightarrow R$.

The result follows by Lemma 2.59.

3. The length is $N + 1$ and the reduction is of the form $P \rightarrow_w^* R' \rightarrow_{cc} R$.

By induction, we have $P \rightarrow^* \rightarrow_{cc}^* R'$. Hence, $P \rightarrow^* \rightarrow_{cc}^* R' \rightarrow_{cc} R$.

The result follows immediately.

□

Lemma 2.59. *If $P \rightarrow_{cc}^* \rightarrow_{cc}^* R$, then $P \rightarrow_{cc}^* R$.*

Proof. By induction on the length of the reduction $P \rightarrow_{cc}^*$.

There are two cases:

1. The length is 0.

The reduction is of the form $P \rightarrow_{cc}^* R$.

The result follows immediately.

2. The length is $N + 1$.

The reduction is of the form $P \rightarrow_{cc}^* P' \rightarrow_{cc} R' \rightarrow_{cc}^* R$.

By Lemma 2.60, one of two cases must hold:

a. There is a reduction $P' \longrightarrow_{cc} R'$.

Hence, $P \longrightarrow_{cc}^* P' \longrightarrow_{cc} R' \longrightarrow_{cc}^* R$.

The result follows from the induction hypothesis.

b. There is a reduction $P' \longrightarrow R'$.

Hence, $P \longrightarrow_{cc}^* P' \longrightarrow R' \longrightarrow_{cc}^* R$.

The result follows from the induction hypothesis.

□

Lemma 2.60. *If $P \longrightarrow_{cc} R$, then either $P \longrightarrow_{cc} R$ or $P \longrightarrow R$.*

Proof. By induction on the cut reduction $\longrightarrow R$ followed by inversion on the commuting conversion $P \longrightarrow_{cc}$. The induction is guarded by the decreasing length of the derivation of the cut reduction. We examine three cases:

1. The cut reduction uses E-SEND and is of the form (reusing P and R)

$$(\nu x \bar{x})(x[y]. (P \parallel Q) \parallel \bar{x}(\bar{y}). R) \longrightarrow (\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel R))$$

By inversion, there are six cases for the commuting conversion. It uses either CC-SEND₁, CC-SEND₂, or CC-RECV under CC-CONG, or the symmetric version of any of these. We examine only the case where the reduction uses CC-RECV under CC-CONG and is of the form (reusing P , Q , and R)

$$\begin{aligned} & (\nu x \bar{x})(x[y]. (P \parallel Q) \parallel (\nu z \bar{z})(\bar{x}(\bar{y}). R \parallel R')) \\ & \longrightarrow_{cc} (\nu x \bar{x})(x[y]. (P \parallel Q) \parallel \bar{x}(\bar{y}). (\nu z \bar{z})(R \parallel R')) \\ & \longrightarrow (\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel (\nu z \bar{z})(R \parallel R'))) \end{aligned}$$

We replace the use of the commuting conversion with two uses of the structural congruence:

$$\begin{aligned} & (\nu x \bar{x})(x[y]. (P \parallel Q) \parallel (\nu z \bar{z})(\bar{x}(\bar{y}). R \parallel R')) \\ & \equiv (\nu z \bar{z})((\nu x \bar{x})(x[y]. (P \parallel Q) \parallel \bar{x}(\bar{y}). R) \parallel R') \\ & \longrightarrow (\nu z \bar{z})((\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel R)) \parallel R') \\ & \equiv (\nu y \bar{y})(P \parallel (\nu x \bar{x})(Q \parallel (\nu z \bar{z})(R \parallel R'))) \end{aligned}$$

2. The cut reduction uses E-CONG and is of the form (reusing P)

$$(\nu x \bar{x})(P \parallel Q) \longrightarrow (\nu x \bar{x})(P' \parallel Q)$$

The premise to E-CONG is of the form $P \longrightarrow P'$. By inversion, there are two cases for the commuting conversion. It uses CC-CONG.

- a. The reduction uses CC-CONG and is of the form

$$(\nu x \bar{x})(P'' \parallel Q) \rightarrow_{cc} (\nu x \bar{x})(P \parallel Q) \rightarrow (\nu x \bar{x})(P' \parallel Q)$$

The premise to CC-CONG is of the form $P'' \rightarrow_{cc} P$. The induction hypothesis gives us either $P'' \rightarrow_{cc} P'$ or $P'' \rightarrow P'$. In either case, the result follows immediately.

- b. The reduction uses CC-CONG and is of the form

$$(\nu x \bar{x})(P \parallel Q') \rightarrow_{cc} (\nu x \bar{x})(P \parallel Q) \rightarrow (\nu x \bar{x})(P' \parallel Q)$$

The premise to CC-CONG is of the form $Q' \rightarrow_{cc} Q$. The two reduction steps act on parallel processes, and can be reordered:

$$(\nu x \bar{x})(P \parallel Q') \rightarrow (\nu x \bar{x})(P' \parallel Q') \rightarrow_{cc} (\nu x \bar{x})(P' \parallel Q)$$

3. The cut reduction uses E-EQUIV and is of the form $Q \equiv Q' \rightarrow R' \equiv R$. The premise to E-EQUIV is of the form $Q' \rightarrow R'$. The induction hypothesis, applied to $P \rightarrow_{cc} Q' \rightarrow R'$, gives us either $P \rightarrow_{cc} R'$ or $P \rightarrow R'$. The result follows as $P \rightarrow_{cc} R$ or $P \rightarrow R$, respectively.

□

Part II

Taking Linear Logic Apart

Chapter 3

Hypersequent Classical Processes

This chapter presents Hypersequent Classical Processes (HCP), a session-typed process calculus based on CP as presented in Chapter 2. Nonetheless, the definitions, statements, and proofs in this chapter are self-contained, and do not reference those in Chapter 2, except where the goal is to relate CP and HCP.

HCP was independently developed by myself and by Fabrizio Montesi & Marco Peressotti [see Montesi and Peressotti, 2018]. Upon finding out about this parallel development, we collaborated on HCP, and published two versions:

- *Taking Linear Logic Apart* at Linearity-TLLA'18 [Kokke et al., 2019b]

The paper presents HCP as presented in this chapter. Unfortunately, the published version contains errors and uses notation and naming conventions that do not match later publications, so I have chosen not to include it in this thesis.¹

The paper refers to the calculus as HCP^{*}, rather than HCP, to distinguish it from the version below.

- *Better Late Than Never* at POPL'19 [Kokke et al., 2019a]

The paper presents Delayed HCP, which extends HCP with non-blocking actions. DHCP is not presented in this thesis.

The paper refers to the calculus as HCP, rather than DHCP, as we believed *at the time* that DHCP would be the preferred version. The authors have since changed their minds. I use HCP to refer to the

¹In May of 2019, Marco Peressotti discovered an error in the published version of Kokke et al. [2019b]. We submitted an erratum to EPTCS in July of 2019 which was approved for publication in October of 2021. Unfortunately, at the time of writing, the erratum has not yet been published. The version of the paper on Marco Peressotti's personal website, linked from the bibliography, contains an erratum.

version *without* delayed actions. Montesi and Peressotti [2021] use the name πLL instead of HCP².

HCP's main innovation is its type system, which lets us safely separate name restriction and parallel composition. Consequently, HCP more closely resembles the π -calculus and is more amenable to standard behavioural theory. Before we delve into the details, let us examine HCP at a glance, and investigate how it relates to Classical Processes, Classical Linear Logic [CLL, Girard, 1987], and the π -calculus [Milner et al., 1992b].

In this chapter, HCP's processes are printed in **red**, and its types are printed in **blue**, and both are rendered in a sans-serif font. To save on accessible colour combinations, the terms and types of *any other system* are printed in **pink** and **green**, respectively, both are rendered in an italicised font with serif, and any relations, such as typing and reduction, are marked by a subscript.

How does HCP relate to CP? The primary difference between CP and HCP is that HCP has a standalone process construct for parallel composition, whereas CP bakes parallel composition into the constructs for name restriction and sending—i.e. $(\nu x\bar{x})P$, $x[y].P$, and $P \parallel Q$ versus $(\nu x\bar{x})(P \parallel Q)$ and $x[y].(P \parallel Q)$. Likewise, HCP has a standalone process construct for the terminated process, which CP bakes into the construct for channel closing—i.e. $x[] . P$ and 0 versus $x[] . 0$.

To match the decomposition of processes, HCP decomposes its typing rules, and adds new typing rules for parallel composition and the terminated process:

$$\frac{P \vdash \mathcal{G} \quad Q \vdash \mathcal{H}}{P \parallel Q \vdash \mathcal{G} \parallel \mathcal{H}} \text{T-PAR} \quad \frac{}{0 \vdash \circ} \text{T-HALT}$$

To ensure that the properties of CP are preserved, only channel endpoints held by different processes are safe to connect. To this end, HCP adds the structural connectives “ $\parallel$ ” and “ $\circ$ ” which separate channel endpoints into groups corresponding to the processes that hold them. (Note that “ $\circ$ ” is distinct from the empty typing environment “ $\emptyset$ ”.)

Why are parallel composition and the terminated process typed by structural connectives rather than logical connectives? A channel endpoint is typed by a session type, and a process is typed by a typing environment that maps each free endpoint to its session type. An action acts on an endpoint, so the typing rules for actions are naturally the logical rules for some logical connective. However, parallel composition acts on processes and the terminated process simply is a process, so their typing rules must be structural rules for some structural connectives.

²Montesi and Peressotti's πLL continues the work on HCP. Of course, πLL is not a mere renaming. For instance, it includes the exponentials and second-order quantifiers.

To match the decomposition of processes, HCP also decompose the rules of its structural congruence, which become the standard rules for the π -calculus—such as associativity and commutativity for parallel composition, the unit laws for the terminated process, and scope extrusion. The reduction rules of HCP are exactly the same as those of CP.

I defer detailed discussion of the relation between HCP and CP to § 3.2. There we will see that HCP is a calculus for analysing collections of CP-like processes, and that HCP connection graphs are forests, rather than trees.

How does HCP relate to Classical Linear Logic? The crucial property of CP is its *exact* correspondence to linear logic. If you erase all that's written in red from the typing rules of CP, you get the inference rules of CLL, as given by Girard [1987]. Certainly, as HCP changes those typing rules, it must weaken that correspondence. I argue that it does the best possible job of preserving the property.

HCP is a conservative extension of CP. The corresponding logic, HCLL, is a conservative extension of CLL. Any CLL proof is an HCLL proof, and HCLL proves no new theorems about existing connectives. Formally,

$$\vdash_c \Gamma \iff \vdash \Gamma$$

where $\vdash_c \Gamma$ and $\vdash \Gamma$ denote sequents in CLL and HCLL, respectively.

HCLL extends CLL with the ability to reason about multiple unrelated proofs using *hypersequents* [Pottinger, 1983, Avron, 1991]. (Hence, Hypersequent CP.) A hypersequent logic has judgements over finite multisets of sequents, rather than single sequents. Avron writes hypersequents as:

$$\Gamma_1 \vdash \Delta_1 \mid \dots \mid \Gamma_n \vdash \Delta_n$$

As CLL and HCLL use one-sided sequents, the turnstile is somewhat superfluous, merely serving to remind us of the polarity of the typing environment, i.e. that commas are structurally pars rather than tensors. Frankly, the turnstile remains in CP only because *something* must separate the process and its typing environment and the colon is already taken. Hence, we might as well use the familiar turnstile. It looks pleasing and it reminds us of the polarity to boot. If we were to use Avron's notation for hypersequents, we would find ourselves writing the turnstile over and over. Worse, the turnstile would lose its primary function—to stand between a process and its typing environment. Instead, I present HCLL as a logic over finite multisets of *environments*, where environments are multisets of formulas. I refer to multisets of environments as *hyper-environments*. The turnstile keeps its position to the left of the hyper-environment, and continues to serve its primary function of separating term from type, red from blue. (Nonetheless, we continue to refer to HCLL as a hypersequent calculus.)

Hyper-environments, denoted by $\mathcal{G}$ and $\mathcal{H}$, are multisets of typing environments. I write $\circ$ for the empty hyper-environment, and $\parallel$ for hyper-environment concatenation. The rules for $(\parallel)$ and $(\circ)$ are as follows:

$$\frac{\vdash \mathcal{G} \quad \vdash \mathcal{H}}{\vdash \mathcal{G} \parallel \mathcal{H}} (\parallel) \quad \frac{}{\vdash \circ} (\circ)$$

What does a proof of a hyper-environment mean? A proof of a multiset of typing environments should imply a multiset of disjoint proofs, with one separate proof for each typing environment. In the larger proof, each of the separate proofs may be entangled—their inference rules mixed together—but each inference rule should belong to exactly one of the separate proofs, and we should be able to disentangle them into a sequence of CLL proofs:

$$\begin{array}{c} p \\ \vdots \\ \vdash \Gamma_1 \parallel \dots \parallel \Gamma_n \end{array} \Rightarrow \begin{array}{c} p_1 \quad \quad p_n \\ \vdots \quad \dots \quad \vdots \\ \vdash \Gamma_1 \quad \quad \vdash \Gamma_n \end{array}$$

There is a deep connection between the logical connective “ $\otimes$ ”, the structural connective “ $\parallel$ ”, and branching—having multiple premises—in multiplicative inference rules. All three capture some notion of *disjointness*:

- For branching, the disjointness is immediately apparent. Sequent calculus proofs are trees, so the premises are disjoint.
- For the structural connective “ $\parallel$ ”, the disjointness follows from the disentanglement property stated above.
- For the logical connective “ $\otimes$ ”, the disjointness follows from the splitting theorem [e.g. Girard, 1987, p. 39, Theorem 2.9.7, for CLL proof nets], which states that a proof of some judgement with subformula $A \otimes B$ can be decomposed into separate proofs of A and its environment, B and its environment, and a common prefix:

$$\begin{array}{c} p \\ \vdots \\ \vdash \Gamma, C[A \otimes B] \end{array} \Rightarrow \begin{array}{c} p_1 \quad \quad p_2 \\ \vdots \quad \quad \vdots \\ \vdash \Delta_1, A \quad \quad \vdash \Delta_2, B \\ \hline \vdash \Delta_1, \Delta_2, A \otimes B \\ \vdots \\ f \\ \vdots \\ \vdash \Gamma, C[A \otimes B] \end{array}$$

In summary, the connective “ $\parallel$ ” is a *structural* “ $\otimes$ ”, and the judgement $\vdash \Gamma \parallel \Delta$ means that the proofs of Γ and Δ are disjoint and can be disentangled into separate proofs. Likewise, the connective “ $\circ$ ” is a *structural* “**1**”, and the judgement “ $\vdash \circ$ ” means that there is nothing to prove.

Since hyper-environments internalise branching for multiplicative rules into the structure of the logic, the multiplicative rules may become unary.

For instance, HCLL's CUT rule has only one premise:

$$\frac{\vdash \mathcal{G} \parallel \Gamma, A \parallel \Delta, \bar{A}}{\vdash \mathcal{G} \parallel \Gamma, \Delta} \text{CUT}$$

The premise requires that the positive and negative occurrences of the cut formula are separated by a “ $\parallel$ ”, which ensures that their respective proofs are disjoint. In the conclusion, the environments are merged, since the proofs have become connected. This bookkeeping maintains the invariant that cut is only allowed to connect disjoint proofs, and avoids admitting rules that fundamentally alter the logic, such as MULTICUT [see, e.g. Atkey et al., 2016].

Likewise, HCLL's rule for tensor has only one premise, but there are two possible versions of the rule, which differ in whether or not the rule applies in the presence of unrelated proofs.

$$\frac{\vdash \Gamma, A \parallel \Delta, B}{\vdash \Gamma, \Delta, A \otimes B} (\otimes) \quad \frac{\vdash \mathcal{G} \parallel \Gamma, A \parallel \Delta, B}{\vdash \mathcal{G} \parallel \Gamma, \Delta, A \otimes B} (\otimes_d)$$

The two are logically equivalent: $(\otimes)$ is derivable from $(\otimes_d)$, and $(\otimes_d)$ is admissible using $(\otimes)$ and disentanglement. However, the two give rise to different process semantics:

- HCP, as presented in this chapter and in Kokke et al. [2019b, with errata], uses $(\otimes)$, which gives rise to the standard semantics of the π -calculus and is compatible with CP.
- DHCP, as presented in Kokke et al. [2019a], uses $(\otimes_d)$, which gives rise to delayed action semantics, as disentanglement requires that the action associated with $(\otimes_d)$ commutes with parallel composition.

How does HCP relate to the π -calculus? There is a slight difference in the semantics of communication channels, name restriction, and parallel composition, between HCP and the π -calculus, as presented by, e.g. Milner et al. [1992b] or Sangiorgi and Walker [2003]. Let us investigate this difference by asking a question: What makes a communication channel? There are at least two answers:

1. A communication channel is a name. Two processes can communicate simply by having access to the same name. Name restriction *restricts* the scope of a name, but communication channels exist irrespective of name restriction.
2. A communication channel is explicitly created. Two processes can only communicate if they have access to the same name bound in the scope of the name restriction that creates it. Unbound names are disconnected endpoints, not attached to any channel.

A bit tongue-in-cheek, I will refer to (1) as the *restrictive* view, and to (2) as the *creative* view. The π -calculus [Milner et al., 1992b] takes the

restrictive³ view, whereas Hypersequent CP takes the creative view. The two can be distinguished by their reduction rules—or, equivalently, their label-transition rules, which we will discuss later. Under the restrictive view, communication happens on free names, e.g. as in rule (a), and can happen on bound names simply by congruence. Under the creative view, communication can only happen on bound names, e.g. as in rule (b).

$$\begin{array}{cc} \text{(a)} & \text{(b)} \\ a\langle c \rangle. P \parallel a(y). Q \longrightarrow_{\pi} P \parallel Q\{c/y\} & (vx)(x\langle c \rangle. P \parallel x(y). Q) \longrightarrow_{\pi} (vx)(P \parallel Q\{c/y\}) \end{array}$$

The view we take has significant consequences for our type system:

1. Under the restrictive view, connection happens coincidentally. When two processes are composed in parallel, any number of connections can happen simultaneously, by the simple coincidence of names.
2. Under the creative view, connection happens intentionally. When two processes are composed in parallel, they are not connected, and they remain unconnected until they are intentionally connected by a v -binder, one channel at a time.

The creative view, due to its simplicity, is significantly more amenable to a correspondence with logic. To illustrate this, let us compare the typing rules for parallel composition from linear π -calculus [L π , Kobayashi et al., 1996], which takes the restrictive view, to Hypersequent CP, which takes the creative view:

$$\frac{P \vdash_{L\pi} \Gamma \quad Q \vdash_{L\pi} \Delta}{P \parallel Q \vdash_{L\pi} \Gamma + \Delta} \text{T-PAR} \quad \frac{P \vdash \mathcal{G} \quad Q \vdash \mathcal{H}}{P \parallel Q \vdash \mathcal{G} \parallel \mathcal{H}} \text{T-PAR}$$

Superficially, the rules are similar, but complexity hides in the details:

1. In L π , all session types are annotated with their capabilities—whether the corresponding channel is used to send, receive, neither, or both. The “+” is a partial function in the meta-language, which merges two typing environments by adding together the uses of a channel name in the two environments. In the well-typed cases, it computes the union of the usages, e.g. if x was used to send in Γ and used to receive in Δ , then it is fully used in $\Gamma + \Delta$. However, in the ill-typed cases, it is undefined, e.g. if x is used to send in both Γ and Δ , then $\Gamma + \Delta$ is undefined, and the typing rule for parallel composition does not apply.
2. In HCP, the “ $\parallel$ ” is a structural connective, much like the comma, and simply means “these resources are used in different processes”. The requirement for wellformedness is simpler: $\mathcal{G} \parallel \mathcal{H}$ is well-formed when the names in $\mathcal{G}$ and $\mathcal{H}$ are all distinct.

³Confusingly, the word “restriction” [as used by Milner et al., 1992a, submitted in 1989] connotes the restrictive view, but the later introduction of the letter “ v ” [pronounced as “new”, coined by Milner, 1991] connotes the creative view.

By taking the creative view, we get a typing rule for parallel composition which is much more logical, in the technical sense. What do we lose? We lose connection by coincidence, which I believe is not much of a loss, and perhaps even a win. So Hypersequent CP takes the creative view.

This chapter proceeds as follows:

- In § 3.1, I introduce Hypersequent CP.
- In § 3.2, I introduce the metatheory for Hypersequent CP.

I prove *preservation* (Proposition 3.30) and *progress* (Proposition 3.35), that its processes are deadlock-free (Corollary 3.42), that its canonical forms are adequate (Corollary 3.50), and that its communication graphs are forests (Proposition 3.52).

Notably, I also define disentanglement (Definition 3.83), which converts Hypersequent CP processes into multisets of CP processes, and prove it preserves typing (Proposition 3.84), structural congruence (Proposition 3.85), and reduction (Proposition 3.87), and define a label-transition system for HCP (Definition 3.88), and prove harmony (Proposition 3.93).

- In § 3.3, I discuss the relation between HCP and other developments in proof theory, and introduce several variants of HCP.

Notably, I introduce variants that interpret the absurd offer as session cancellation (§ 3.3.1), assign synchronous semantics to link (§ 3.3.2), and separate actions into their own syntactic sort by deconstructing the offer into a guarded summation (§ 3.3.3).

- Finally, § 3.5 contains all omitted proofs.

This chapter shares its structure with Chapter 2, and since the two systems share a fair amount of their structure, an equally fair amount of exposition is repeated, often verbatim, in Sections 3.1.2 to 3.1.6 and Sections 3.2.1 to 3.2.3. If you have read Chapter 2, it is worth reading § 3.1 up to and including § 3.1.1, the definition of configuration contexts in § 3.2.1, and then reading from § 3.2.5 onwards.

3.1 Hypersequent Classical Processes

In this section, I introduce Hypersequent Classical Processes (HCP), a session-typed process calculus based on CP. HCP's process calculus resembles the π -calculus more closely than CP, and its type system corresponds to a logic that is a slight variation of Classical Linear Logic, which I call Hypersequent Classical Linear Logic (HCLL). To the best of my knowledge, HCLL does not exist in the literature, and is only implicitly defined in this thesis by means of squinting and ignoring the

red parts. HCLL is a well-behaved and interesting logic, and, moreover, a conservative extension of CLL.

The fundamental notions of programs and computation in HCP, as in CP, are processes and message-passing communication. Processes communicate by passing messages over channels. Communication channels are *binary*, which means that each channel has exactly two endpoints, and each endpoint is held by exactly one process. Names refer to channel endpoints, rather than channels.

Processes (ranged over by P, Q, R) are defined by the following grammar:

$P, Q, R ::=$	$x \leftrightarrow y$	link
	$ (vx\bar{x})P$	new
	$ P \parallel Q$	parallel
	$ 0$	terminated process
	$ x[y]. P$	send
	$ x(y). P$	receive
	$ x[] . P$	close
	$ x(). P$	wait
	$ x \triangleleft \text{inl}. P$	select left
	$ x \triangleleft \text{inr}. P$	select right
	$ x \triangleright \{\text{inl}: P; \text{inr}: Q\}$	choice
	$ x \not\downarrow N$	absurd

The names x, y, z , and w range over the endpoints of communication channels—‘channel endpoints’ or ‘endpoints’ for short. The names $\bar{x}, \bar{y}, \bar{z}$, and $\bar{w}$ as well as a, b , and c also range over endpoints, and are used with the same conventions as in Chapter 2, which we revisit shortly. The names N and M range over sets of endpoints.

An endpoint is *bound* in the following cases:

- In $(vx\bar{x})P$, x and $\bar{x}$ are bound in P .
- In $x[y]. P$, y is bound in P .
- In $x(y). P$, y is bound in P .

An endpoint is *free* if it is not bound. Notably, for $x \not\downarrow N$, x and all names in N are free. I write $\text{fn}(P)$ to denote the set of free endpoints in P . By convention, the names a, b , and c are used as a shorthand to imply to the reader that the endpoint is free.

Two endpoints are *dual* if they are bound by the same name restriction, e.g. in $(vx\bar{x})P$, x and $\bar{x}$ are bound in P , and are dual. By convention, the names $\bar{x}, \bar{y}, \bar{z}$, and $\bar{w}$ are used as a shorthand to imply duality to the reader, e.g. I use x and $\bar{x}$ when they are dual endpoints of the same channel.

In HCP, as in CP, actions are not a well-defined syntactic sort. Nonetheless, I will informally write “action” in reference to the bit before the dot, e.g. the action for $x[y]. (P \parallel Q)$ is $x[y]$. For the troublemakers without a dot, $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ and $x \not\downarrow N$, I write $x \triangleright \text{inl}$, $x \triangleright \text{inr}$, and $x \not\downarrow$, respectively.

Types (ranged over by A, B) are the formulas of CLL, as defined by the following grammar:

$$A, B ::= A \otimes B \mid \mathbf{1} \mid A \oplus B \mid \mathbf{0} \\ \mid A \wp B \mid \perp \mid A \& B \mid \top$$

Duality plays an important role in HCP, as it does in CP and CLL. Viewed from the perspective of a logic, it corresponds to negation. Viewed from the perspective of a process calculus, it guarantees session fidelity, i.e. that processes act on dual endpoints of the same channel in dual ways, e.g. one process sends when the other receives. As in CP and CLL, duality is not defined as a type constructor, but as a function on types:

$$\begin{array}{lll} \overline{A \otimes B} & \triangleq & \overline{A} \wp \overline{B} \quad \overline{\mathbf{1}} \triangleq \perp \\ \overline{A \wp B} & \triangleq & \overline{A} \otimes \overline{B} \quad \overline{\perp} \triangleq \mathbf{1} \\ \overline{A \oplus B} & \triangleq & \overline{A} \& \overline{B} \quad \overline{\mathbf{0}} \triangleq \top \\ \overline{A \& B} & \triangleq & \overline{A} \oplus \overline{B} \quad \overline{\top} \triangleq \mathbf{0} \end{array}$$

As we will see, dual endpoints have dual types. The notation for duality ($A, \overline{A}$) and the naming convention for dual endpoints ($x, \overline{x}$) were chosen to emphasize this. Duality is involutive.

Lemma 3.1. $\overline{\overline{A}} = A$

Typing environments (ranged over by Γ, Δ) are sets of type assignments, as defined by the following grammar:

$$\Gamma, \Delta ::= \emptyset \mid \Gamma, x : A$$

The set of free endpoint names in a typing environment, written $\text{fn}(\Gamma)$, is defined by recursion on the environment, i.e. $\text{fn}(\emptyset) \triangleq \emptyset$ and $\text{fn}(\Gamma, x : A) \triangleq \text{fn}(\Gamma) \cup \{x\}$. The extension $\Gamma, x : A$ is only defined when x is not free in Γ , i.e. $x \notin \text{fn}(\Gamma)$.

We write Γ, Δ for the concatenation of typing environments Γ and Δ . The concatenation Γ, Δ is only defined when the names in Γ and Δ are unique, i.e. $\text{fn}(\Gamma) \cap \text{fn}(\Delta) = \emptyset$.

Hyper-environments (ranged over by $\mathcal{G}, \mathcal{H}$) are multisets of typing environments, as defined by the following grammar:

$$\mathcal{G}, \mathcal{H} ::= \circ \mid \mathcal{G} \parallel \Gamma$$

The set of free endpoint names in a hyper environment, written $\text{fn}(\mathcal{G})$, is defined by recursion on the hyper environment, i.e. $\text{fn}(\circ) \triangleq \emptyset$ and $\text{fn}(\mathcal{G} \parallel \Gamma) \triangleq \text{fn}(\mathcal{G}) \cup \text{fn}(\Gamma)$. The extension $\mathcal{G} \parallel \Gamma$ is only defined when the names in all typing environments in $\mathcal{G}$ and Γ are unique, i.e. $\text{fn}(\mathcal{G}) \cap \text{fn}(\Gamma) = \emptyset$.

$$\begin{array}{c}
\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \text{T-LINK} \qquad \frac{P \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}}{(vx\bar{x})P \vdash \mathcal{G} \parallel \Gamma, \Delta} \text{T-NEW} \\
\\
\frac{P \vdash \mathcal{G} \quad Q \vdash \mathcal{H}}{P \parallel Q \vdash \mathcal{G} \parallel \mathcal{H}} \text{T-PAR} \qquad \frac{}{0 \vdash \circ} \text{T-HALT} \\
\\
\frac{P \vdash \Gamma, x : A \parallel \Delta, y : B}{x[y]. P \vdash \Gamma, \Delta, x : A \otimes B} \text{T-SEND} \qquad \frac{P \vdash \Gamma, y : A, x : B}{x(y). P \vdash \Gamma, x : A \wp B} \text{T-RECV} \\
\\
\frac{P \vdash \circ}{x[] . P \vdash x : \mathbf{1}} \text{T-CLOSE} \qquad \frac{P \vdash \Gamma}{x(). P \vdash \Gamma, x : \perp} \text{T-WAIT} \\
\\
\frac{P \vdash \Gamma, x : A}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_1 \qquad \frac{P \vdash \Gamma, x : B}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_2 \\
\\
\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B}{x \triangleright \{\text{inl} : P; \text{inr} : Q\} \vdash \Gamma, x : A \& B} \text{T-OFFER} \qquad \frac{N = \text{fn}(\Gamma)}{x \not\triangleleft N \vdash \Gamma, x : \top} \text{T-ABSURD}
\end{array}$$

Figure 3.1: Typing Rules for Hypersequent CP

We write $\mathcal{G} \parallel \mathcal{H}$ for the concatenation of hyper-environments $\mathcal{G}$ and $\mathcal{H}$. The concatenation $\mathcal{G} \parallel \mathcal{H}$ is only defined when the names in all typing environments in $\mathcal{G}$ and $\mathcal{H}$ are unique, i.e. $\text{fn}(\mathcal{G}) \cap \text{fn}(\mathcal{H}) = \emptyset$.

We write $\mathcal{G}^k$ to mean that the hyper-environment $\mathcal{G}$ consists of k typing environments, i.e. $\mathcal{G}^k$ is of the form $\Gamma_1 \parallel \dots \parallel \Gamma_k$.

The typing judgment $P \vdash \mathcal{G}$ means that P is well-typed if, for each typing environment Γ in $\mathcal{G}$, and for each type assignment $x : A$ in Γ , exactly one process in P uses the endpoint x according to the session type A and that process only uses free endpoints that are in Γ .

Definition 3.2 (Typing). *A process P is well-typed under typing environment $\mathcal{G}$ if there exists a derivation with conclusion $P \vdash \mathcal{G}$ that uses the typing rules in Figure 3.1.*

Processes are considered equivalent up to structural congruence, written $\equiv$, which ensures that, e.g. the direction of a link and the order of a parallel composition are irrelevant.

Definition 3.3 (Structural Congruence). *Structural congruence, written $P \equiv Q$, is the congruence closure over processes which satisfies the rules in Figure 3.2.*

Definition 3.4 (Evaluation Context). *Evaluation contexts are one-hole process contexts, as defined by the following grammar:*

$$\mathcal{E}, \mathcal{F} ::= \square \mid (vx\bar{x})\mathcal{E} \mid \mathcal{E} \parallel Q \mid Q \parallel \mathcal{E}$$

$x \leftrightarrow y$	$\equiv y \leftrightarrow x$	SC-LINKCOMM
$P \parallel 0$	$\equiv P$	SC-PARNIL
$P \parallel Q$	$\equiv Q \parallel P$	SC-PARCOMM
$P \parallel (Q \parallel R)$	$\equiv (P \parallel Q) \parallel R$	SC-PARASSOC
$(\nu x \bar{x})P$	$\equiv (\nu \bar{x}x)P$	SC-NEWCOMM
$(\nu x \bar{x})(\nu y \bar{y})P$	$\equiv (\nu y \bar{y})(\nu x \bar{x})P$	SC-NEWASSOC
$(\nu x \bar{x})(P \parallel Q)$	$\equiv P \parallel (\nu x \bar{x})Q$	if $x, \bar{x} \notin \text{fn}(P)$ SC-SCOPEEXT

Figure 3.2: Structural Congruence for Hypersequent CP

Plugging is defined by replacing the one hole with a process:

$$\begin{aligned}
\Box & [P] \triangleq P \\
(\nu x \bar{x})\mathcal{E} [P] & \triangleq (\nu x \bar{x})\mathcal{E}[P] \\
(\mathcal{E} \parallel Q)[P] & \triangleq \mathcal{E}[P] \parallel Q \\
(Q \parallel \mathcal{E})[P] & \triangleq Q \parallel \mathcal{E}[P]
\end{aligned}$$

We write $\text{fn}(\mathcal{E})$ for the free endpoints in $\mathcal{E}$.

We write $\text{bn}(\mathcal{E})$ for the endpoints bound by $\mathcal{E}$.

We write $\mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H}$ to mean that the evaluation context $\mathcal{E}$ is well-typed under input typing context $\mathcal{G}$ and output typing context $\mathcal{H}$.

$$\begin{array}{c}
\frac{}{\Box \vdash \mathcal{G} \rightarrow \mathcal{G}} \quad \frac{\mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}}{(\nu x \bar{x})\mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H} \parallel \Gamma, \Delta} \\
\\
\frac{\mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H}_1 \quad Q \vdash \mathcal{H}_2}{\mathcal{E} \parallel Q \vdash \mathcal{G} \rightarrow \mathcal{H}_1 \parallel \mathcal{H}_2} \quad \frac{Q \vdash \mathcal{H}_1 \quad \mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H}_2}{Q \parallel \mathcal{E} \vdash \mathcal{G} \rightarrow \mathcal{H}_1 \parallel \mathcal{H}_2}
\end{array}$$

The semantics of HCP processes is given by *reduction*, written $\longrightarrow$. Reduction is closed over evaluation contexts, and structural congruence is embedded in reduction by allowing pre- and post-composition using E-CONG, written as $\equiv \longrightarrow$, $\longrightarrow \equiv$, or $\equiv \longrightarrow \equiv$.

Definition 3.5 (Reduction). *Reduction is the smallest relation on processes defined by the rules in Figure 3.3.*

In the remainder of the section, I discuss each process construct together with its typing rule and operational semantics, either by itself—e.g. *link*—or together with its dual—e.g. *send* and *receive*.

3.1.1 Process Structure

The process $(\nu x \bar{x})P$ denotes a *name restriction*, which creates a communication channel with two endpoints, x and $\bar{x}$, and the process $P \parallel Q$ denotes a *parallel composition*, which—well—composes two processes in parallel.

$$\begin{array}{lll}
(\nu x \bar{x})(x \leftrightarrow w \parallel P) & \longrightarrow P\{w/\bar{x}\} & \text{E-LINK} \\
(\nu x \bar{x})(x[y]. P \parallel \bar{x}(\bar{y}). Q) & \longrightarrow (\nu x \bar{x})(\nu y \bar{y})(P \parallel Q) & \text{E-SEND} \\
(\nu x \bar{x})(x[\cdot]. 0 \parallel \bar{x}(). Q) & \longrightarrow Q & \text{E-CLOSE} \\
(\nu x \bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) & \longrightarrow (\nu x \bar{x})(P \parallel Q) & \text{E-SELECT}_1 \\
(\nu x \bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) & \longrightarrow (\nu x \bar{x})(P \parallel R) & \text{E-SELECT}_2
\end{array}$$

$$\begin{array}{c}
\text{E-EQUIV} \\
\frac{P \equiv P' \quad P' \longrightarrow Q' \quad Q' \equiv Q}{P \longrightarrow Q}
\end{array}
\qquad
\begin{array}{c}
\text{E-CONG} \\
\frac{P \longrightarrow P'}{\mathcal{E}[P] \longrightarrow \mathcal{E}[P']}
\end{array}$$

Figure 3.3: Reduction for Hypersequent CP

The typing rules for name restriction and parallel composition are as follows:

$$\frac{P \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}}{(\nu x \bar{x})P \vdash \mathcal{G} \parallel \Gamma, \Delta} \text{T-NEW} \qquad
\frac{P \vdash \mathcal{G} \quad Q \vdash \mathcal{H}}{P \parallel Q \vdash \mathcal{G} \parallel \mathcal{H}} \text{T-PAR}$$

For communication safety, it is important that the two endpoints x and $\bar{x}$ have dual types, and that every endpoint is used exactly once.

- The T-NEW rule guarantees the former. It requires that $x : A$ and $\bar{x} : \bar{A}$.
- The T-PAR rule guarantees the latter. It requires that the processes P and Q do not share any free endpoints, as the concatenation $\mathcal{G} \parallel \mathcal{H}$ is only defined when the hyper-environments $\mathcal{G}$ and $\mathcal{H}$ do not share any endpoint names.

For deadlock freedom—as well as a close correspondence to CLL—it is important that the two endpoints of a channel are used in separate processes, and that any two processes share *at most* one channel. Both of these requirements are guaranteed by the T-NEW rule, but the T-PAR rule does the bookkeeping that enables this. Therefore, let us start our discussion with the T-PAR rule:

- In a parallel composition $P \parallel Q$, the two processes P and Q are not connected by any channel, as there is no name restriction that takes scope over both. Hence, there cannot be any dependency between any of their endpoints. The two processes are disjoint. The T-PAR rule registers this disjointness in the sequent by separating their respective hyper-environments with a “ $\parallel$ ”.
- In a name restriction $(\nu x \bar{x})P$, it is important that the two endpoints of a channel are used in separate processes, and that any two processes share *at most* one channel. The T-NEW rule guarantees both:
 - For the former, it requires that the two endpoints x and $\bar{x}$ are

separated by “ $\parallel$ ” in the premise, which means that they will *eventually* be split between two separate processes.

- For the latter, it removes the “ $\parallel$ ” in the conclusion, to register the fact that the two processes that use x and $\bar{x}$ are now connected by that channel.

The process 0 denotes the *terminated process*, which does nothing. The typing rule for the terminated process is as follows:

$$\frac{}{0 \vdash \circ} \text{T-HALT}$$

Hyper-environments are multisets. They are considered equal up to the unit rule for $\parallel$ and $\circ$, and the commutativity and associativity rules for $\parallel$.

$$\begin{aligned} \mathcal{G} \parallel \circ &= \mathcal{G} \\ \mathcal{G}_1 \parallel \mathcal{G}_2 &= \mathcal{G}_2 \parallel \mathcal{G}_1 \\ \mathcal{G}_1 \parallel (\mathcal{G}_2 \parallel \mathcal{G}_3) &= (\mathcal{G}_1 \parallel \mathcal{G}_2) \parallel \mathcal{G}_3 \end{aligned}$$

These structural rules for hypersequents are used implicitly, e.g. the following typing derivation implicitly uses the right unit rule.

$$\frac{\begin{array}{c} \vdots \\ P \vdash \mathcal{G} \end{array} \quad \frac{}{0 \vdash \circ}}{P \parallel 0 \vdash \mathcal{G}}$$

Processes are equivalent up to structural congruence, which includes the unit rule for $\parallel$ and 0 and the commutativity and associativity rules for $\parallel$. These explicit structural rules for processes mirror the implicit structural rules for hyper-environments.

$$\begin{array}{lll} P \parallel 0 & \equiv & P & \text{SC-PARNIL} \\ P \parallel Q & \equiv & Q \parallel P & \text{SC-PARCOMM} \\ P \parallel (Q \parallel R) & \equiv & (P \parallel Q) \parallel R & \text{SC-PARASSOC} \end{array}$$

The structural congruence also permits flipping the direction of a channel, permuting two name restrictions, and scope extrusion—which permits a process to move out of and into the scope of a name restriction, as long as it does not use the channel bound by that name restriction.

$$\begin{array}{lll} (vx\bar{x})P & \equiv & (v\bar{x}x)P & \text{SC-NEWCOMM} \\ (vx\bar{x})(vy\bar{y})P & \equiv & (vy\bar{y})(vx\bar{x})P & \text{SC-NEWASSOC} \\ (vx\bar{x})(P \parallel Q) & \equiv & P \parallel (vx\bar{x})Q & \text{if } x, \bar{x} \notin \text{fn}(P) \text{ SC-SCOPEEXT} \end{array}$$

There are no reduction rules associated with name restriction, parallel composition, or the terminated process. None of these constructs perform any action. Rather, they manage the structure of connections and parallel processes that facilitate message-passing communication.

Name restriction and parallel composition appear in every reduction rule, as communication can only happen over a channel, which requires

the presence of a name restriction, and between parallel processes, which requires the presence of a parallel composition. Furthermore, the E-CONG rule permits communication to occur under an arbitrary evaluation context of name restrictions, parallel compositions, and unrelated processes.

$$\frac{\text{E-CONG} \quad P \rightarrow P'}{\mathcal{E}[P] \rightarrow \mathcal{E}[P']}$$

3.1.2 Link

The process $x \leftrightarrow y$ denotes a *link*. It forwards any messages received on x to y , and vice versa. For communication safety, the two endpoints must have dual types. Hence, $x : A$ and $y : \bar{A}$.

$$\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \text{T-LINK}$$

HCP's semantics follows CP's semantics for link, and does not explicitly forward messages, but treats links as suspended α -renaming. When a link $x \leftrightarrow y$ reduces, it renames all occurrences of the dual of x to y , or all occurrences of the dual of y to x . In essence, this updates all the processes connected to the one side of the link to point directly at the other side, circumventing the link.

$$(\nu x \bar{x})(x \leftrightarrow w \parallel P) \rightarrow P\{w/\bar{x}\} \text{ E-LINK}$$

The renaming targets a bound name. Hence, there cannot be any other occurrences of that name, and the link can be removed.

Links are *commutative*. If two channels are connected by a link, the order in which they are connected is irrelevant. This property is captured by the following equivalence:

$$x \leftrightarrow y \equiv y \leftrightarrow x \text{ SC-LINKCOMM}$$

3.1.3 Send and Receive

The send and receive actions are dual:

- The process $x[y].P$ denotes a *send* action. It creates a fresh channel, names one endpoint of that channel y , sends the other endpoint over x , then continues as P .
- The process $x(y).P$ denotes a *receive* action. It receives an endpoint over x , names it y , then continues as P .

The typing rules for send and receive are as follows:

$$\frac{P \vdash \Gamma, x : A \parallel \Delta, y : B}{x[y].P \vdash \Gamma, \Delta, x : A \otimes B} \text{T-SEND} \quad \frac{P \vdash \Gamma, y : A, x : B}{x(y).P \vdash \Gamma, x : A \wp B} \text{T-RECV}$$

The behaviour of send and receive is given by the following rule:

$$(\nu x \bar{x})(x[y]. P \parallel \bar{x}(\tilde{y}). Q) \longrightarrow (\nu x \bar{x})(\nu y \tilde{y})(P \parallel Q) \quad \text{E-SEND}$$

The continuation of a send action is P , as opposed to CP's send, $x[y].(P \parallel Q)$, which requires that the endpoints y and x are immediately split between the processes P and Q . This is misleading. HCP *still* requires that the channels y and x are split between parallel processes. It does not require that the split happens *immediately*, settling for *eventually*. However, HCP's *eventually* is not as lenient as one might assume. As discussed in the introduction to Chapter 3, the T-SEND rule does not permit the presence of any unrelated typing environments—i.e. there is no generic $\mathcal{E}$ in the rule, as there is in the T-NEW rule. The same is true for the typing rules for the other actions. Consequently, no other action may come between the send action $x[y]$ and the parallel composition that splits x and y . The only process constructs that may come between those are unrelated name restrictions and parallel compositions. In effect, the only possible forms for a send action and its continuation are

$$x[y]. \mathcal{E}[P \parallel_y Q] \quad \text{and} \quad x[y]. \mathcal{E}[P \parallel_x Q].$$

Where $P \parallel_y Q$ denotes the parallel composition that splits x and y such that $x \in \text{fn}(P)$ and $y \in \text{fn}(Q)$. This may seem restrictive—and it is—but the purpose of HCP is to preserve the semantics of CP.

(The alternative—permitting unrelated typing environments—is explored in DHCP [Kokke et al., 2019a]. DHCP preserves the semantics of CLL, but not of CP. It uses *delayed* actions, where the actions in a process may to some extent resolve out of order.)

3.1.4 Close and Wait

The close and wait actions are dual:

- The process $x[] . P$ denotes a *close* action.
It sends a ping over x , then continues as P .
- The process $x(). P$ denotes a *wait* action.
It receives a ping over x , then continues as P .

(I say ‘ping’ to imply the interaction between close and wait is merely a synchronisation, and does not transmit any information.)

The typing rules for close and wait are as follows:

$$\frac{P \vdash \circ}{x[] . P \vdash x : \mathbf{1}} \quad \text{T-CLOSE} \qquad \frac{P \vdash \Gamma}{x(). P \vdash \Gamma, x : \perp} \quad \text{T-WAIT}$$

The behaviour of these two actions is given by the following rule:

$$(\nu x \bar{x})(x[] . P \parallel \bar{x}(). Q) \longrightarrow Q \quad \text{E-CLOSE}$$

The continuation of the close action is P , as opposed to CP's close, $x[] . 0$, which requires that the process immediately terminates. This is misleading. HCP *still* requires that the process immediately terminates, but as a consequence of its type system rather than its process syntax. The T-CLOSE rule requires that the continuation P is typed under the empty hyper-environment, and—as I will show Lemma 3.55—all processes typed under the empty hyper-environment are equivalent to the terminated process. In effect, the only possible forms for a close action and its continuation are variations of $x[] . 0 \parallel \dots \parallel 0$.

3.1.5 Select and Offer

The select and offer actions are dual:

- The process $x \triangleleft \text{inl}. P$ denotes a *left selection* action. It sends the label inl over x , then continues as P .
- The process $x \triangleleft \text{inr}. P$ denotes a *right selection* action. It sends the label inr over x , then continues as P .
- The process $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ denotes a *choice* action. It receives a label over x , and then continues as either P or Q , depending on which label was received.

The typing rules for select and offer are as follows:

$$\frac{P \vdash \Gamma, x : A}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_1 \quad \frac{P \vdash \Gamma, x : B}{x \triangleleft \text{inl}. P \vdash \Gamma, x : A \oplus B} \text{T-SELECT}_2$$

$$\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B}{x \triangleright \{\text{inl}: P; \text{inr}: Q\} \vdash \Gamma, x : A \& B} \text{T-OFFER}$$

The behaviour of these actions is given by the following rules:

$$\begin{aligned} (\nu x \bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) &\longrightarrow (\nu x \bar{x})(P \parallel Q) \quad \text{E-SELECT}_1 \\ (\nu x \bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) &\longrightarrow (\nu x \bar{x})(P \parallel R) \quad \text{E-SELECT}_2 \end{aligned}$$

As discussed in § 2.1.5, this syntax was adapted from Dardha and Gay [2018], because it is more easily generalized to variant types.

3.1.6 The Absurd Offer

The process $x \not\triangleright N$ denotes the *absurd offer*. It waits to receive a choice between *zero* alternatives. Such a choice cannot be made, which means that there is no corresponding select action, and no corresponding reduction rule. In essence, an absurd offer is inert. The absurd offer is the sole process that is allowed to leave endpoints unused, and the set of those unused endpoints is denoted by N .

$$\frac{N = \text{fn}(\Gamma)}{x \not\triangleright N \vdash \Gamma, x : \top} \text{T-ABSURD}$$

For an in-depth discussion of the absurd offer, its syntax, its inert semantics, and its relation to Wadler’s absurd offer, see § 2.1.6. I present an exceptional semantics for the absurd offer in § 3.3.1.

3.2 Metatheory

In this section, I introduce the metatheory for Hypersequent Classical Processes. The principal developments are as follows:

- In § 3.2.1, I give several preliminary definitions that are used throughout the discussion of the metatheory.
- In § 3.2.2, I prove preservation (Proposition 3.30).
- In § 3.2.3, I define canonical form (Definition 3.33) and prove progress (Proposition 3.35). The proof of progress is adapted from the proof for CP (Proposition 2.32) and first appeared in Kokke et al. [2019b].
- In § 3.2.4, I define dependency graphs for HCP processes (Definition 3.36). I prove that HCP is deadlock-free, as its dependency graphs are always acyclic (Corollary 3.42), and I prove that my definition of canonical form is adequate (Corollary 3.50).
- In § 3.2.5, I define connection graphs for HCP processes (Definition 2.49) and prove that HCP’s connection graphs are always forests (Proposition 3.52). The validity of right-branching forest form for HCP follows as a corollary.
- In § 3.2.6, I define Multiset CP, a process calculus whose processes are multisets of parallel CP processes. The principal use of Multiset CP is to clarify the correspondence between HCP and CP.
- In § 3.2.7, I define fission, the translation from Multiset CP into HCP, and its inverse, disentanglement-and-fusion, the translation from HCP into Multiset CP, and prove that these translations preserve types, structural congruence, and reduction.

3.2.1 Preliminaries

Configuration Contexts

Configuration contexts are multi-hole process contexts that consist only of name restrictions, parallel compositions, and the terminated process.

The notion of configuration contexts generalises neatly from CP to HCP. We follow the changes made to the process language, and decompose the cut into name restriction and parallel composition, and add a new constructor for the terminated process. The latter leads to an important

distinction between CP and HCP configuration contexts, which is that HCP configuration contexts can have zero holes.

Definition 3.6 (Configuration Context). Configuration contexts are n -hole process contexts, as defined by the following grammar:

$$\mathcal{C}[\cdot], \mathcal{D}[\cdot] ::= \square \mid 0 \mid (\nu x \bar{x})\mathcal{C}[\cdot] \mid \mathcal{C}[\cdot] \parallel \mathcal{D}[\cdot]$$

If there is risk of ambiguity, we explicitly write the number of holes in a configuration context with a superscript, e.g. as $\mathcal{C}^n[\cdot]$.

Plugging is defined by replacing the n holes with n processes, left to right:

$$\begin{aligned} \square & \quad [P] & \triangleq P \\ 0 & \quad [& \triangleq 0 \\ (\nu x \bar{x})\mathcal{C}^n[\cdot] & \quad [P_1, \dots, P_n] & \triangleq (\nu x \bar{x})(\mathcal{C}^n[P_1, \dots, P_n]) \\ (\mathcal{C}^n[\cdot] \parallel \mathcal{D}^k[\cdot]) & \quad [P_1, \dots, P_n, P_{n+1}, \dots, P_{n+k}] & \triangleq (\mathcal{C}^n[P_1, \dots, P_n] \parallel \mathcal{D}^k[P_{n+1}, \dots, P_{n+k}]) \end{aligned}$$

I write $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n]$ for the evaluation context focused on the i 'th hole in $\mathcal{C}[\cdot]$ such that $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n][P_i] = \mathcal{C}[P_1, \dots, P_i, \dots, P_n]$.

I write $\text{dn}(\mathcal{C}[\cdot])$ for the unordered pairs of dual endpoints bound by $\mathcal{C}[\cdot]$.

$$\begin{aligned} \text{dn}(\square) & \triangleq \emptyset \\ \text{dn}(0) & \triangleq \emptyset \\ \text{dn}((\nu x \bar{x})\mathcal{C}[\cdot]) & \triangleq \{\{x, \bar{x}\}\} \cup \text{dn}(\mathcal{C}[\cdot]) \\ \text{dn}(\mathcal{C}[\cdot] \parallel \mathcal{D}[\cdot]) & \triangleq \text{dn}(\mathcal{C}[\cdot]) \cup \text{dn}(\mathcal{D}[\cdot]) \end{aligned}$$

I write $\text{bn}(\mathcal{C}[\cdot])$ for the endpoints bound by $\mathcal{C}[\cdot]$, i.e. $\text{bn}(\mathcal{C}[\cdot]) \triangleq \bigcup \text{dn}(\mathcal{C}[\cdot])$.

I write $\mathcal{C}^n[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \rightarrow \mathcal{G}$ to mean that the configuration context $\mathcal{C}^n[\cdot]$ is well-typed under input hyper-environments $\mathcal{G}_1, \dots, \mathcal{G}_n$ and output hyper-environment $\mathcal{G}$. I use “.” and “ $\mathcal{G}_1 \mid \dots \mid \mathcal{G}_n$ ” to denote sequences of hyper-environments, which correspond, in left-to-right order, to the holes in the configuration context.

$$\frac{\frac{\frac{\square \vdash \mathcal{G} \rightarrow \mathcal{G} \quad 0 \vdash \cdot \rightarrow \circ}{\mathcal{C}^n[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \rightarrow \mathcal{G} \parallel \Gamma, x:A \parallel \Delta, \bar{x}:\bar{A}}}{(\nu x \bar{x})\mathcal{C}^n[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \rightarrow \mathcal{G} \parallel \Gamma, \Delta}}{\frac{\mathcal{C}^n[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \rightarrow \mathcal{G} \quad \mathcal{D}^k[\cdot] \vdash \mathcal{H}_{n+1} \mid \dots \mid \mathcal{H}_k \rightarrow \mathcal{H}}{\mathcal{C}^n[\cdot] \parallel \mathcal{D}^k[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \mid \mathcal{H}_{n+1} \mid \dots \mid \mathcal{H}_k \rightarrow \mathcal{G}, \mathcal{H}}}$$

Shallow Structural Congruence

Reduction is closed under evaluation contexts, not under arbitrary process contexts, and only acts on the topmost actions. As such, reduction only needs a structural congruence that is similarly closed under evaluation contexts—a *shallow* structural congruence. Therefore, it is useful to distinguish different kinds of structural congruence, based on which portions of the process they are permitted to rewrite.

Definition 3.7 (Shallow Structural Congruence). Shallow structural congruence, written $\dot{\equiv}$, is the smallest symmetric relation over processes that satisfies the rules in Figure 3.2 and is closed under evaluation contexts, as per the following rule:

$$\frac{\text{SC-CONG} \quad P \dot{\equiv} P'}{\mathcal{E}[P] \dot{\equiv} \mathcal{E}[P']}$$

Most rules of the structural congruence target name restriction and parallel composition. The odd one out is SC-LINKCOMM, which rewrites a link action. It will be useful to single out the portions of a structural congruence that rewrite links.

Definition 3.8 (Link-Preserving Structural Congruence). Link-preserving structural congruence, written $\dot{\equiv}_L$, is the congruence closure over processes that satisfies the rules in Figure 3.2 except for SC-LINKCOMM.

Finally, it will be useful to have variants which combine these restrictions. In practice, I only need link-preserving shallow structural congruence and deep structural congruence.

Definition 3.9 (Link-Preserving Shallow Structural Congruence). Link-preserving shallow structural congruence, written $\dot{\equiv}_L$, is the intersection of link-preserving and shallow structural congruence.

Definition 3.10 (Deep Structural Congruence). Deep structural congruence, written $\equiv$, is the equivalence closure of the complement of $\dot{\equiv}_L$ with respect to $\equiv$.

Any structural congruence can be decomposed into its link-preserving shallow and deep structural components.

Lemma 3.11. If $P \equiv Q$, then there exists some R such that $P \dot{\equiv}_L R$ and $R \equiv Q$.

Ready Processes and Threads

A process is ready if it is ready to perform some communication action, i.e. if it is a link or it is prefixed by an action. (The definition of ready remains essentially unchanged from CP, though, of course, the forms of send and close actions has changed.)

Definition 3.12 (Ready). A process P is ready to act on x , written $\text{ready}(P, x)$, if it is of one of the forms:

$$\begin{aligned} & x \leftrightarrow y \quad x[y].P \quad x[].P \quad x \triangleleft \text{inl}.P \quad x \triangleright \{\text{inl}: P; \text{inr}: Q\} \\ & y \leftrightarrow x \quad x(y).P \quad x().P \quad x \triangleleft \text{inr}.P \quad x \not\triangleleft N \end{aligned}$$

A process is ready if it ready to act on some endpoint.

In particular, that links $x \leftrightarrow y$ are considered ready to act on both x and y , and absurd $x \not\vdash N$ is *not* considered ready to act on the channels $y \in N$.

A process can be decomposed into a prefix of its cuts, and a series of threads connected by those cuts. Such a prefix is the *maximum configuration context*, in the sense that no further cuts can be added. (The definition of maximal configuration contexts is adjusted to permit configuration contexts with zero holes, i.e. the maximum configuration contexts for terminated processes, but otherwise unchanged.)

Definition 3.13 (Maximum Configuration Context). *The maximum configuration context $\mathcal{C}^n[\cdot]$ of a process P is the configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 0$) and (for $1 \leq i \leq n$) each P_i is ready. The processes P_i are the threads of P . Every process has a unique maximum configuration context.*

Likewise, evaluation contexts are maximal if no further cuts can be added. Informally, maximal evaluation contexts are paths to the threads contained within some process, so each maximum configuration context $\mathcal{C}^n[\cdot]$ gives us n distinct maximal evaluation contexts. (The definition of maximal evaluation contexts is unchanged from CP.)

Definition 3.14 (Maximal Evaluation Context). *A maximal evaluation context $\mathcal{E}$ of a process P is an evaluation context such that $P = \mathcal{E}[Q]$ and Q is ready. If $\mathcal{C}[\cdot]$ is the maximum configuration context of P , then $\mathcal{C}[P_1, \dots, \square_i, \dots, P_n]$ is a maximal evaluation context of P .*

Finally, we refer to the top-level ready processes as *threads*. A significant portion of HCP's metatheory deals with threads. Let the metavariable T range over threads.

Definition 3.15 (Thread). *A process P is a thread of Q if there exists some evaluation context $\mathcal{E}$ of Q such that $Q = \mathcal{E}[P]$ and P is ready. We say that Q contains the thread P to mean that P is a thread of Q . We say P is a thread when the process Q that P is a thread of can be inferred from context. Let T range over threads.*

The use of the thread metavariable allows us to succinctly decompose any process P into its maximum configuration context and its threads, by stating “ P is of the form $\mathcal{C}[T_1, \dots, T_n]$ ”, as the notation implies that each thread T_i is a ready process, and therefore that $\mathcal{C}[\cdot]$ is the maximum configuration context.

Process Contexts

For the occasional convenience, I define full n -hole process contexts, which are arbitrary processes with any number of holes. Process contexts may contain any process construct, and may contain holes in any position, including nested under actions. As such, process contexts generalise evaluation and configuration contexts.

Definition 3.16 (Process Context). *Process contexts are defined by the grammar for processes, extended with the hole constructor, written $\square$. A process context may have any number of holes.*

The names $P[\cdot]$, $Q[\cdot]$, and $R[\cdot]$ range over process contexts, where the trailing $[\cdot]$ is intended to help distinguish between process and process context metavariables, and denotes the position of the arguments for plugging. We write $P^n[\cdot]$ to denote that the process context $P[\cdot]$ has n holes.

Plugging, written $P[P_1, \dots, P_n]$, is defined by replacing the n holes in the process context $P[\cdot]$ with the processes $P_1, \dots, P_n$ in order from left to right.

Linearity

HCP has a *linear* type system. It ensures that resources are always used exactly once, and never copied or dropped. However, due to HCP's reuse of endpoint names, it may appear that resources are used multiple times. (For instance, the process $x().x[\cdot].0$ appears to use the endpoint x twice.)

I define linearity by means of a free name count—taking the sum across a parallel composition, the union across an offer, and counting the absurd offer as using all available resources. The definition of the free name count is easily adapted to HCP. (For an in-depth discussion, see § 2.2.1.)

Definition 3.17 (Free Name Count). *The multiset of free endpoints in P , written $\underline{\text{fn}}(P)$, is a multiset (see Definition A.3) with support set $\text{fn}(P)$ and multiplicity function $\mu_{\underline{\text{fn}}(P)}$.*

$$\begin{array}{ll}
 \underline{\text{fn}}(x \leftrightarrow y) & \triangleq \{x, y\} \\
 \underline{\text{fn}}((\nu x \bar{x})P) & \triangleq \underline{\text{fn}}(P) \setminus \{x, \bar{x}\} \\
 \underline{\text{fn}}(P \parallel Q) & \triangleq \underline{\text{fn}}(P) + \underline{\text{fn}}(Q) \\
 \underline{\text{fn}}(0) & \triangleq \{\} \\
 \underline{\text{fn}}(x[y].P) & \triangleq \{x\} + (\underline{\text{fn}}(P) \setminus \{x, y\}) \\
 \underline{\text{fn}}(x[\cdot].0) & \triangleq \{x\} \\
 \underline{\text{fn}}(x(y).P) & \triangleq \{x\} + (\underline{\text{fn}}(P) \setminus \{x, y\}) \\
 \underline{\text{fn}}(x().P) & \triangleq \{x\} + \underline{\text{fn}}(P) \\
 \underline{\text{fn}}(x \triangleleft \text{inl}.P) & \triangleq \{x\} + (\underline{\text{fn}}(P) \setminus \{x\}) \\
 \underline{\text{fn}}(x \triangleleft \text{inr}.P) & \triangleq \{x\} + (\underline{\text{fn}}(P) \setminus \{x\}) \\
 \underline{\text{fn}}(x \triangleright \{\text{inl}: P; \text{inr}: Q\}) & \triangleq \{x\} + ((\underline{\text{fn}}(P) \setminus \{x\}) \cup (\underline{\text{fn}}(Q) \setminus \{x\})) \\
 \underline{\text{fn}}(x \not\prec N) & \triangleq \{x\} + \{w \mid w \in N\}
 \end{array}$$

Note that the operation $\mathcal{X} \setminus X$ removes all occurrences of the elements in the set X from the multiset $\mathcal{X}$ (see Definition A.3).

Linearity states that, for well-typed processes, each endpoint in the typing environment is used exactly once in the process, and vice versa.

Proposition 3.18 (Linearity). *If $P \vdash \mathcal{G}$, then:*

- $x \in^k \underline{\text{fn}}(P) \implies k = 1$

- $x \in^1 \underline{\text{fn}}(\mathcal{P}) \iff x \in \text{fn}(\mathcal{G})$

Proof. By induction on the derivation of $\mathcal{P} \vdash \mathcal{G}$.

□

For any well-typed process $\mathcal{E}[\mathcal{P}]$, each endpoint bound by $\mathcal{E}$ is used exactly once in the process $\mathcal{P}$, and vice versa.

Corollary 3.19. *If $\mathcal{E}[\mathcal{P}] \vdash \mathcal{G}$, then:*

- $x \in^k \underline{\text{fn}}(\mathcal{P}) \implies k = 1$
- $x \in^1 \underline{\text{fn}}(\mathcal{P}) \iff x \in \text{bn}(\mathcal{E})$
- $\text{bn}(\mathcal{E}) \subseteq \text{fn}(\mathcal{P})$

For any well-typed configuration $\mathcal{C}^n[\mathcal{P}_1, \dots, \mathcal{P}_n]$, the processes $\mathcal{P}_1, \dots, \mathcal{P}_n$ must collectively use all the endpoints bound by $\mathcal{C}^n[\cdot]$ exactly once.

Corollary 3.20. *If $\mathcal{C}^n[\mathcal{P}_1, \dots, \mathcal{P}_n] \vdash \mathcal{G}$, then:*

- $x \in^k \bigcup_{1 \leq i \leq n} \underline{\text{fn}}(\mathcal{P}_i) \implies k = 1$
- $x \in^1 \bigcup_{1 \leq i \leq n} \underline{\text{fn}}(\mathcal{P}_i) \iff x \in \text{bn}(\mathcal{C}^n[\cdot])$
- $\text{bn}(\mathcal{C}^n[\cdot]) \subseteq \bigcup_{1 \leq i \leq n} \text{fn}(\mathcal{P}_i)$

Separation

Separation relates configuration contexts and evaluation contexts—it ‘zooms in’, from viewing a process as a series of connected processes, to viewing two specific processes and the cut connecting them. Separation also captures an essential property of HCP’s type system: dual endpoints must be in distinct processes, separated by a parallel composition.

Whereas CP has one separation lemma, HCP has two. One for parallel composition, and one for name restriction.

Lemma 3.21. *If $\mathcal{P} \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}, \mathcal{C}^n[\cdot]$ is a configuration context such that $\mathcal{P} = \mathcal{C}^n[\mathcal{P}_1, \dots, \mathcal{P}_n]$ (for some $n \geq 2$), and there exists some $\{x, \bar{x}\} \in \text{fn}(\mathcal{P})$ such that $x \in \text{fn}(\mathcal{P}_i)$ and $\bar{x} \in \text{fn}(\mathcal{P}_j)$ (for some $1 \leq i, j \leq n$), there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that*

1. $\mathcal{P} = \mathcal{E}[\mathcal{F}_i[\mathcal{P}_i] \parallel \mathcal{F}_j[\mathcal{P}_j]]$, or
2. $\mathcal{P} = \mathcal{E}[\mathcal{F}_j[\mathcal{P}_j] \parallel \mathcal{F}_i[\mathcal{P}_i]]$.

Proof. By induction on the structure of $\mathcal{C}^n[\cdot]$, followed by inversion on $\mathcal{P}$ and the corresponding typing derivation.

There are two cases:

- Case $\mathcal{C}^n[\cdot]$ is of the form $(\text{vy}\bar{y})\mathcal{C}^n[\cdot]$ (reusing $\mathcal{C}^n[\cdot]$).

By the induction on $\mathcal{C}^n[\cdot]$, there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $(\text{vy}\bar{y})\square$ to $\mathcal{E}$.

- Case $\mathcal{C}^n[\cdot]$ is of the form $\mathcal{D}_1^n[\cdot] \parallel \mathcal{D}_2^m[\cdot]$ (reusing n).

By inversion, $P = \mathcal{D}_1^n[P_1, \dots, P_n] \parallel \mathcal{D}_2^m[P_{n+1}, \dots, P_{n+m}]$. There are four subcases, depending on whether P_i and P_j are on the same or different sides of the parallel composition.

- Subcase $1 \leq i, j \leq n$.

By induction on $\mathcal{D}_1^n[\cdot]$, there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $\square \parallel \mathcal{D}_2^m[P_{n+1}, \dots, P_{n+m}]$ to $\mathcal{E}$.

- Subcase $n < i, j \leq m$.

By induction on $\mathcal{D}_2^m[\cdot]$, there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $\mathcal{D}_1^n[P_1, \dots, P_n] \parallel \square$ to $\mathcal{E}$.

- Subcase $1 \leq i \leq n < j \leq m$.

Let $\mathcal{E} = \square$, $\mathcal{F}_i = \mathcal{D}_1[P_1, \square_i, P_n]$, and $\mathcal{F}_j = \mathcal{D}_2[P_{n+1}, \square_j, P_{n+m}]$.

The result follows as (1).

- Subcase $1 \leq j \leq n < i \leq m$.

Let $\mathcal{E} = \square$, $\mathcal{F}_i = \mathcal{D}_1[P_{n+1}, \square_i, P_{n+m}]$, and $\mathcal{F}_j = \mathcal{D}_2[P_1, \square_j, P_n]$.

The result follows as (2).

□

Lemma 3.22 (Separation). *If $P \vdash \mathcal{G}$, and $\mathcal{C}^n[\cdot]$ is a configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 2$), and there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$ such that $x \in \text{fn}(P_i)$ and $\bar{x} \in \text{fn}(P_j)$ (for some $1 \leq i, j \leq n$), there exist $\mathcal{E}_1, \mathcal{E}_2, \mathcal{F}_i$, and $\mathcal{F}_j$ such that either*

1. $P = \mathcal{E}_1[(vx\bar{x})(\mathcal{E}_2[\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j]])]$,
2. $P = \mathcal{E}_1[(v\bar{x}x)(\mathcal{E}_2[\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j]])]$,
3. $P = \mathcal{E}_1[(vx\bar{x})(\mathcal{E}_2[\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i]])]$, or
4. $P = \mathcal{E}_1[(v\bar{x}x)(\mathcal{E}_2[\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i]])]$.

Proof. By induction on the structure of $\mathcal{C}[\cdot]$, followed by inversion on P and the corresponding typing derivation.

There are two cases:

- Case $\mathcal{C}^n[\cdot]$ is of the form $(vy\bar{y})\mathcal{C}^n[\cdot]$ (reusing $\mathcal{C}[\cdot]$).

By inversion, $P = (vy\bar{y})(\mathcal{D}_1^n[P_i, \dots, P_n])$.

There are three subcases:

- Subcase $x = y$ and $\bar{x} = \bar{y}$.

By Lemma 3.21, there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that:

- a. $P = \mathcal{E}[\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j]]$, or

$$\text{b. } P = \mathcal{E}[\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i]].$$

Let $\mathcal{E}_1 = \square$ and $\mathcal{E}_2 = \mathcal{E}$.

In case (a), the result follows as case (1).

In case (b), the result follows as case (3).

– Subcase $x = \bar{y}$ and $\bar{x} = y$.

By Lemma 3.21, there exist $\mathcal{E}$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that:

$$\text{a. } P = \mathcal{E}[\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j]], \text{ or}$$

$$\text{b. } P = \mathcal{E}[\mathcal{F}_j[P_j] \parallel \mathcal{F}_i[P_i]].$$

Let $\mathcal{E}_1 = \square$ and $\mathcal{E}_2 = \mathcal{E}$.

In case (a), the result follows as case (2).

In case (b), the result follows as case (4).

– Subcase $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$.

By induction on $\mathcal{C}[\cdot]$, there exist $\mathcal{E}_1$, $\mathcal{E}_2$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $(vy\bar{y})\square$ to $\mathcal{E}_1$.

• Case $\mathcal{C}^n[\cdot]$ is of the form $\mathcal{D}_1^n[\cdot] \parallel \mathcal{D}_2^k[\cdot]$ (reusing n).

By inversion, $P = \mathcal{D}_1^n[P_1, \dots, P_n] \parallel \mathcal{D}_2^k[P_{n+1}, \dots, P_k]$. By inversion on the fact that $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$, P_i and P_j must be on the same side of the parallel composition. There are two subcases:

– Subcase $1 \leq i, j \leq n$.

By induction on $\mathcal{D}_1^n[\cdot]$, there exist $\mathcal{E}_1$, $\mathcal{E}_2$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $\square \parallel \mathcal{D}_2^k[P_{n+1}, \dots, P_k]$ to $\mathcal{E}_1$.

– Subcase $n < i, j \leq k$.

By induction on $\mathcal{D}_2^k[\cdot]$, there exist $\mathcal{E}_1$, $\mathcal{E}_2$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that (1) or (2). The result follows by prepending $\mathcal{D}_1^n[P_1, \dots, P_n] \parallel \square$ to $\mathcal{E}_1$.

□

The separation lemma is rather precise, and gives us one of four equalities. However, all cases are equivalent up to structural congruence. Usually, it is easier to forget the exact case.

Corollary 3.23 (Separation). *If $P \vdash \mathcal{G}$, $\mathcal{C}^n[\cdot]$ is a configuration context such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 2$), and there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$ such that $x \in \text{fn}(P_i)$ and $\bar{x} \in \text{fn}(P_j)$ (for some $1 \leq i, j \leq n$), there exist $\mathcal{E}_1$, $\mathcal{E}_2$, $\mathcal{F}_i$, and $\mathcal{F}_j$ such that $P \equiv \mathcal{E}_1[(vx\bar{x})(\mathcal{E}_2[\mathcal{F}_i[P_i] \parallel \mathcal{F}_j[P_j]])]$.*

Proof. By Lemma 3.22, SC-NEWCOMM, and SC-PARCOMM.

□

Evaluation contexts commute with name restriction and parallel composition.

Lemma 3.24. *If $x, \bar{x} \notin \text{fn}(\mathcal{E}) \cup \text{bn}(\mathcal{E})$, then $\mathcal{E}[(\nu x \bar{x})P] \equiv (\nu x \bar{x})(\mathcal{E}[P])$.*

Proof. By induction on the structure of the evaluation context $\mathcal{E}$. □

Lemma 3.25. *If $\text{fn}(P) \cap \text{bn}(\mathcal{E}) = \emptyset$, then $\mathcal{E}[P \parallel Q] \equiv P \parallel \mathcal{E}[Q]$.*

Proof. By induction on the structure of the evaluation context $\mathcal{E}$. □

Corollary 3.26. *If $\text{fn}(P) \cap \text{bn}(\mathcal{E}) = \emptyset$, then $\mathcal{E}[P] \equiv \mathcal{E}[0] \parallel P$.*

Proof. By SC-PARNIL and Lemma 3.25. □

3.2.2 Preservation

Structural congruence preserves typing. If some process P is well-typed and is equivalent to some process Q under structural congruence, then Q is well-typed under the same typing environment.

Lemma 3.27. *If $P \equiv Q$, then $P \vdash \mathcal{G}$ if and only if $Q \vdash \mathcal{G}$.*

Proof. By induction on the derivation of the equivalence $P \equiv Q$.

The case for reflexivity follows immediately. The cases for symmetry and transitivity follow immediately by induction. The case for congruence closure follows by induction and the injectivity of the type derivation rules. The cases for applications of SC-LINKCOMM, SC-PARNIL, SC-PARCOMM, SC-PARASSOC, SC-NEWCOMM, and SC-SCOPEEXT are as follows, presented as equivalences on type derivations:

- Case SC-LINKCOMM:

$$\frac{}{x \leftrightarrow y \vdash x : A, y : \bar{A}} \equiv \frac{\frac{}{y \leftrightarrow x \vdash x : \bar{A}, y : A}}{y \leftrightarrow x \vdash x : A, y : \bar{A}} \text{ Lemma 2.1}$$

- Case SC-PARNIL.

$$\frac{P \vdash \mathcal{G} \quad \overline{0 \vdash \circ}}{P \parallel 0 \vdash \mathcal{G}} \equiv P \vdash \mathcal{G}$$

- Case SC-PARCOMM.

$$\frac{P_1 \vdash \mathcal{G}_1 \quad P_2 \vdash \mathcal{G}_2}{P_1 \parallel P_2 \vdash \mathcal{G}_1 \parallel \mathcal{G}_2} \equiv \frac{P_2 \vdash \mathcal{G}_2 \quad P_1 \vdash \mathcal{G}_1}{P_2 \parallel P_1 \vdash \mathcal{G}_1 \parallel \mathcal{G}_2}$$

- Case SC-PARASSOC.

$$\frac{\frac{P_1 \vdash \mathcal{G}_1 \quad \frac{P_2 \vdash \mathcal{G}_2 \quad P_3 \vdash \mathcal{G}_3}{P_2 \parallel P_3 \vdash \mathcal{G}_2 \parallel \mathcal{G}_3}}{P_1 \parallel (P_2 \parallel P_3) \vdash \mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \mathcal{G}_3}}{\frac{P_1 \vdash \mathcal{G}_1 \quad P_2 \vdash \mathcal{G}_2}{P_1 \parallel P_2 \vdash \mathcal{G}_1 \parallel \mathcal{G}_2} \quad P_3 \vdash \mathcal{G}_3}{P_1 \parallel (P_2 \parallel P_3) \vdash \mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \mathcal{G}_3}} \equiv$$

- Case SC-NEWCOMM.

$$\frac{P \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}}{(vx\bar{x})P \vdash \mathcal{G} \parallel \Gamma, \Delta} \equiv \frac{P \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}}{P \vdash \mathcal{G} \parallel \Gamma, x : \bar{A} \parallel \Delta, \bar{x} : \bar{A}} \text{ Lemma 2.1} \\ \frac{}{(vx\bar{x})P \vdash \mathcal{G} \parallel \Gamma, \Delta}$$

- Case SC-NEWASSOC.

$$\frac{\frac{P \vdash \mathcal{G} \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A} \parallel \Delta_1, y : B \parallel \Delta_2, \bar{y} : \bar{B}}{(vy\bar{y})P \vdash \mathcal{G} \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A} \parallel \Delta_1, \Delta_2}}{(vx\bar{x})(vy\bar{y})P \vdash \mathcal{G} \parallel \Gamma_1, \Gamma_2 \parallel \Delta_1, \Delta_2}}{\equiv} \\ \frac{P \vdash \mathcal{G} \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A} \parallel \Delta_1, y : B \parallel \Delta_2, \bar{y} : \bar{B}}{(vx\bar{x})P \vdash \mathcal{G} \parallel \Gamma_1, \Gamma_2 \parallel \Delta_1, y : B \parallel \Delta_2, \bar{y} : \bar{B}}}{(vy\bar{y})(vx\bar{x})P \vdash \mathcal{G} \parallel \Gamma_1, \Gamma_2 \parallel \Delta_1, \Delta_2}}$$

- Case SC-SCOPEEXT.

$$\frac{\frac{P \vdash \mathcal{G}_1 \quad Q \vdash \mathcal{G}_2 \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A}}{P \parallel Q \vdash \mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A}}}{(vx\bar{x})(P \parallel Q) \vdash \mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \Gamma_1, \Gamma_2}}{\equiv} \\ \frac{Q \vdash \mathcal{G}_2 \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A}}{(vx\bar{x})Q \vdash \mathcal{G}_2 \parallel \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A}}}{P \vdash \mathcal{G}_1 \quad (vx\bar{x})Q \vdash \mathcal{G}_2 \parallel \Gamma_1, \Gamma_2}}{P \parallel (vx\bar{x})Q \vdash \mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \Gamma_1, \Gamma_2}}$$

By inversion, these derivations are the only ones, as $x, \bar{x} \notin \text{fn}(P)$.

□

Renaming preserves typing. If a process is well-typed, then renaming any free endpoint does not affect its typing.

Lemma 3.28. *If $P \vdash \mathcal{G} \parallel \Gamma, x : A$, then $P\{w/x\} \vdash \mathcal{G} \parallel \Gamma, w : A$.*

Proof. The result follows by induction.

□

Plugging with any form of process context preserves typing.

Lemma 3.29. *If $P^n[\cdot] \vdash \mathcal{G}_1 \mid \dots \mid \mathcal{G}_n \rightarrow \mathcal{G}$ and $P_i \vdash \mathcal{G}_i$ (for $1 \leq i \leq n$), $P^n[P_1, \dots, P_n] \vdash \mathcal{G}$.*

Proof. By induction on the derivation of $P^n[\cdot] \vdash \mathcal{G}_1 \mid \cdots \mid \Gamma_n \rightarrow \mathcal{G}$.

□

Reduction preserves typing. If a process P is well-typed and reduces to some other process Q , then Q is well-typed under the same typing environment.

Proposition 3.30 (Preservation). *If $P \vdash \mathcal{G}$ and $P \rightarrow Q$, then $Q \vdash \mathcal{G}$.*

Proof. By induction on the derivation of the reduction $P \rightarrow Q$.

The case for E-CONG follows by induction and Lemma 3.29. The case for E-EQUIV follows by induction and Lemma 3.27. The cases for the rules E-LINK, E-SEND, E-CLOSE, E-SELECT₁, and E-SELECT₂ are as follows, presented as reductions on type derivations:

- Case E-LINK:

$$\begin{array}{c}
 \frac{x \leftrightarrow w \vdash x : A, w : \bar{A} \quad P \vdash \Gamma, \bar{x} : \bar{A}}{x \leftrightarrow w \parallel P \vdash x : A, w : \bar{A} \parallel \Gamma, \bar{x} : \bar{A}} \\
 \frac{}{(vx\bar{x})(x \leftrightarrow w \parallel P) \vdash \Gamma, w : \bar{A}} \\
 \downarrow \\
 \frac{P \vdash \Gamma, \bar{x} : \bar{A}}{P\{w/\bar{x}\} \vdash \Gamma, w : \bar{A}} \text{ Lemma 2.25}
 \end{array}$$

- Case E-SEND:

$$\begin{array}{c}
 \frac{P \vdash \Gamma_1, y : A \quad Q \vdash \Gamma_2, x : B}{P \parallel Q \vdash \Gamma_1, y : A \parallel \Gamma_2, x : B} \quad \frac{R \vdash \Gamma_3, \bar{y} : \bar{A}, \bar{x} : \bar{B}}{\bar{x}(\bar{y}). R \vdash \Gamma_3, \bar{x} : \bar{A} \wp \bar{B}} \\
 \frac{x[y].(P \parallel Q) \vdash \Gamma_1, \Gamma_2, x : A \otimes B \quad \bar{x}(\bar{y}). R \vdash \Gamma_3, \bar{x} : \bar{A} \wp \bar{B}}{x[y].(P \parallel Q) \parallel \bar{x}(\bar{y}). R \vdash \Gamma_1, \Gamma_2, x : A \otimes B \parallel \Gamma_3, \bar{x} : \bar{A} \wp \bar{B}} \\
 \frac{}{(vx\bar{x})(x[y].(P \parallel Q) \parallel \bar{x}(\bar{y}). R) \vdash \Gamma_1, \Gamma_2, \Gamma_3} \\
 \downarrow \\
 \frac{Q \vdash \Gamma_2, x : B \quad R \vdash \Gamma_3, \bar{y} : \bar{A}, \bar{x} : \bar{B}}{Q \parallel R \vdash \Gamma_2, x : B \parallel \Gamma_3, \bar{y} : \bar{A}, \bar{x} : \bar{B}} \\
 \frac{P \vdash \Gamma_1, y : A \quad (vx\bar{x})(Q \parallel R) \vdash \Gamma_2, \Gamma_3, \bar{y} : \bar{A}}{P \parallel (vx\bar{x})(Q \parallel R) \vdash \Gamma_1, y : A \parallel \Gamma_2, \Gamma_3, \bar{y} : \bar{A}} \\
 \frac{}{(vy\bar{y})(P \parallel (vx\bar{x})(Q \parallel R)) \vdash \Gamma_1, \Gamma_2, \Gamma_3}
 \end{array}$$

- Case E-CLOSE:

$$\begin{array}{c}
 \frac{x[\cdot].0 \vdash x : \mathbf{1} \quad Q \vdash \Gamma_2}{x[\cdot].0 \parallel \bar{x}().Q \vdash x : \mathbf{1} \parallel \Gamma_2, \bar{x} : \perp} \rightarrow Q \vdash \Gamma_2 \\
 \frac{}{(vx\bar{x})(x[\cdot].0 \parallel \bar{x}().Q) \vdash \Gamma_2}
 \end{array}$$

- Case E-SELECT₁:

$$\begin{array}{c}
\frac{P \vdash \Gamma_1, x : A}{x \triangleleft \text{inl}. P \vdash \Gamma_1, x : A \oplus B} \quad \frac{Q \vdash \Gamma_2, \bar{x} : \bar{A} \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{\bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_2, \bar{x} : \bar{A} \& \bar{B}} \\
\hline
\frac{x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_1, x : A \oplus B \parallel \Gamma_2, \bar{x} : \bar{A} \& \bar{B}}{(\nu x \bar{x})(x \triangleleft \text{inl}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) \vdash \Gamma_1, \Gamma_2} \\
\downarrow \\
\frac{P \vdash \Gamma_1, x : A \quad Q \vdash \Gamma_2, \bar{x} : \bar{A}}{P \parallel Q \vdash \Gamma_1, x : A \parallel \Gamma_2, \bar{x} : \bar{A}} \\
\hline
(\nu x \bar{x})(P \parallel Q) \vdash \Gamma_1, \Gamma_2
\end{array}$$

- Case E-SELECT₂:

$$\begin{array}{c}
\frac{P \vdash \Gamma_1, x : B}{x \triangleleft \text{inr}. P \vdash \Gamma_1, x : A \oplus B} \quad \frac{Q \vdash \Gamma_2, \bar{x} : \bar{A} \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{\bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_2, \bar{x} : \bar{A} \& \bar{B}} \\
\hline
\frac{x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\} \vdash \Gamma_1, x : A \oplus B \parallel \Gamma_2, \bar{x} : \bar{A} \& \bar{B}}{(\nu x \bar{x})(x \triangleleft \text{inr}. P \parallel \bar{x} \triangleright \{\text{inl}: Q; \text{inr}: R\}) \vdash \Gamma_1, \Gamma_2} \\
\downarrow \\
\frac{P \vdash \Gamma_1, x : B \quad R \vdash \Gamma_2, \bar{x} : \bar{B}}{P \parallel R \vdash \Gamma_1, x : B \parallel \Gamma_2, \bar{x} : \bar{B}} \\
\hline
(\nu x \bar{x})(P \parallel R) \vdash \Gamma_1, \Gamma_2
\end{array}$$

□

3.2.3 Progress

What should the canonical forms be for processes in HCP?

Unlike CP, HCP has syntax for the terminated process. We could make the distinction between *normal form* and *neutral form*. The normal form is the terminated process. The neutral forms are processes stuck on a free name, i.e. processes whose communication is blocked on a free channel.

Is that distinction meaningful? No, not in the context of HCP's reduction semantics. HCP processes do not *reduce to* normal form. After all, HCP preserves CP's semantics. Any process that does *anything* is never done.

Let us stick with the terminology “canonical form” for now. CP's definition for canonical form is adapted almost effortlessly. An HCP process P is in canonical form if and only if

1. P contains no link thread ready to act on a bound endpoint; and
2. P contains no two processes ready to act on dual endpoints.

Case (1) rules out α -reduction—any link evaluating as α -renaming. Case (2) rules out β -reduction—any other reduction. (The definitions of α - and β -reduction are unchanged from § 2.2.3.)

Definition 3.31 (α -Reduction). *A process P α -reduces to Q , written $P \rightarrow_\alpha Q$, if there exists a reduction $P \rightarrow Q$ which only uses the rules E-LINK, E-CONG, and E-EQUIV.*

Definition 3.32 (β -Reduction). *A process P β -reduces to Q , written $P \rightarrow_\beta Q$, if there exists a reduction $P \rightarrow Q$ that does not use the rule E-LINK.*

I divide reduction into α -reduction, which captures link evaluating as α -renaming, and β -reduction, which captures all other reduction. In essence, α -reduction captures asynchronous reduction, and β -reduction captures synchronous reduction.

Canonical forms are defined, as described above, by the forms of processes that cannot α - or β -reduce. The definition does not require that terminated processes are normalised to 0 , e.g. $0 \parallel 0$ is considered canonical. This simplifies the statement of progress, as we do not need to make allowances for processes such as $0 \parallel 0$ that do not reduce, but are not exactly 0 , and lets us defer proving Lemma 3.55. (The definition for canonical forms is adjusted to permit maximum configuration contexts with zero holes, i.e. the canonical forms of terminated processes.)

Definition 3.33 (Canonical Form). *A process P is in canonical form, written $\text{canonical}(P)$, if P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$) and (for $1 \leq i, j \leq n$) and*

1. *no T_i is a link thread ready to act on an endpoint $x \in \text{bn}(\mathcal{C}[\cdot])$; and*
2. *no T_i and T_j are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$.*

If condition (1) holds, $P \not\rightarrow_\alpha$. If condition (2) holds, $P \not\rightarrow_\beta$.

If a process contains two processes ready to act on dual endpoints, then it can reduce. The following lemma abstracts over the reduction rules for HCP's many dual actions, and is unchanged from CP.

Lemma 3.34 (Reduction). *If $(\nu x \bar{x})(P \parallel Q) \vdash \mathcal{G}$, and P and Q are ready to act on x and $\bar{x}$, respectively, there exists some R such that $(\nu x \bar{x})(P \parallel Q) \rightarrow R$.*

Proof. By inversion on the derivation of $(\nu x \bar{x})(P \parallel Q) \vdash \mathcal{G}$.

There are two cases, where either P or Q is a link, which correspond to E-LINK and its symmetric variant under E-EQUIV with SC-PARCOMM.

There are eight cases that correspond exactly to E-SEND, E-CLOSE, E-SELECT₁, E-SELECT₂, and their symmetric variants under E-EQUIV with SC-PARCOMM.

□

Progress states that any process is either in canonical form or can reduce. In essence, the proof shows that conditions (1) and (2) of the definition of canonical form correspond to the absence of α - and β -reduction.

Proposition 3.35 (Progress). *If $P \vdash \mathcal{G}$, then either P is in canonical form, or there exists some Q such that $P \rightarrow Q$.*

Proof. The process P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$).

If P is in canonical form, the result follows.

Otherwise, $n \geq 2$, and there are two cases:

Condition (1) does not hold. Some T_i (for $1 \leq i \leq n$) is a link thread ready to act on an endpoint $x \in \text{bn}(\mathcal{C}[\cdot])$. By SC-LINKCOMM, $T_i \stackrel{s}{=} x \leftrightarrow y$. By definition, there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$. By Corollary 3.23 and Proposition 3.18, there exists some T_j (for $1 \leq j \leq n$ and $i \neq j$) such that $\bar{x} \in \text{fn}(T_j)$.

$$\begin{aligned}
 P &\stackrel{s}{=} \mathcal{C}_1[(vx\bar{x})(\mathcal{C}_2[\mathcal{F}_1[x \leftrightarrow y] \parallel \mathcal{F}_2[T_j]])] && \langle \text{by Corollary 3.23 and SC-LINKCOMM} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[(vx\bar{x})(\mathcal{F}_1[x \leftrightarrow y] \parallel \mathcal{F}_2[T_j])]] && \langle \text{by Lemma 3.24} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_1[(vx\bar{x})(x \leftrightarrow y \parallel \mathcal{F}_2[T_j])]]] && \langle \text{by Lemma 3.25} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_1[\mathcal{F}_2[(vx\bar{x})(x \leftrightarrow y \parallel T_j)]]]] && \langle \text{by Lemma 3.25} \rangle \\
 &\rightarrow \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_1[\mathcal{F}_2[T_j\{y/\bar{x}\}]]]] && \langle \text{by E-LINK and E-CONG} \rangle
 \end{aligned}$$

Condition (2) does not hold. Some T_i and T_j (for $1 \leq i, j \leq n$) are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$.

$$\begin{aligned}
 P &\stackrel{s}{=} \mathcal{C}_1[(vx\bar{x})(\mathcal{C}_2[\mathcal{F}_i[T_i] \parallel \mathcal{F}_j[T_j]])] && \langle \text{by Corollary 3.23} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[(vx\bar{x})(\mathcal{F}_i[T_i] \parallel \mathcal{F}_j[T_j])]] && \langle \text{by Lemma 3.24} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_i[(vx\bar{x})(T_i \parallel \mathcal{F}_j[T_j])]]] && \langle \text{by Lemma 3.25} \rangle \\
 &\stackrel{s}{=} \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_i[\mathcal{F}_j[(vx\bar{x})(T_i \parallel T_j)]]]] && \langle \text{by Lemma 3.25} \rangle \\
 &\rightarrow \mathcal{C}_1[\mathcal{C}_2[\mathcal{F}_i[\mathcal{F}_j[R]]]] && \langle \text{by Lemma 3.34 and E-CONG} \rangle
 \end{aligned}$$

□

3.2.4 Duality, Dependency, and Deadlock

The definition of canonical form (Definition 3.33) requires justification. It defines canonical forms as “processes that do not reduce”, which is an easy way to get in trouble and admit processes that are stuck for undesirable reasons such as deadlock as “canonical”.

The section follows the same structure as the corresponding section for CP (§ 2.2.4):

- *Dependency Graph and Deadlock Freedom.*

I define the dependency graph for HCP processes, and prove that HCP processes are deadlock-free because their dependency graph is always essentially acyclic.

- *Adequacy.*

I define when a process is blocked on a set of endpoints, and prove

that an HCP process cannot β -reduce if and only if it is blocked on a set of free endpoints.

The dependency graph is a mixed graph. I informally revisit the relevant definitions. For a detailed discussion, see § A.1.

- A *mixed graph* G has a set of vertices (denoted V_G , ranged over by u, v), a set of edges (denoted E_G), a set of arcs (denoted A_G). Edges are unordered pairs denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$. The set of edges may not contain loops uu . Arcs are ordered pairs denoted by juxtaposition overset with an arrow to indicate the direction, i.e. $\overrightarrow{uv} \triangleq (u, v)$. The set of arcs may not contain loops $\overrightarrow{uu}$.
- For any graph G with vertices $u, v \in V_G$, u is *adjacent* to v when there exists some edge $uv \in E_G$ or some arc $\overrightarrow{uv} \in A_G$.
- A *walk* w is a sequence of pairwise adjacent vertices.
- A *path* p is a walk with no repeated vertices, except possibly the first and last.
- A *cycle* c is a path that begins and ends at the same vertex.
- A walk is *essentially directed* when it contains at least one arc.
- A graph is *essentially acyclic* when if it contains no essentially directed cycles.
- A graph is *strongly connected* when if there exists a path between any two vertices.
- A graph is *connected* when if the graph formed by replacing all arcs with edges is strongly connected.
- The *subgraph of G induced by U* (denoted by $G[U]$) is the graph formed by taking the subset of vertices U and restricting the edges and arcs according to their correctness criteria, i.e. $E_{G[U]} \triangleq E_G \cap \{uv \mid u, v \in U \wedge u \neq v\}$.
- A *component* of a graph is a maximal connected subgraph.
- The *undirected reachability* relation (denoted by $\sim_G$) is the equivalence closure over E_G .
- The *essentially directed reachability* relation (denoted by $<_G$) is the transitive closure over A_G quotiented by $\sim_G$.

The vertices of the dependency graph are endpoint names, which are a proxy for the first action on that endpoint. The edges represent channels, created by either links or name restrictions. The arcs represent dependencies, created by prefixing. For instance, in $a(). b[] . 0$, the process $b[] . 0$ is prefixed with the action $a()$. Hence, the action on b depends on the action on a . (The definition of the dependency graph is unchanged from CP, except to account for the terminated process.)

Definition 3.36 (Dependency Graph). *The shallow dependency graph of a process P , written $\text{Dep}(P)$, is a mixed graph (see § A.1). The process P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$). The shallow dependency*

graph $\text{Dep}(\mathbf{P})$ is defined as:

$$\begin{aligned} V_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \text{fn}(T_i) \\ E_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \{xy \mid T_i = x \leftrightarrow y\} \cup \{x\bar{x} \mid \{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])\} \\ A_{\text{Dep}(\mathbf{P})} &\triangleq \bigcup_{1 \leq i \leq n} \{\bar{x}y \mid x, y \in \text{fn}(T_i), \text{ready}(T_i, x) \wedge \neg \text{ready}(T_i, y)\} \end{aligned}$$

By Lemma 3.22, $E_{\text{Dep}(\mathbf{P})}$ and $A_{\text{Dep}(\mathbf{P})}$ contain no loops. If G is (the subgraph of) some dependency graph, I write $\text{fn}(G)$ for its vertices, i.e. $\text{fn}(G) = V_G$.

The dependency graph gives us *duality* on actions, which is undirected reachability in the dependency graph. (The definition of duality is unchanged from CP.)

Definition 3.37 (Duality). An endpoint x is dual to some endpoint y in $\mathbf{P}$, written $x \sim_{\mathbf{P}} y$, if and only if there exists an undirected path from x to y in $\text{Dep}(\mathbf{P})$.

If $x \sim_{\mathbf{P}} y$, the corresponding path in $\text{Dep}(\mathbf{P})$ may be arbitrarily long, as undirected edges arise from both cuts and links. Consider the process

$$(vx\bar{x})(a \leftrightarrow x \parallel \bar{x} \leftrightarrow b)$$

The duality $a \sim b$ is witnessed by the path $(ax, x\bar{x}, \bar{x}b)$. However, while the paths arising from cuts and links may be arbitrarily long, they must alternate between cut-edges and link-edges and can never branch.

The dependency graph also gives us *dependency* on actions, which is the converse of essentially directed reachability in the dependency graph. (The definition of dependency is unchanged from CP.)

Definition 3.38 (Dependency). An endpoint x depends on some endpoint y , written $x >_{\mathbf{P}} y$, if and only if there exists an essentially directed path from y to x in $\text{Dep}(\mathbf{P})$.

One quirk of using endpoints as a proxy for actions is that the duality and dependency appear to “leak” restricted names, i.e. $\sim_{\mathbf{P}}$ and $>_{\mathbf{P}}$ are not relations over $\text{fn}(\mathbf{P})$, but relations over $\text{fn}(\mathbf{P}) \cup \text{bn}(\mathcal{C}[\cdot])$, where $\mathcal{C}[\cdot]$ is the maximum configuration context of $\mathbf{P}$. However, as stated, these relations should be viewed as relations on the first actions on those endpoints, not the endpoints themselves.

A process is in deadlock if the dependency relation is not antisymmetric, or, equivalently, if there is a cycle in the dependency graph that contains at least one arc.

Definition 3.39 (Deadlock). A process $\mathbf{P}$ is in deadlock, written $\text{deadlock}(\mathbf{P})$, if $\text{Dep}(\mathbf{P})$ contains an essential cycle.

For HCP, the statement that well-typed processes are deadlock free is too weak as an induction hypothesis. Instead, we prove the stronger proposition that (1) the dependency graph is essentially acyclic, and (2) that the components of the dependency graph correspond one-to-one to

the typing environments in the hyper-environment. When specialised to CP, this property becomes deadlock freedom, since every CP process is typed under exactly one typing environment, and its dependency graph is always connected.

Let us start with the base case. If a process is ready, then its dependency graph is essentially acyclic and connected. The latter suffices, since threads are always typed under a single typing environment.

Lemma 3.40. *If P is ready, then $\text{Dep}(P)$ is essentially acyclic and connected.*

Proof. By case analysis on P .

(For the full proof, see § 3.5.)

□

While the statement of deadlock freedom is stronger, the actual proof does not differ significantly from CP. There is only the small additional burden of maintaining the isomorphism between components and typing environments. The interesting case is the case for name restriction, which, as for CP, relies on Lemma A.2, the property that connecting two essentially acyclic graphs with a single edge yields another essentially acyclic graph.

Proposition 3.41. *If $P \vdash \mathcal{G}$, then the dependency graph $\text{Dep}(P)$ is essentially acyclic, and there is an isomorphism f between the typing environments in $\mathcal{G}$ and the components of $\text{Dep}(P)$ that preserves fn, i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. By induction on the derivation of $P \vdash \mathcal{G}$. The case where P is ready follows by Lemma 3.40. The case where P is of the form 0 follows vacuously. The case where P is of the form $P_1 \parallel P_2$ follows by taking the union of the induction hypotheses. The case where P is of the form $(\nu x\bar{x})P'$, is the interesting case. By typing, x and $\bar{x}$ are in distinct typing environments. By f , these typing environments correspond to distinct components of the dependency graph. By Lemma A.2, connecting distinct components with a single edge preserves acyclicity.

(For the full proof, see § 3.5.)

□

Every well-typed HCP process is deadlock-free. This follows immediately from Proposition 3.41 by forgetting the isomorphism.

Corollary 3.42. *If $P \vdash \mathcal{G}$, then $\neg \text{deadlock}(P)$.*

A *blocking action* is an action that blocks a process from making progress. For instance, in the process

$$(\nu x \bar{x})(a(). x[] . 0 \parallel \bar{x}(). P)$$

the action $a()$ is *blocking*. However, not every ready action is *blocking*. The action $\bar{x}()$ is ready, but not blocking. Rather, it is *blocked*: its dual $x[]$ depends on $a()$, so it cannot reduce until $a()$ does.

As I did with dependency, I approximate blocking actions with blocking endpoints. *Blocking endpoints* are the maxima of the dependency relation, or, equivalently, the sources of the dependency graph. The blocking set of a process is the set of all sources of its dependency graph. Every ready action in a process is blocked on one of the endpoints in the blocking set. (The definition of the blocking set is unchanged from CP.)

Definition 3.43 (Blocking Set). *The blocking set of endpoints of a process P , written $\text{blocking}(P)$, is the set of sources of $\text{Dep}(P)$, i.e. $\{x \in V_{\text{Dep}(P)} \mid \nexists y. x >_P y\}$.*

The blocking set is closed under duality.

Lemma 3.44. *If $P \vdash \Gamma$ and $x \sim_P y$, then $x \in \text{blocking}(P) \implies y \in \text{blocking}(P)$.*

Each endpoint in the blocking set corresponds to a ready action.

Lemma 3.45. *If $P \vdash \Gamma$ and $x \in \text{blocking}(P)$, then $P = \mathcal{C}[T]$ and $\text{ready}(T, x)$.*

Due to the dualities generated by links, the blocking set may contain more endpoints than necessary. For instance, the blocking set of the process

$$(\nu x \bar{x})(x \leftrightarrow a \parallel \bar{x}(). P)$$

is $\{x, \bar{x}, a\}$. An action in P is blocked on all of these endpoints. However, I want to be able to say that any action in is blocked on a free name, and, in this case, the set $\{a\}$ suffices. A process is blocked on a set of endpoints if any action is blocked on at least one endpoint in that set. (The definition of blocking is unchanged from CP.)

Definition 3.46 (Blocked). *A process P is blocked on a set of endpoints X , written $\text{blocked}(P, X)$, if closing X under duality yields the blocking set $\text{blocking}(P)$, i.e. if $x \in X$ and $x \sim_P y$, then $y \in \text{blocked}(P)$.*

Any process is blocked on its blocking set.

Lemma 3.47. *If $P \vdash \Gamma$, then $\text{blocked}(P, \text{blocking}(P))$.*

If a process is blocked on some set of endpoints, it is blocked on the set formed by replacing any endpoint in that set with its dual.

Lemma 3.48. *If $P \vdash \Gamma$ and $x \sim_P y$, then $\text{blocked}(P, X) \implies \text{blocked}(P, X\{y/x\})$.*

If a process cannot β -reduce, then it is blocked on some set of free names.

Proposition 3.49. *If $P \vdash \Gamma$, then $P \not\rightarrow_\beta \iff \exists A \subseteq \text{fn}(P). \text{blocked}(P, A)$.*

Proof. There are two cases:

- Case $(\Rightarrow)$.

By contradiction. Assume $x \in \text{blocking}(P)$ and $\nexists a \in \text{fn}(P). x \sim_P a$.

There are two cases:

- If $x \in \text{fn}(P)$, then $x \sim_P x$.
- The process $P = \mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$).

If $x \in \text{bn}(\mathcal{C}[\cdot])$, then there exists some $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$.

There are two cases:

- * If $\bar{x} \in \text{blocking}(P)$, there exist processes T_i and T_j that are ready to act on dual endpoints. By Lemma 3.34, P is not β -free.
- * If $\bar{x} \notin \text{blocking}(P)$, there exists some y such that $\bar{x} >_P y$. By definition, $x\bar{x} \in E_{\text{Dep}(P)}$. Hence, $x >_P y$ and $x \notin \text{blocking}(P)$.

- Case $(\Leftarrow)$.

By contradiction.

Assume $P \rightarrow_\beta$. By inversion, there exist some T_i and T_j (for $1 \leq i, j \leq n$) that are ready to act on dual endpoints x and $\bar{x}$.

By definition, x and $\bar{x}$ only have outgoing arcs in $\text{Dep}(P)$, and the only edge connected to either is $x\bar{x}$. Hence, $\{x, \bar{x}\} \subseteq \text{blocking}(P)$ and $\nexists a \in \text{fn}(P). x \sim_P a \vee \bar{x} \sim_P a$.

□

Corollary 3.50. *If $P \vdash \Gamma$, then $\text{canonical}(P) \implies \text{blocking}(P) \subseteq \text{fn}(P)$.*

Proof. As $\text{canonical}(P)$, $P \not\rightarrow_a$ and $P \not\rightarrow_\beta$. By Proposition 3.49, there exists some $A \subseteq \text{fn}(P)$ such that $\text{blocked}(P, A)$. By definition, $A \subseteq \text{blocking}(P)$. It remains to show that $\text{blocking}(P) \subseteq A$.

By contradiction.

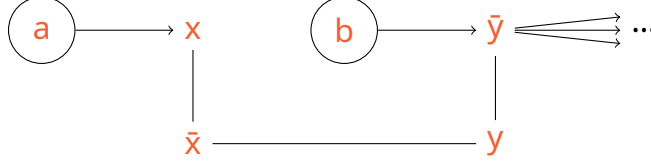
Assume $x \in \text{blocking}(P)$ and $x \notin A$. By definition, there exists some $a \in A$ such that $x \sim_P a$. The duality $x \sim_P a$ corresponds to some undirected path $p_{xa} = (x, \dots, a)$ in $\text{Dep}(P)$, which must contain at least one edge that connects some *bound* name y to some *free* name b , say, yb . By definition, any edge in $\text{Dep}(P)$ generated by a cut connects two *bound* names. Therefore, yb must be generated by a link. By Lemma 3.45, $P = \mathcal{C}[y \leftrightarrow b]$. Hence, $P \rightarrow_a$.

□

Unfortunately, as with CP, “blocked on free endpoints” does not characterise canonical forms, as $\text{blocking}(P) \subseteq \text{fn}(P) \not\Rightarrow P \not\rightarrow_\alpha$. For instance, the process

$$(\nu x \bar{x})(a().x[] . 0 \parallel (\nu y \bar{y})(\bar{x} \leftrightarrow y \parallel b().\bar{y}().P))$$

can α -reduce. Its dependency graph, with blocking endpoints circled, is



In conclusion, my definition of canonical form (Definition 3.33) is *adequate*: any process in canonical form is blocked on some set of free endpoints. Due to the behaviour of links, “blocked on free endpoints” is insufficient to characterise canonical forms. As this would be a desirable property to have, I consider alternative semantics for the link construct in § 3.3. I have not adopted any of these alternatives as standard to maintain backwards compatibility with the work based on Wadler’s CP.

3.2.5 Connection and Disentanglement

In this section, I formalise the notion of the *connection graph* of a process, and prove *disentanglement*, the property that any HCP process can be rewritten into the parallel composition of CP-like processes, where I use “CP-like” to mean processes that would be syntactically well-formed and well-typed CP processes, i.e. each name restriction is followed by its corresponding parallel composition, each send action is followed by its corresponding parallel composition in the correct order, and each close action is followed by the terminated process.

The first part of this section follows the same structure as the corresponding section for CP (§ 2.2.5):

- *Connection Forest.*
I define the connection graph for HCP processes, and prove that the connection graph is always a forest.
- *Right-Branching Forest Form.*
I define the equivalent of right-branching form for HCP, which is right-branching *forest* form, and prove that any process can be rewritten to right-branching forest form.

The second part of this section proves disentanglement for HCP. In the introduction, I defined disentanglement as a function that converts HCLL proofs into a sequence of CLL proofs:

$$\begin{array}{c} p \\ \vdots \\ \vdash \Gamma_1 \parallel \dots \parallel \Gamma_n \end{array} \quad \Longrightarrow \quad \begin{array}{c} p_1 \quad \dots \quad p_n \\ \vdots \quad \dots \quad \vdots \\ \vdash \Gamma_1 \quad \dots \quad \vdash \Gamma_n \end{array}$$

For the process calculus HCP, I prefer a slightly stronger property. Disentanglement should convert HCP processes to the parallel composition of CP-like processes *and* should be justified by the structural congruence. For any well-typed process $P \vdash \mathcal{G}^k$:

- If $k = 0$, then $P \equiv 0$ and $\mathcal{G}^0 = \circ$.
- If $k \geq 1$, then $P \equiv P_1 \parallel \dots \parallel P_k$ and $\mathcal{G}^k = \Gamma_1 \parallel \dots \parallel \Gamma_k$ such that $P_i \vdash \Gamma_i$ (for $1 \leq i \leq k$) and each P_i is CP-like.

The conversion to right-branching forest form satisfies part of this definition. It converts an HCP process to a sequence of processes that are typed under a single typing environment, *and* it is justified by the structural congruence. However, the conversion is *shallow*. It only rewrites the maximum configuration context. Hence, the resulting processes are *shallowly* CP-like. Only the name restrictions and parallel compositions in their maximum configuration contexts are arranged as CP cuts.

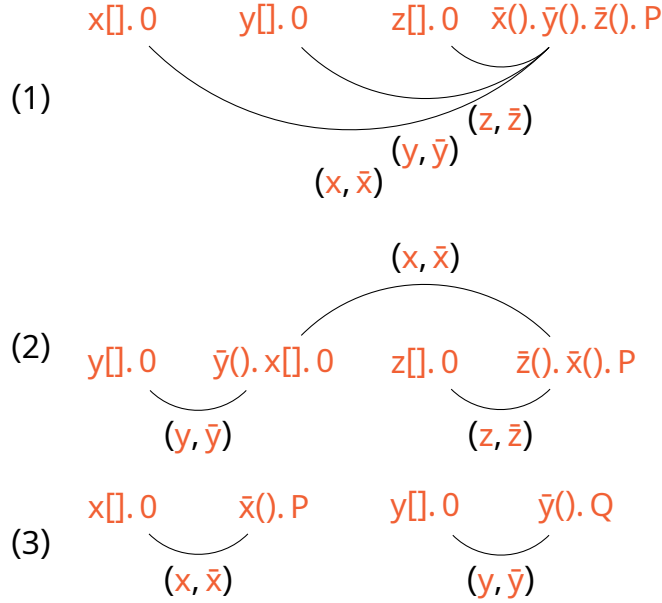
I define disentanglement by iterating the conversion to right-branching forest form, which arranges *all* name restrictions and parallel compositions as CP cuts, and subsequently arranging the continuations of send and close actions to match those of CP send and close actions. (I defer the proof that disentangled processes are CP-like to § 3.2.7, where I discuss the translation between HCP processes and multisets of CP processes.)

Let us revisit connection graphs by examining three example processes:

- (1) $(vx\bar{x})(vy\bar{y})(vz\bar{z})(x[].0 \parallel y[].0 \parallel z[].0 \parallel \bar{x}().\bar{y}().\bar{z}().P)$
- (2) $(vx\bar{x})(vy\bar{y})(vz\bar{z})(y[].0 \parallel \bar{y}().x[].0 \parallel z[].0 \parallel \bar{z}().\bar{x}().P)$
- (3) $(vx\bar{x})(vy\bar{y})(x[].0 \parallel \bar{x}().P \parallel y[].0 \parallel \bar{y}().Q)$

The connection graph of a process is the graph formed of all ready sub-processes and the channels that connect them. Whereas CP's connection graphs are trees, HCP's connection graphs are forests. The connection graphs for processes (1) and (2) are equivalent to those for the example processes discussed in § 2.2.5. Both are fully connected, consisting of four threads, connected by three channels. Process (3) is *not* fully connected.

It consists of four processes, connected in pairs by two channels.



We can use the connection graph to rewrite any process into right-branching forest form, which is the parallel composition of a sequence of processes in right-branching tree form, which is more or less CP's right-branching form:

$$(\nu x_1^1 \bar{x}_1^1)(P_1^1 \parallel \dots (\nu x_n^1 \bar{x}_n^1)(P_n^1 \parallel P_{n+1}^1) \dots) \parallel \dots \parallel (\nu x_1^k \bar{x}_1^k)(P_1^k \parallel \dots (\nu x_n^k \bar{x}_n^k)(P_n^k \parallel P_{n+1}^k) \dots)$$

As mentioned, CP's right-branching form is often used to write nice and concise proofs. Unfortunately, HCP's right-branching forest form is a bit too verbose to fill the same niche. Worse, the presentation above does not communicate the case where $k = 0$ and the process is 0 .

The procedure to convert processes to right-branching forest form picks a tree from the connection graph, moves *all* the corresponding name restrictions and threads to the top-level, arranges them in right-branching tree form, removes the tree, and continues until the graph is empty.

The procedure to convert processes to right-branching tree form is the same as the procedure for CP. It picks a leaf from the tree, moves the corresponding name restriction and thread to the topmost, leftmost position, removes the leaf, and continues until all of the tree is empty.

A process may have multiple different right-branching forest forms. The following are one possible right-branching forest form for each of the example processes above:

- (1) $(\nu x \bar{x})(x[] . 0 \parallel (\nu y \bar{y})(y[] . 0 \parallel (\nu z \bar{z})(z[] . 0 \parallel \bar{x}(). \bar{y}(). \bar{z}(). P)))$
- (2) $(\nu y \bar{y})(y[] . 0 \parallel (\nu x \bar{x})(\bar{y}(). x[] . 0 \parallel (\nu z \bar{z})(z[] . 0 \parallel \bar{z}(). \bar{x}(). P)))$
- (3) $(\nu x \bar{x})(x[] . 0 \parallel \bar{x}(). P) \parallel (\nu y \bar{y})(y[] . 0 \parallel \bar{y}(). Q)$

As in § 2.2.5, the definition of connection graph is *shallow*, rather than deep, as we only account for the connections up to the maximum configuration context. However, in the second part of this section, I will demonstrate that we can use the shallow connection graph to reason about the deep connection structure of a process.

The connection graph is a undirected edge-labelled graph. I informally revisit the relevant definitions. For a detailed discussion, see § A.1.

- A *undirected edge-labelled graph* G has a set of vertices (denoted V_G , ranged over by u, v), a set of edges (denoted E_G), a set of edge labels (denoted $\mathcal{L}_G$), and an edge-labeling function (denoted ℓ_G) that assigns labels to edges. Edges are unordered pairs denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$. The set of edges may not contain loops uu . (It suffices to define V_G and ℓ_G , since $E_G \triangleq \text{dom}(\ell_G)$ and $\mathcal{L}_G \triangleq \text{cod}(\ell_G)$.)
- Two vertices $u, v \in V_G$ are *adjacent* when there exists an edge $uv \in E_G$.
- A *walk* w is a sequence of pairwise adjacent vertices.
- A *path* p is a walk with no repeated vertices, except possibly the first and last.
- A *cycle* c is a path that begins and ends at the same vertex.
- The *subgraph of G induced by U* (denoted by $G[U]$) is the graph formed by taking the subset of vertices U and restricting the edges, edge labels, and edge-labelling function according to their correctness criteria, i.e. $E_{G[U]} \triangleq E_G \cap \{uv \mid u, v \in U \wedge u \neq v\}$.
- A graph is *acyclic* when it does not contain a cycle.
- A graph is *connected* when there is a path between any two vertices.
- A *component* C of a graph is a maximal connected subgraph.
- A *tree* T is a graph that is connected and acyclic.
- A *forest* F is a graph whose components are trees.

(The definition of the connection graph is unchanged from CP, except to account for the terminated process.)

Definition 3.51 (Connection Graph). *The shallow connection graph of a well-typed process P , written $\text{Con}(P)$, is an undirected edge-labelled graph (see § A.1) where the vertices are threads, the edges are the channels that connect those threads, and the edges are labelled by unordered pairs of their endpoints. The process P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$). The shallow connection graph $\text{Con}(P)$ is defined as:*

$$\begin{aligned} V_{\text{Con}(P)} &\triangleq \{T_1, \dots, T_n\} \\ \ell_{\text{Con}(P)} &\triangleq \{T_i T_j \mapsto (\bar{x}, \bar{x}) \mid T_i, T_j \in V_{\text{Con}(P)}, \{\bar{x}, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot]), x \in T_i \wedge \bar{x} \in T_j\} \end{aligned}$$

By Lemma 2.21, $E_{\text{Con}(P)}$ contains no loops. By Proposition 2.18, $\ell_{\text{Con}(P)}$ is a function. If G is a subgraph of some connection graph, I write $\text{fn}(G)$ for the free names in the vertices of G , i.e. $\text{fn}(G) \triangleq \bigcup_{P \in V_G} \text{fn}(P)$.

In HCP, the connection graph of a process is always a forest. However, that statement by itself is too weak to prove by induction. Instead, we

prove the stronger proposition that (1) the connection graph is a forest, and (2) that the components in the connection graph correspond one-to-one to the typing environments in the hyper-environment. When specialised to CP, this property becomes the property that every connection graph is a tree, since every CP process is typed under exactly one typing environment, and its connection graph is always connected.

The proof does not differ significantly from CP. There is only the small additional burden of maintaining the isomorphism between component and typing environments. The interesting case is the case for name restriction, which, as for CP, relies on the property that connecting two trees with a single edge yields another tree.

Proposition 3.52. *If $P \vdash \mathcal{G}$, then $\text{Con}(P)$ is a forest, and there is an isomorphism f between the typing environments in $\mathcal{G}$ and the trees of $\text{Con}(P)$ that preserves fn , i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. By induction on the maximum configuration context of P and inversion on P and its typing derivation. The interesting case is for name restriction. The endpoints connected by the name restriction occur in different typing environment. Hence, by the isomorphism f , those endpoints occur in disjoint trees in the connection graph and, by Lemma A.1, the connection graph for the result, formed by connecting those trees with the single edge arising from the name restriction, remains a forest.

(For the full proof, see § 3.5.)

□

A process is in right-branching forest form when it is either 0 or it is the parallel composition of processes in right-branching tree form.

Definition 3.53 (Right-branching Tree Form). *A process P is in right-branching tree form if P is ready or if P is of the form (for some $n \geq 1$)*

$$(v x_1 \bar{x}_1)(T_1 \parallel \cdots (v x_n \bar{x}_n)(T_n \parallel T_{n+1}) \cdots)$$

A process is in right-branching tree form if it is a right-branching list of CP cuts connecting threads.

(The definition of right-branching tree form is unchanged, except in name, from CP's definition of right-branching form.)

Definition 3.54 (Right-branching Forest Form). *A process P is in right-branching forest form if P is of the form 0 , or if P is of the form $P_1 \parallel \cdots \parallel P_n$ (for some $n \geq 1$) such that (for $1 \leq i \leq n$) each P_i is in right-branching tree form.*

Any well-typed process can be rewritten to right-branching forest form. The proof is given by induction on the structure of the hypersequent, and

is decomposed into two lemmas, which correspond to the two cases of the induction:

- If $P \vdash \circ$, then P can be rewritten to the terminated process.
- If $P \vdash \mathcal{G} \parallel \Gamma$, then P can be rewritten to pull one process in right-branching tree form out to the top level, using a procedure similar to the one used for converting CP processes to right-branching form.

Lemma 3.55. *If $P \vdash \circ$, then $P \equiv 0$.*

Proof. Let G be $\text{Con}(P)$. By Proposition 3.52, there is an isomorphism f between the typing environments in $\circ$ and the trees of G that preserves fn , i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$. As there are no typing environments in $\circ$, f is the empty function, and G is the null graph. Let $\mathcal{C}^n[\cdot]$ be maximal for P such that $P = \mathcal{C}^n[P_1, \dots, P_n]$ (for some $n \geq 0$). By definition, $V_G = \{P_1, \dots, P_n\}$. As G is the null graph, $V_G = \emptyset$. Hence, $n = 0$, i.e. the maximum configuration context $\mathcal{C}^0[\cdot]$ contains zero holes. By induction on $\mathcal{C}^0[\cdot]$, the process P is equal to the parallel composition of terminated processes, and is equivalent to 0 by repeated application of SC-PARNIL. □

Lemma 3.56. *If $P \vdash \mathcal{G} \parallel \Gamma$, then there exist processes Q and R such that $Q \vdash \mathcal{G}$, $R \vdash \Gamma$, $P \equiv Q \parallel R$, and R is in right-branching tree form.*

Proof. The proof proceeds by picking the tree corresponding to Γ from the connection forest, using the isomorphism constructed by Proposition 3.52, and then iteratively picking its leaves and rewriting the corresponding process to right-branching form, as in Proposition 2.51.

(For the full proof, see § 3.5.) □

Proposition 3.57. *If $P \vdash \mathcal{G}$, then there exists a process Q , such that $P \equiv Q$, and Q is in right-branching forest form, and there is an isomorphism f between the processes in right-branching tree form in Q and the typing environments in $\mathcal{G}$ that preserves fn , i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. The proof proceeds by iteratively picking typing environments in the hyper-environment and moving the corresponding process to the top-level in right-branching tree form using Lemma 3.56.

(For the full proof, see § 3.5.) □

Connection graphs are an alternative representation for processes, with the interesting property that they represent the maximum configuration context of a process without any spurious ambiguity. Connection

graphs correspond to shallow proof nets for HCLL. (For a more detailed discussion of the correspondence between connection graphs and proof nets, see § 2.2.5.)

Proposition 3.58. *If P is well-typed, then $P \equiv Q \iff \text{Con}(P) = \text{Con}(Q)$.*

Proof. There are two cases:

- Case $(\Rightarrow)$.

By induction on the proof of the structural congruence $P \equiv Q$. The cases for reflexivity, transitivity, symmetry, and SC-CONG follow by induction and those same properties of equality. The remaining cases follow immediately.

- Case $(\Leftarrow)$.

Let $\text{Proc}(T)$ be the set of processes in right-branching forest form obtained from the connection graph G of a well-typed process by Proposition 3.57. (This is a set because Proposition 3.57 defines a non-deterministic procedure.)

Pick any $R \in \text{Proc}(\text{Con}(P))$. As $\text{Con}(P) = \text{Con}(Q)$, $\text{Proc}(\text{Con}(P)) = \text{Proc}(\text{Con}(Q))$. Hence, $R \in \text{Proc}(\text{Con}(Q))$. By definition, $P \equiv R$ and $Q \equiv R$. Hence, $P \equiv Q$.

□

An HCP process is disentangled when it is the terminated process, or the parallel composition of a sequence of CP-like processes. The formal definition does not make explicit reference to CP syntax, and I defer the proof disentangled HCP processes correspond to CP processes to § 3.2.7.

Definition 3.59 (Disentangled). *A process P is disentangled when every terminated process matches the form of a CP close, and every parallel composition matches the form of a CP cut or CP send, or is at the top-level.*

Formally, P is disentangled when the following conditions hold:

1. If P is of the form $Q[P_1 \parallel P_2]$, then one of the following holds:
 - a. $Q[\cdot]$ is of the form $R[(\nu x \bar{x})\square]$ such that $x \in \text{fn}(P_1)$ and $\bar{x} \in \text{fn}(P_2)$.
 - b. $Q[\cdot]$ is of the form $R[x[y].\square]$ such that $y \in \text{fn}(P_1)$ and $x \in \text{fn}(P_2)$.
 - c. $Q[\cdot]$ is of the form $\mathcal{E}$ such that $\text{bn}(\mathcal{E}) = \emptyset$.
2. If P is of the form $Q[0]$, then one of the following holds:
 - a. $Q[\cdot]$ is of the form $\square$.
 - b. $Q[\cdot]$ is of the form $R[x[].\square]$.

The conditions of disentangled form in Definition 3.59 are sufficient to guarantee that the process is CP-like.

Lemma 3.60. *If $P \vdash \mathcal{G}$ and P is disentangled, then:*

1. $\mathcal{G}$ is of the form $\circ$ $\implies P$ is of the form 0
2. $\mathcal{G}$ is of the form $\mathcal{G} \parallel \Gamma$ $\implies P$ is of the form $Q \parallel R$
3. P is of the form $Q[(\nu x\bar{x})R]$ $\implies R$ is of the form $R_1 \mathrel{x \parallel_x} R_2$
4. P is of the form $Q[x[y].R]$ $\implies R$ is of the form $R_1 \mathrel{y \parallel_x} R_2$
5. P is of the form $Q[x[] . R]$ $\implies R$ is of the form 0

Proof. There are five cases:

1. By induction on the structure of P . If P is of the form 0 , the result follows immediately. If P is of the form $Q \parallel R$, the induction hypotheses give us $Q = 0$ and $R = 0$, which contradicts condition (2) of disentangled form. The remaining cases, where P is ready or P is of the form $(\nu x\bar{x})P'$ are impossible, by inversion on the typing derivation.
2. By induction on the structure of P . If P is of the form $Q \parallel R$, the result follows immediately. If P is of the form $(\nu x\bar{x})P'$, the induction hypothesis gives us that P' is of the form $Q' \parallel R'$, which contradicts condition (1) of disentangled form. The remaining cases, where P is ready or P is of the form 0 are impossible, by inversion on the typing derivation.
3. By case (2) and condition (1a) of disentangled form.
4. By case (2) and condition (1b) of disentangled form.
5. By case (1).

□

Every HCP process can be disentangled.

Right-branching forest form is useful for disentangling processes, but not quite sufficient. If we convert every sub-process to right-branching forest form, the resulting process is *nearly* disentangled. Every terminated process matches the form of a CP close, and every parallel composition is at the top-level, matches a CP cut, or *nearly* matches a CP send. The continuation of a send action must be the parallel composition that splits the corresponding endpoints, but the processes need not be in the correct order, i.e. we may have $x[y]. Q \parallel P$, where Q handles x and P handles y .

To disentangle a process, we convert every sub-process to right-branching forest form, and take special care to correct the order in the continuation of send actions.

Proposition 3.61. *If $P \vdash \mathcal{G}$, then there exists some Q such that $P \equiv Q$ and Q is disentangled.*

Proof. By induction on the derivation of $P \vdash \mathcal{G}$ and inversion on P .

If P is ready, the result follows by induction. Most cases can be handled uniformly, but send, offer and the absurd offer are handled separately.

- Case P is of the form $x[y].P'$.

By induction, $P' \equiv Q'$ for some disentangled process Q' . By Lemma 3.60 (2), Q' is of the form $Q'_1 \parallel Q'_2$. By condition (2) of disentangled form, $Q'_1 \neq 0$ and $Q'_2 \neq 0$. There are two cases:

- Subcase $y \in \text{fn}(Q'_1)$ and $x \in \text{fn}(Q'_2)$.

Let Q be $x[y].Q'_1 \parallel Q'_2$. The result follows.

- Subcase $x \in \text{fn}(Q'_1)$ and $y \in \text{fn}(Q'_2)$.

Let Q be $x[y].Q'_2 \parallel Q'_1$. The result follows by SC-PARCOMM.

- Case P is of the form $x(y).P'$, $x[\cdot].P'$, $x().P'$, $x \triangleleft \text{inl}.P'$, or $x \triangleleft \text{inr}.P'$.

By induction, $P' \equiv Q'$ for some disentangled process Q' . Let Q be $x(y).Q'$, $x[\cdot].Q'$, $x().Q'$, $x \triangleleft \text{inl}.Q'$, or $x \triangleleft \text{inr}.Q'$, respectively. The result follows.

- Case P is of the form $x \triangleright \{\text{inl}: P'_1; \text{inr}: P'_2\}$.

By induction on P'_1 , $P'_1 \equiv Q'_1$ for some disentangled process Q'_1 . By induction on P'_2 , $P'_2 \equiv Q'_2$ for some disentangled process Q'_2 . Let Q be $x \triangleright \{\text{inl}: Q'_1; \text{inr}: Q'_2\}$. The result follows.

- Case P is of the form $x \leftrightarrow y$ or $x \not\leq N$.

The result follows immediately.

Otherwise, the result follows by converting P to right-branching forest form, taking the maximum configuration context of the resulting process, and converting each ready subprocess by induction.

- Case P is of the form $(\nu x\bar{x})P'$, $P_1 \parallel P_2$, or 0 .

By Proposition 3.57, there exists some Q' such that $P \equiv Q'$ and Q' is in right-branching forest form. By case analysis on Q' :

- If Q' is of the form 0 , the result follows immediately.
- Otherwise, let $\mathcal{C}^n[\cdot]$ be the maximum configuration context of Q' such that $Q' = \mathcal{C}^n[Q'_1, \dots, Q'_n]$ (for some $n \geq 1$) and (for $1 \leq i \leq n$) each Q'_i is ready. By induction, there exist processes $Q_1, \dots, Q_n$ such that (for $1 \leq i \leq n$) each $Q'_i \equiv Q_i$ and Q_i is disentangled. Let Q be $\mathcal{C}^n[Q_1, \dots, Q_n]$. By transitivity and congruence, $P \equiv Q$. By induction on the structure of $\mathcal{C}^n[\cdot]$ and inversion on the fact that Q' is in right-branching forest form, Q is disentangled. The result follows.

□

Disentanglement is the procedure defined by the proof of Proposition 3.61. The procedure is non-deterministic and follows the outline in the introduction to this section. (The non-determinism arises from the

arbitrary choice of tree in Proposition 3.61 and the arbitrary choice of leaf in Lemma 3.56.)

The result is a function that converts a process to a set of disentangled processes, all of which are equivalent to the original process and to each other by structural congruence.

Definition 3.62 (Disentanglement). *The disentanglement of P , written $\{P\}$, is the set of processes obtained from P by Proposition 3.61.*

All elements of $\{P\}$ are equivalent to each other and to P under link-preserving structural congruence. I use $\{\cdot\}$ as functional under structural congruence, and write $\{P\} \triangleq Q$ to mean Q is an arbitrary element of $\{P\}$. I extend $\{\cdot\}$ to configuration and evaluation contexts by preserving holes and otherwise acting as on the corresponding process terms.

The disentanglement $\{P\}$ does not contain *all* disentangled processes equivalent to P , only those which are in right-branching form.

Disentanglement distributes over maximal evaluation contexts.

Lemma 3.63. *If $\mathcal{C}[T] \vdash \mathcal{G}$, then $\{\mathcal{C}[T]\} = \{\mathcal{C}'[T'] \mid \mathcal{C}' \in \{\mathcal{C}\}, T' \in \{T\}\}$.*

Disentanglement distributes over the maximum configuration context.

Lemma 3.64. *If $\mathcal{C}[T_1, \dots, T_n] \vdash \mathcal{G}$, then*

$$\{\mathcal{C}[T_1, \dots, T_n]\} = \{\mathcal{C}'[T'_{\rho(1)}, \dots, T'_{\rho(n)}] \mid \exists p. \mathcal{C}'[\cdot] \in \{\mathcal{C}[\cdot]\}, T'_1 \in \{T_1\}, \dots, T'_n \in \{T_n\}\}$$

where p is a permutation on the indices $[1, n]$.

The set of right-branching forms of a process is closed under link-preserving shallow structural congruence, as link-preserving shallow structural congruence preserves the connection graph. Likewise, the disentanglement is closed under any link-preserving structural congruence.

Lemma 3.65. *If $P \vdash \mathcal{G}$ and $P \triangleq Q$, then $\{P\} = \{Q\}$.*

More importantly, disentanglement is closed under CP's structural congruence, as we will see in § 3.2.7. First, however, we must pause our discussion of disentanglement to introduce Multiset CP, the calculus of multisets of CP processes.

3.2.6 Multiset CP

In this section, I define Multiset CP. At several points in this chapter, I have mentioned that an HCP process corresponds to a multiset of CP processes. However, I have found that it is tedious to relate HCP directly to multisets of CP processes, and much clearer to relate it to another session-typed processes calculus which captures that notion. Multiset CP is that calculus. Its processes are multisets of unconnected CP processes,

presented using process notation. Its processes are typed under hyper-environments, which are multisets of CP typing environments. Its typing derivation, structural congruence, and reduction relations are all some sort of pointwise lifting of the corresponding CP relations. (I will not abbreviate Multiset CP, and continue to write it in full, to avoid confusion with Multiparty CP [MCP, Carbone et al., 2016].)

To avoid syntax highlighting whiplash between sections, the terms and types of *both* Multiset CP and CP are printed in *pink* and *green*, respectively, both are rendered in an italicised font with serif, and any relations, such as typing and reduction, are marked by subscript “ $\mathcal{C}$ ” or “ $\mathcal{C}$ ”, respectively.

The processes of Multiset CP are the parallel composition of CP processes, as defined by the following grammar:

$$\mathcal{P}, \mathcal{Q}, \mathcal{R} ::= P \mid 0 \mid \mathcal{P} \parallel \mathcal{Q}$$

The metavariables $\mathcal{P}$, $\mathcal{Q}$, and $\mathcal{R}$ refer to Multiset CP processes, and the metavariables P , Q , and R refer to CP processes, as defined in § 2.1. This matches the convention to use the calligraphic font for multisets.

Multiset CP’s parallel composition and terminated process may only occur at the top-level. The syntax for CP processes is not permitted to recurse back into the syntax for Multiset CP processes. The parallel composition and terminated process that are syntactically part of CP’s cut, send, and close constructs are distinct from Multiset CP’s parallel composition and terminated process.

Multiset CP has no top-level name restriction. Hence, there is no possibility to connect the individual CP processes, and no possibility for them to communicate amongst each other.

The processes of Multiset CP are equivalent up to structural congruence. Concretely, the individual CP processes—i.e. the elements of the multiset—are equivalent up to CP’s structural congruence, as defined in § 2.1, and Multiset CP processes—i.e. the multisets themselves—are equivalent up to reordering.

$$\begin{array}{lll}
 \mathcal{P} \parallel 0 & \equiv_{\mathcal{C}} \mathcal{P} & \text{SC-PARNIL} \\
 \mathcal{P} \parallel \mathcal{Q} & \equiv_{\mathcal{C}} \mathcal{Q} \parallel \mathcal{P} & \text{SC-PARCOMM} \\
 \mathcal{P} \parallel (\mathcal{Q} \parallel \mathcal{R}) & \equiv_{\mathcal{C}} (\mathcal{P} \parallel \mathcal{Q}) \parallel \mathcal{R} & \text{SC-PARASSOC} \\
 \\
 \text{SC-EMBED} & \text{SC-PARCONG} & \\
 \frac{P \equiv_{\mathcal{C}} Q}{P \equiv_{\mathcal{C}} Q} & \frac{\mathcal{P} \equiv_{\mathcal{C}} \mathcal{P}' \quad \mathcal{Q} \equiv_{\mathcal{C}} \mathcal{Q}'}{\mathcal{P} \parallel \mathcal{Q} \equiv_{\mathcal{C}} \mathcal{P}' \parallel \mathcal{Q}'} &
 \end{array}$$

In essence, the rules SC-PARNIL, SC-PARCOMM, and SC-PARASSOC ensure the process syntax defines a multiset, and the rules SC-EMBED and SC-PARCONG define the universal pointwise lifting of CP’s structural congruence to multisets of equal length. I likewise lift the restricted

versions of structural congruence— $\stackrel{s}{\equiv}_e$, $\stackrel{l}{\equiv}_e$, $\stackrel{ls}{\equiv}_e$, and $\stackrel{p}{\equiv}_e$ —to Multiset CP (see § 2.2.1 for details).

Reduction for Multiset CP processes is the existential pointwise lifting of CP's reduction, i.e. if $\mathcal{P} \rightarrow_e \mathcal{Q}$, then one of the processes $P \in \mathcal{P}$ reduces to one of the processes $Q \in \mathcal{Q}$, and all the other processes are preserved up to structural congruence.

$$\begin{array}{c}
\text{E-EMBED} \quad \frac{P \rightarrow_c Q}{P \rightarrow_e Q} \qquad \text{E-EQUIV} \quad \frac{\mathcal{P} \equiv_e \mathcal{P}' \quad \mathcal{P}' \rightarrow_e \mathcal{Q}' \quad \mathcal{Q}' \equiv_e \mathcal{Q}}{\mathcal{P} \rightarrow_e \mathcal{Q}} \\
\\
\text{E-PARCONG}_1 \quad \frac{\mathcal{P} \rightarrow_e \mathcal{P}'}{\mathcal{P} \parallel \mathcal{Q} \rightarrow_e \mathcal{P}' \parallel \mathcal{Q}} \qquad \text{E-PARCONG}_2 \quad \frac{\mathcal{Q} \rightarrow_e \mathcal{Q}'}{\mathcal{P} \parallel \mathcal{Q} \rightarrow_e \mathcal{P} \parallel \mathcal{Q}'}
\end{array}$$

Finally, the typing judgement for Multiset CP uses hyper-environments. The definition of hyper-environments follows the definition for HCP. Recall that hyper-environments are multisets of CP typing environments with the additional restriction that endpoint names must be unique across all typing environments.

$$\mathcal{G}, \mathcal{H} ::= \circ \mid \mathcal{G} \parallel \mathcal{I}$$

The typing judgement itself is defined, similar to the structural congruence, as the universal pointwise lifting of CP's typing judgement to multisets of equal size.

$$\begin{array}{c}
\text{T-EMBED} \quad \frac{P \vdash_c \mathcal{I}}{P \vdash_e \mathcal{I}} \qquad \text{T-PAR} \quad \frac{}{0 \vdash_e \circ} \qquad \text{T-HALT} \quad \frac{\mathcal{P} \vdash_e \mathcal{G} \quad \mathcal{Q} \vdash_e \mathcal{H}}{\mathcal{P} \parallel \mathcal{Q} \vdash_e \mathcal{G} \parallel \mathcal{H}}
\end{array}$$

The metatheory for Multiset CP can be constructed from the metatheory for CP, borrowing the occasional bit from Hypersequent CP. (I only discuss preservation and progress, but these should suffice to reassure the reader of that claim.)

Preservation follows from CP's preservation proofs (Lemma 2.24 and Proposition 2.27), and the additional cases for the new rules follow those from HCP's preservation proofs (Lemma 3.27 and Proposition 3.30).

Lemma 3.66. *If $\mathcal{P} \equiv_e \mathcal{Q}$, then $\mathcal{P} \vdash_e \mathcal{G}$ if and only if $\mathcal{Q} \vdash_e \mathcal{G}$.*

Proposition 3.67 (Preservation). *If $\mathcal{P} \vdash_e \mathcal{G}$ and $\mathcal{P} \rightarrow_e \mathcal{Q}$, then $\mathcal{Q} \vdash_e \mathcal{G}$.*

Canonical form and progress are lifted directly from CP. A Multiset CP process is in canonical form when each CP process is in canonical form, and progress follows by pointwise application of CP's progress (Proposition 2.32).

Definition 3.68 (Canonical Form). *A process $\mathcal{P}$ is in canonical form when each CP process $P \in \mathcal{P}$ is in CP's canonical form (Definition 2.30).*

Proposition 3.69 (Progress). *If $\mathcal{P} \vdash_{\mathcal{C}} \mathcal{G}$, then either $\mathcal{P}$ is in canonical form, or there exists some $\mathcal{Q}$ such that $\mathcal{P} \rightarrow_{\mathcal{C}} \mathcal{Q}$.*

I sketched the correspondence between Multiset CP and CP in the introduction to this section. Propositions 3.70 to 3.72 formalise it. The proof of each proposition follows by induction. I believe the verbosity of these statements illustrates the advantage of working with Multiset CP over working directly with multisets of CP processes.

Multiset CP's typing relation is the universal pointwise lifting of CP's typing relation, i.e. a Multiset CP process is well-typed if each individual CP process is well-typed under one of the typing environments.

Proposition 3.70. *If $\mathcal{P} \vdash_{\mathcal{C}} \mathcal{G}$, then there is an isomorphism $f : \mathcal{P} \rightarrow \mathcal{G}$ between the elements of the multisets $\mathcal{P}$ and $\mathcal{G}$ such that $P \vdash_{\mathcal{C}} f(P)$ for every $P \in \mathcal{P}$ and $f^{-1}(T) \vdash_{\mathcal{C}} T$ for every $T \in \mathcal{G}$.*

Multiset CP's structural congruence is the universal pointwise lifting of the CP's structural congruence, i.e. one Multiset CP process is equivalent to another if each individual CP process in the one is equivalent to one in the other.

Proposition 3.71. *If $\mathcal{P} \equiv_{\mathcal{C}} \mathcal{Q}$, then there is an isomorphism $f : \mathcal{P} \rightarrow \mathcal{Q}$ between the elements of the multisets $\mathcal{P}$ and $\mathcal{Q}$ such that $P \equiv_{\mathcal{C}} f(P)$ for every $P \in \mathcal{P}$ and $f^{-1}(Q) \equiv_{\mathcal{C}} Q$ for every $Q \in \mathcal{Q}$.*

Multiset CP's reduction is the existential pointwise lifting of CP's reduction, i.e. one Multiset CP process reduces to another if one of its component CP processes reduces to one in the other, and all other component CP processes are preserved.

Proposition 3.72. *If $\mathcal{P} \rightarrow_{\mathcal{C}} \mathcal{Q}$, then there is exactly one pair of processes $(P, Q) \in \mathcal{P} \times \mathcal{Q}$ such that $P \rightarrow_{\mathcal{C}} Q$, and an isomorphism $f : \mathcal{P}' \rightarrow \mathcal{Q}'$ between the remaining elements of the multisets, $\mathcal{P}' = \mathcal{P} - \{P\}$ and $\mathcal{Q}' = \mathcal{Q} - \{Q\}$, such that $P' \equiv_{\mathcal{C}} f(P')$ for every $P' \in \mathcal{P}'$ and $f^{-1}(Q') \equiv_{\mathcal{C}} Q'$ for every $Q' \in \mathcal{Q}'$.*

3.2.7 Fission, Fusion, and Disentanglement

In this section, I discuss the correspondence between Hypersequent CP and Multiset CP. The correspondence has three components—the titular components of this section. Fission is a translation from Multiset CP to Hypersequent CP, and, together, fusion and disentanglement are a translation from Hypersequent CP to Multiset CP. I discuss each component of the correspondence in order, starting with the easy one.

Fission

Fission is a translation from Multiset CP to Hypersequent CP, which splits CP's cut, send, and close into the corresponding HCP processes:

$$\begin{array}{ll}
 (\nu x\bar{x})(P \parallel Q) & \text{to } (\nu x\bar{x}) _ (P \parallel Q) \\
 x[y].(P \parallel Q) & \text{to } x[y]. _ (P \parallel Q) \\
 x[] \cdot 0 & \text{to } x[] \cdot _ 0
 \end{array}$$

Visually, its definition resembles an elaborate identity function.

Definition 3.73. *Fission, $\llbracket P \rrbracket_H$, translates CP processes into HCP processes by splitting each cut into a separate name restriction and parallel composition, and likewise splitting send and close.*

$$\begin{array}{ll}
 \llbracket x \leftrightarrow y \rrbracket_H & \triangleq x \leftrightarrow y \\
 \llbracket (\nu x\bar{x})(P \parallel Q) \rrbracket_H & \triangleq (\nu x\bar{x})(\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H) \\
 \llbracket x[y].(P \parallel Q) \rrbracket_H & \triangleq x[y].(\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H) \\
 \llbracket x(y).P \rrbracket_H & \triangleq x(y). \llbracket P \rrbracket_H \\
 \llbracket x[] \cdot 0 \rrbracket_H & \triangleq x[] \cdot 0 \\
 \llbracket x().P \rrbracket_H & \triangleq x(). \llbracket P \rrbracket_H \\
 \llbracket x \triangleleft \text{inl}.P \rrbracket_H & \triangleq x \triangleleft \text{inl}. \llbracket P \rrbracket_H \\
 \llbracket x \triangleleft \text{inr}.P \rrbracket_H & \triangleq x \triangleleft \text{inr}. \llbracket P \rrbracket_H \\
 \llbracket x \triangleright \{\text{inl}: P; \text{inr}: Q\} \rrbracket_H & \triangleq x \triangleright \{\text{inl}: \llbracket P \rrbracket_H; \text{inr}: \llbracket Q \rrbracket_H\} \\
 \llbracket x \not\leq N \rrbracket_H & \triangleq x \not\leq N
 \end{array}$$

Fission is extended to Multiset CP as follows:

$$\begin{array}{ll}
 \llbracket P \rrbracket_H & \triangleq \llbracket P \rrbracket_H \\
 \llbracket \mathcal{P} \parallel \mathcal{Q} \rrbracket_H & \triangleq \llbracket \mathcal{P} \rrbracket_H \parallel \llbracket \mathcal{Q} \rrbracket_H
 \end{array}$$

Fission is extended to process contexts, where it preserves holes and otherwise acts as it does on processes. Fission is the identity on types and typing environments, mapping each connective in CP to the same connective in HCP.

Fission distributes over arbitrary process contexts.

Proposition 3.74. $\llbracket (P[Q]) \rrbracket_H = (\llbracket P[_] \rrbracket_H) \llbracket [Q] \rrbracket_H$

Proof. By induction on the structure of the evaluation context $P[_]$. □

Fission preserves types, structural congruence, and reduction.

Proposition 3.75. *If $P \vdash_c T$, then $\llbracket P \rrbracket_H \vdash \llbracket T \rrbracket_H$.*

Proof. By induction on the structure of the derivation of $P \vdash_c T$.

- Case T-CUT.

We have

$$\frac{P \vdash_c \Gamma, x : A \quad Q \vdash_c \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(P \parallel Q) \vdash_c \Gamma, \Delta} \text{T-CUT}$$

By induction, $\llbracket P \rrbracket_H \vdash \llbracket \Gamma \rrbracket_H, x : \llbracket A \rrbracket_H$ and $\llbracket Q \rrbracket_H \vdash \llbracket \Delta \rrbracket_H, \bar{x} : \llbracket \bar{A} \rrbracket_H$.

The result follows as

$$\frac{\frac{\llbracket P \rrbracket_H \vdash \llbracket \Gamma \rrbracket_H, x : \llbracket A \rrbracket_H \quad \llbracket Q \rrbracket_H \vdash \llbracket \Delta \rrbracket_H, \bar{x} : \llbracket \bar{A} \rrbracket_H}{\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H \vdash \llbracket \Gamma \rrbracket_H, x : \llbracket A \rrbracket_H \parallel \llbracket \Delta \rrbracket_H, \bar{x} : \llbracket \bar{A} \rrbracket_H} \text{T-PAR}}{(\text{vx}\bar{x})(\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H) \vdash \llbracket \Gamma \rrbracket_H, \llbracket \Delta \rrbracket_H} \text{T-CUT}$$

- Case T-SEND.

We have

$$\frac{P \vdash_c \Gamma, y : A \quad Q \vdash_c \Delta, x : B}{x[y].(P \parallel Q) \vdash_c \Gamma, \Delta, x : A \otimes B} \text{T-SEND}$$

By induction, $\llbracket P \rrbracket_H \vdash \Gamma, y : A$ and $\llbracket Q \rrbracket_H \vdash \Delta, x : B$.

The result follows as

$$\frac{\frac{\llbracket P \rrbracket_H \vdash \Gamma, y : A \quad \llbracket Q \rrbracket_H \vdash \Delta, x : B}{\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H \vdash \Gamma, y : A \parallel \Delta, x : B} \text{T-PAR}}{x[y].(\llbracket P \rrbracket_H \parallel \llbracket Q \rrbracket_H) \vdash \Gamma, \Delta, x : A \otimes B} \text{T-SEND}$$

- Case T-CLOSE. We have

$$\frac{}{x[\cdot].0 \vdash_c x : \mathbf{1}} \text{T-CLOSE}$$

The result follows as

$$\frac{0 \vdash \circ}{x[\cdot].0 \vdash x : \mathbf{1}}$$

- Cases T-LINK, T-RECV, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, T-ABSURD.

The result follows immediately from the HCP rules T-LINK, T-RECV, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, T-ABSURD and the induction hypotheses (if any).

□

Proposition 3.76. *If $P \equiv_c Q$, then $\llbracket P \rrbracket_H \equiv \llbracket Q \rrbracket_H$.*

Proof. The cases for reflexivity, symmetry, transitivity, and congruence follow from the same properties of the structural congruence of HCP. The cases for the axioms are as follows:

- Case SC-LINKCOMM.

Immediately, by SC-LINKCOMM.

- Case SC-CUTCOMM.

Immediately, by SC-NEWCOMM and SC-PARCOMM.

- Case SC-CUTASSOC.

Immediately, by SC-SCOPEEXT, SC-NEWASSOC, and SC-PARASSOC.

□

Proposition 3.77. *If $P \rightarrow_c Q$, then $\llbracket P \rrbracket_H \rightarrow \llbracket Q \rrbracket_H$.*

Proof. By induction on the structure of the derivation $P \rightarrow_c Q$.

- Cases E-LINK, E-SEND, E-CLOSE, E-SELECT₁, and E-SELECT₂.

Immediately, by E-LINK, E-SEND, E-CLOSE, E-SELECT₁, and E-SELECT₂, respectively.

- Case E-EQUIV.

By E-EQUIV, the induction hypothesis, and Proposition 3.76.

- Case E-CONG.

By E-CONG, the induction hypothesis, and Proposition 3.74.

□

Fusion

Fusion is the partial inverse of fission. It fuses an HCP name restriction and parallel composition into a CP cut, an HCP send and parallel composition into a CP send, and an HCP close and parallel composition into a CP close, but *only* if the HCP process already has the form of a CP process.

$$\begin{array}{ll}
 (vx\bar{x})._ (P \parallel Q) & \text{to } (vx\bar{x})(P \parallel Q) \\
 x[y]._ (P \parallel Q) & \text{to } x[y].(P \parallel Q) \\
 x[]._ 0 & \text{to } x[].0
 \end{array}$$

Visually, its definition likewise resembles an elaborate identity function.

Definition 3.78. *Fusion, $\llbracket P \rrbracket_c$, translates HCP processes into CP processes by fusing adjacent name restrictions and parallel compositions into cuts,*

and likewise fusing send and close. Fusion is partial.

$$\begin{array}{ll}
\llbracket x \leftrightarrow y \rrbracket_c & \triangleq x \leftrightarrow y \\
\llbracket (vx\bar{x})(P \parallel Q) \rrbracket_c & \triangleq (vx\bar{x})(\llbracket P \rrbracket_c \parallel \llbracket Q \rrbracket_c) \\
\llbracket x[y].(P \parallel Q) \rrbracket_c & \triangleq x[y].(\llbracket P \rrbracket_c \parallel \llbracket Q \rrbracket_c) \\
\llbracket x(y).P \rrbracket_c & \triangleq x(y).\llbracket P \rrbracket_c \\
\llbracket x[].0 \rrbracket_c & \triangleq x[].0 \\
\llbracket x().P \rrbracket_c & \triangleq x().\llbracket P \rrbracket_c \\
\llbracket x \triangleleft \text{inl}.P \rrbracket_c & \triangleq x \triangleleft \text{inl}.\llbracket P \rrbracket_c \\
\llbracket x \triangleleft \text{inr}.P \rrbracket_c & \triangleq x \triangleleft \text{inr}.\llbracket P \rrbracket_c \\
\llbracket x \triangleright \{\text{inl}:P; \text{inr}:Q\} \rrbracket_c & \triangleq x \triangleright \{\text{inl}:\llbracket P \rrbracket_c; \text{inr}:\llbracket Q \rrbracket_c\} \\
\llbracket x \not\downarrow N \rrbracket_c & \triangleq x \not\downarrow N
\end{array}$$

Fusion is extended to Multiset CP as follows:

$$\llbracket P \rrbracket_e \triangleq \begin{cases} \llbracket Q \rrbracket_e \parallel \llbracket R \rrbracket_e & \text{if } P = Q \parallel R \\ \llbracket P \rrbracket_c & \text{otherwise} \end{cases}$$

Fusion is extended to process contexts, where it preserves holes and otherwise acts as it does on processes. Fusion is the identity on types and typing environments, mapping each connective in HCP to the same connective in CP.

Fusion is the partial inverse of fission.

Fusion distributes over arbitrary process contexts.

Lemma 3.79. $\llbracket P[Q] \rrbracket_c = \llbracket P[\cdot] \rrbracket_c \llbracket [Q]_c \rrbracket_c$.

Fusion preserves types.

Proposition 3.80. If $\llbracket P \rrbracket_H \vdash \llbracket \Gamma \rrbracket_H$, then $P \vdash_c \Gamma$.

Proof. By induction on the structure of the derivation $\llbracket P \rrbracket_H \vdash \llbracket \Gamma \rrbracket_H$.

- Cases T-NEW, T-SEND, and T-CLOSE.

The result follows from the CP rules T-CUT, T-SEND, and T-CLOSE and inversion on the structure of the CP process P , which establishes that the derivations constructed in Case $(\Rightarrow)$ are the only derivations.

- Cases T-PAR and T-HALT.

Impossible, by inversion on the structure of the CP process P .

- Cases T-LINK, T-RECV, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, and T-ABSURD.

The result follows from the induction hypotheses (if any) and the CP rule of the same name, i.e. T-LINK, T-RECV, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, and T-ABSURD, respectively.

□

Fusion is defined for all disentangled processes.

Lemma 3.81. *If $P \vdash \mathcal{G}$ and $Q \in \mathbb{P}$, then $\llbracket Q \rrbracket_{\mathcal{G}}$ is defined.*

Proof. By induction on the structure of Q with Lemma 3.60. □

The fusions of the disentangled processes are equivalent up to CP's link-preserving structural congruence.

Lemma 3.82. *If $P \vdash \mathcal{G}$ and $Q, R \in \mathbb{P}$, then $\llbracket Q \rrbracket_{\mathcal{G}} \stackrel{1}{\equiv}_{\mathcal{G}} \llbracket R \rrbracket_{\mathcal{G}}$.*

Proof. The proof is an extension to the proof of Proposition 3.61, which proves the following two facts:

1. In the proof for Lemma 3.56, under the case $V_T > 1$, we pick an arbitrary leaf of the connection tree T . Any choice is equivalent up to CP's structural congruence.

Assume that P_i and P_j are both leaves of T . If we pick P_i and then P_j , we get (a), but if we pick P_j and then P_i , we get (b).

$$\begin{array}{cc} \text{(a)} & \text{(b)} \\ Q \parallel (vx\bar{x})(P_i \parallel (vy\bar{y})(P_j \parallel R')) & Q \parallel (vy\bar{y})(P_j \parallel (vx\bar{x})(P_i \parallel R')) \end{array}$$

As P_i and P_j are both leaves of T , $x, \bar{x} \notin \text{fn}(P_j)$ and $y, \bar{y} \notin \text{fn}(P_i)$, which means the two processes are equivalent by SC-CUTCOMM and SC-CUTASSOC.

If we pick P_i , then pick a number of processes that have become leaves by the removal of P_i and successive choices, *and then* pick P_j , we get (a), where there is some (right-branching) evaluation context $\mathcal{E}$ between P_i and P_j , but if we pick P_j , then pick P_i , and then pick the processes in $\mathcal{E}$, we get (b).

$$\begin{array}{cc} \text{(a)} & \text{(b)} \\ Q \parallel (vx\bar{x})(P_i \parallel \mathcal{E}[(vy\bar{y})(P_j \parallel R')]) & Q \parallel (vy\bar{y})(P_j \parallel (vx\bar{x})(P_i \parallel \mathcal{E}[R'])) \end{array}$$

As P_j is a leaf of T , $\text{fn}(P_j) \cap \text{bn}(\mathcal{E}) = \emptyset$, which means the two processes are equivalent by Lemma 3.24 and Lemma 3.25.

2. In the proof for Proposition 3.57, under the case where $\mathcal{G}$ is of the form $\mathcal{G} \parallel \Gamma$, we pick an arbitrary tree of the connection forest G . Any choice is equivalent up to Multiset CP's structural congruence.

The result follows immediate from the multiset structure of Multiset CP's processes, i.e. by SC-PARCOMM and SC-PARASSOC. □

Disentanglement into CP is the composition of disentanglement and fusion.

Definition 3.83 (Disentanglement into CP). *The disentanglement of P into CP, $\llbracket P \rrbracket_{\mathcal{C}}$, is the composition of disentanglement and fusion. Formally,*

$$\llbracket P \rrbracket_{\mathcal{C}} \triangleq \{ \llbracket Q \rrbracket_{\mathcal{C}} \mid Q \in \{P\} \}$$

By Lemma 3.81 and Lemma 3.82, the set $\llbracket P \rrbracket_{\mathcal{C}}$ is non-empty for any P , and all elements of $\llbracket P \rrbracket_{\mathcal{C}}$ are equivalent to each other under Multiset CP's link-preserving structural congruence. I use $\llbracket \cdot \rrbracket_{\mathcal{C}}$ as functional up to structural congruence, and write $\llbracket P \rrbracket_{\mathcal{C}} \triangleq_{\mathcal{C}} \mathcal{Q}$ to mean $\mathcal{Q}$ is an arbitrary element of $\llbracket P \rrbracket_{\mathcal{C}}$.

Disentanglement into CP does not distribute over process contexts, as disentanglement normalises the entire maximum configuration context of a process in a single step. Consequently, when doing proof by induction over the structure of a process, we must generalise the cases for name restriction and parallel composition to cover the entire maximum configuration context, as done in, e.g. the proof of Proposition 3.85.

Disentanglement into CP preserves types, structural congruence, and reduction.

Proposition 3.84. *If $P \vdash \mathcal{G}$ and $\llbracket P \rrbracket_{\mathcal{C}} \triangleq_{\mathcal{C}} \mathcal{P}$, then $\mathcal{P} \vdash_{\mathcal{C}} \llbracket \mathcal{G} \rrbracket_{\mathcal{C}}$.*

Proof. By Proposition 3.61, Lemma 3.27, Proposition 3.80, and Lemma 2.24. □

Proposition 3.85. *If $P \vdash \mathcal{G}$ and $P \equiv Q$, then $\llbracket P \rrbracket_{\mathcal{C}} \equiv_{\mathcal{C}} \llbracket Q \rrbracket_{\mathcal{C}}$.*

Proof. By induction on the derivation of $P \vdash \mathcal{G}$ and inversion on P . There are three distinct cases, based on whether P is a link, some other ready process, or some other process construct.

- Case P is of the form $x \leftrightarrow y$.

By inversion on $P \equiv Q$, there are two cases. Either $Q = x \leftrightarrow y$, in which case the result follows by reflexivity, or $Q = y \leftrightarrow x$, in which case the result follows by SC-LINKCOMM.

- Case P is of the form $x[y].P'$, $x(y).P'$, $x[] . P'$, $x().P'$, $x \triangleleft \text{inl}.P'$, or $x \triangleleft \text{inr}.P'$.

By inversion on $P \equiv Q$, Q is of the form $x[y].Q'$, $x(y).Q'$, $x[] . Q'$, $x().Q'$, $x \triangleleft \text{inl}.Q'$, or $x \triangleleft \text{inr}.Q'$. By the induction hypothesis, $\llbracket P' \rrbracket_{\mathcal{C}} \equiv_{\mathcal{C}} \llbracket Q' \rrbracket_{\mathcal{C}}$. The result follows by congruence.

- Case P is of the form $x \triangleright \{\text{inl}: P_1; \text{inr}: P_2\}$.

As above, but with two appeals to the induction hypothesis.

- Case P is of the form $x \not\leq N$.

By inversion on $P \equiv Q$, $P = Q$. The result follows by reflexivity.

- Case **P** is of the form $(\nu x\bar{x})P'$, $P_1 \parallel P_2$, or **0**.

By Lemma 3.11, there exists some **R** such that $P \equiv R$ and $R \equiv Q$. By Lemma 3.65, $\llbracket P \rrbracket_{\mathcal{C}} = \llbracket R \rrbracket_{\mathcal{C}}$. Hence, $\llbracket P \rrbracket_{\mathcal{C}} = \llbracket R \rrbracket_{\mathcal{C}}$. We have **R** is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ for some $n \geq 0$. By inversion on the structure of $\mathcal{C}[\cdot]$, decompose $R \equiv Q$ into separate structural congruences for each thread $T_i \equiv T'_i$ for $1 \leq i \leq n$. (The “deep” restriction is not lost. For ready processes, $\equiv$ is exactly $\equiv$.) By the induction hypothesis, $\llbracket T_i \rrbracket_{\mathcal{C}} \equiv_{\mathcal{C}} \llbracket T'_i \rrbracket_{\mathcal{C}}$ for $1 \leq i \leq n$. The result follows by composing these intermediate results using Lemma 3.64.

□

Lemma 3.86. Any derivation of a reduction $P \rightarrow Q$ can be rewritten to

$$\frac{P \equiv \mathcal{C}[(\nu x\bar{x})(P_1 \parallel P_2)] \quad \frac{\frac{(\nu x\bar{x})(P_1 \parallel P_2) \rightarrow R}{\mathcal{C}[(\nu x\bar{x})(P_1 \parallel P_2)] \rightarrow \mathcal{C}[R]} \quad \mathcal{C}[R] \equiv Q}{P \rightarrow Q} r$$

where r is E-LINK, E-SEND, E-CLOSE, E-SELECT₁, or E-SELECT₂.

Proof. By induction on the derivation of the reduction $P \rightarrow Q$. In the cases for E-LINK, E-SEND, E-CLOSE, E-SELECT₁, or E-SELECT₂, the result follows immediately. In the case for E-CONG, the result follows by composing the two evaluation contexts. In the case for E-EQUIV, the result follows by composing the structural congruences using transitivity.

□

Proposition 3.87. If $P \vdash \mathcal{G}$ and $P \rightarrow Q$, then $\llbracket P \rrbracket_{\mathcal{C}} \rightarrow_{\mathcal{C}} \llbracket Q \rrbracket_{\mathcal{C}}$.

Proof. By Lemma 3.86, rewrite the derivation of the reduction $P \rightarrow Q$ to

$$\frac{P \equiv (\nu x\bar{x})(P_1 \parallel P_2) \quad \frac{\frac{(\nu x\bar{x})(P_1 \parallel P_2) \rightarrow R}{\mathcal{C}[(\nu x\bar{x})(P_1 \parallel P_2)] \rightarrow \mathcal{C}[R]} \quad \mathcal{C}[R] \equiv Q}{P \rightarrow Q} r$$

where r is E-LINK, E-SEND, E-CLOSE, E-SELECT₁, or E-SELECT₂.

By Proposition 3.85, $\llbracket P \rrbracket_{\mathcal{C}} \equiv_{\mathcal{C}} \llbracket \mathcal{C}[(\nu x\bar{x})(P_1 \parallel P_2)] \rrbracket_{\mathcal{C}}$ and $\llbracket \mathcal{C}[R] \rrbracket_{\mathcal{C}} \equiv_{\mathcal{C}} \llbracket Q \rrbracket_{\mathcal{C}}$.

Pick an arbitrary $P' \in \llbracket \mathcal{C}[(\nu x\bar{x})(P_1 \parallel P_2)] \rrbracket_{\mathcal{C}}$.

If r is not E-LINK, both P_1 and P_2 are ready, and, by Lemma 3.63, preserved up to link-preserving deep structural congruence, by disentanglement. By Lemma 2.21 and the fact that P' is in right-branching form, P' is either of the form $\mathcal{F}_1[(\nu x\bar{x})(P'_1 \parallel \mathcal{F}_2[P'_2])]$ or of the form $\mathcal{F}_2[(\nu x\bar{x})(P'_2 \parallel \mathcal{F}_1[P'_1])]$,

where $P'_1 \in \llbracket P_1 \rrbracket_{\mathcal{C}}$ and $P'_2 \in \llbracket P_2 \rrbracket_{\mathcal{C}}$. In the first case, the CP reduction is constructed as:

$$\begin{aligned} \mathcal{F}_1[(vx\bar{x})(P'_1 \parallel \mathcal{F}_2[P'_2])] &\stackrel{\text{ls}}{\equiv}_{\mathcal{C}} \langle \text{by Lemma 2.23} \rangle \\ \mathcal{F}_1[\mathcal{F}_2[(vx\bar{x})(P'_1 \parallel P'_2)]] &\longrightarrow_{\mathcal{C}} \langle \text{by } r \rangle \\ \mathcal{F}_1[\mathcal{F}_2[R']] & \end{aligned}$$

By case analysis on r , it follows that $\mathcal{F}_1[\mathcal{F}_2[R']] \stackrel{\text{ls}}{\equiv}_{\mathcal{C}} \llbracket \mathcal{E}[R] \rrbracket_{\mathcal{C}}$.

In the second case, the CP reduction is constructed similarly.

If r is E-LINK, the CP reduction is constructed similarly, since, if P_1 is of the form $x \leftrightarrow y$, the thread in P_2 that contains $\bar{x}$ is preserved up to fusion and link-preserving deep structural congruence in P' .

□

3.2.8 Label-Transition System and Harmony

In the previous sections, we got HCP up to speed. We adapted CP's metatheory and showed that HCP enjoys all the same properties, and characterised the exact correspondence between CP and HCP by disentanglement.

In this section, we get to enjoy some of the spoils of all that work. We introduce the beginnings of a behavioural theory for HCP, *just enough* to relate the use of the label-transition system in the following chapter to the reduction semantics, and to use as a springboard to find and fix further flaws in our present formulation of HCP in the discussion.

The label-transition semantics for HCP, first introduced by Montesi and Peressotti [2018], are given by labelled transition, which is a ternary relation on two processes and a label. A label is either the internal action label τ , an observable action label π , or a pair of parallel observable action labels $\pi \parallel \bar{\pi}$.

$$\ell ::= \tau \mid \pi \mid \pi \parallel \bar{\pi}$$

The observable action labels correspond to the action prefixes of processes, except for the offer labels $x \triangleright \text{inl}$ and $x \triangleright \text{inr}$, which denote the offer of inl or inr , respectively, in $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$.

$$\pi ::= x \leftrightarrow y \mid x[y] \mid x(y) \mid x[] \mid x() \mid x \triangleleft \text{inl} \mid x \triangleleft \text{inr} \mid x \triangleright \text{inl} \mid x \triangleright \text{inr}$$

An endpoint occurs free in an action label in the same cases where it would occur free in a process prefixed by that action. Likewise, an endpoint is bound by an action label in the same cases where it would occur bound in a process prefixed by that action. For instance, $x \in \text{fn}(x[y])$ and $y \in \text{bn}(x[y])$. An endpoint occurs free in $\pi \parallel \bar{\pi}$ when it occurs free in either π or $\bar{\pi}$. An endpoint is bound by $\pi \parallel \bar{\pi}$ when it is bound by either π or $\bar{\pi}$. The label τ has no free or bound names.

Definition 3.88 (Label-Transition). *Label-transition, written $P \xrightarrow{\ell} Q$, is the smallest ternary relation on two processes and one label defined by the rules in Figure 3.4.*

What does a label-transition mean? If a process has an observable action, i.e. a transition with an observable action label π , it means the action π is observable—some unblocked sub-process is prefixed by the action π on a free endpoint. If a process has pair of parallel observable actions, i.e. a transition with a label $\pi \parallel \bar{\pi}$, it means both actions are observable and independent. If a process has a τ -transition, it has some internal, unobservable reduction. To use familiar terminology:

$$\begin{aligned} P \xrightarrow{\pi} Q &\iff P \equiv \pi.P' \parallel R \\ P \xrightarrow{\pi \parallel \bar{\pi}} Q &\iff P \equiv \pi.P_1 \parallel \bar{\pi}.P_2 \parallel R \\ P \xrightarrow{\tau} Q &\iff P \rightarrow Q \end{aligned}$$

These properties are intended to provide an *informal* intuition. They are not well-formed for the offer, since, e.g. “ $x \triangleright \text{inl}. P$ ” is a syntactically ill-formed process. However, the intuition still applies:

$$\begin{aligned} P \xrightarrow{x \triangleright \text{inl}} Q &\iff P \equiv x \triangleright \{\text{inl}: P_1; \text{inr}: P_2\} \parallel R \\ P \xrightarrow{x \triangleright \text{inr}} Q &\iff P \equiv x \triangleright \{\text{inl}: P_1; \text{inr}: P_2\} \parallel R \\ P \xrightarrow{x \triangleleft \text{inl} \parallel x \triangleright \text{inl}} Q &\iff P \equiv x \triangleleft \text{inl}. P_1 \parallel x \triangleright \{\text{inl}: P_2; \text{inr}: P_3\} \parallel R \\ P \xrightarrow{x \triangleleft \text{inr} \parallel x \triangleright \text{inr}} Q &\iff P \equiv x \triangleleft \text{inr}. P_1 \parallel x \triangleright \{\text{inl}: P_2; \text{inr}: P_3\} \parallel R \end{aligned}$$

Label-transitions are closed under independent evaluation contexts by STR-CONG, which requires that the bound and free endpoints of the evaluation context and the label are disjoint. Consequently, τ -transitions, whose labels have neither free nor bound endpoints, are closed under arbitrary evaluation contexts.

The label-transition system does not use the structural congruence. As a consequence of a label-transition, a process may change *locally* and it may close the channel or open a new channel. However, the structure of any name restrictions and parallel compositions *unrelated* to the communication does not change. On the other hand, the structural congruence does not essentially alter actions. Label-transition and structural congruence are mutually invariant:

$$P \xRightarrow{\ell} Q \iff P \xrightarrow{\ell} Q$$

Reductions may contain arbitrary structural congruence, and change the process structure in ways that are not strictly necessary for the reduction itself. Hence, while a τ -transition corresponds to a reduction, a reduction corresponds to a τ -transition *and* a structural congruence.

Label-transition and structural congruence are mutually invariant.

Action Rules

$x \leftrightarrow y \xrightarrow{x \leftrightarrow y} 0$	ACT-LINK ₁	$x \leftrightarrow y \xrightarrow{y \leftrightarrow x} 0$	ACT-LINK ₂
$x[y].P \xrightarrow{x[y]} P$	ACT-SEND	$x \triangleleft \text{inl}.P \xrightarrow{x \triangleleft \text{inl}} P$	ACT-SELECT ₁
$x(y).P \xrightarrow{x(y)} P$	ACT-RECV	$x \triangleleft \text{inr}.P \xrightarrow{x \triangleleft \text{inr}} P$	ACT-SELECT ₂
$x[].P \xrightarrow{x[]} P$	ACT-CLOSE	$x \triangleright \{\text{inl}: P; \text{inr}: Q\} \xrightarrow{x \triangleright \text{inl}} P$	ACT-OFFER ₁
$x().P \xrightarrow{x()} P$	ACT-WAIT	$x \triangleright \{\text{inl}: P; \text{inr}: Q\} \xrightarrow{x \triangleright \text{inr}} Q$	ACT-OFFER ₂

Structural Rules

STR-PAR

$$\frac{P \xrightarrow{\pi} P' \quad Q \xrightarrow{\tilde{\pi}} Q' \quad \text{bn}(\pi) \cap \text{bn}(\tilde{\pi}) = \emptyset}{P \parallel Q \xrightarrow{\pi \parallel \tilde{\pi}} P' \parallel Q'}$$

STR-CONG

$$\frac{P \xrightarrow{\pi} P' \quad (\text{bn}(\mathcal{E}) \cup \text{fn}(\mathcal{E})) \cap (\text{bn}(\pi) \cup \text{fn}(\pi)) = \emptyset}{\mathcal{E}[P] \xrightarrow{\pi} \mathcal{E}[P']}$$

Communication Rules

TAU-LINK

$$\frac{P \xrightarrow{x \leftrightarrow y} P'}{(vx\bar{x})P \xrightarrow{\tau} P'\{y/\bar{x}\}}$$

TAU-SEND-RECV

$$\frac{P \xrightarrow{x[y]\|\bar{x}(\bar{y})} P'}{(vx\bar{x})P \xrightarrow{\tau} (vx\bar{x})(vy\bar{y})P'}$$

TAU-CLOSE-WAIT

$$\frac{P \xrightarrow{x[]\|\bar{x}()} P'}{(vx\bar{x})P \xrightarrow{\tau} P'}$$

TAU-SELECT-OFFER₁

$$\frac{P \xrightarrow{x \triangleleft \text{inl}\|\bar{x} \triangleright \text{inl}} P'}{(vx\bar{x})P \xrightarrow{\tau} (vx\bar{x})P'}$$

TAU-SELECT-OFFER₂

$$\frac{P \xrightarrow{x \triangleleft \text{inr}\|\bar{x} \triangleright \text{inr}} P'}{(vx\bar{x})P \xrightarrow{\tau} (vx\bar{x})P'}$$

Figure 3.4: Label-Transition System for Hypersequent CP

Lemma 3.89. *If $P \xRightarrow{\ell} Q$, then $P \xrightarrow{\ell} \equiv Q$.*

Proof. By induction on the derivation of the structural congruence $P \equiv Q$ and inversion on the transition $P \xrightarrow{\ell} P'$.

(For the full proof, see § 3.5.)

□

Any derivation of a label-transition can be normalised, by combining successive uses of the STR-CONG rule. Lemmas 3.90 to 3.92 normalise action transitions, parallel action transitions, and τ -transitions, respectively.

Lemma 3.90. *Any derivation of a transition $P \xrightarrow{\pi} Q$ can be rewritten to*

$$\frac{\frac{\overline{\quad} a}{P_1 \xrightarrow{\pi} P'_1} \text{STR-CONG}}{\mathcal{F}_1[P_1] \xrightarrow{\pi} \mathcal{F}_1[P'_1]} \text{STR-CONG}$$

where a is ACT-LINK_1 , ACT-LINK_2 , ACT-SEND , ACT-RECV , ACT-CLOSE , ACT-WAIT , ACT-SELECT_1 , ACT-SELECT_2 , ACT-OFFER_1 , or ACT-OFFER_2 .

Proof. By induction on the derivation of the transition $P \xrightarrow{\pi} Q$. The cases for ACT-LINK_1 , ACT-LINK_2 , ACT-SEND , ACT-RECV , ACT-CLOSE , ACT-WAIT , ACT-SELECT_1 , ACT-SELECT_2 , ACT-OFFER_1 , or ACT-OFFER_2 follow immediately. The case for STR-CONG follows by induction, composing the evaluation context introduced by STR-CONG with the evaluation context in the induction hypothesis.

□

Lemma 3.91. *Any derivation of a transition $P \xrightarrow{\tau \parallel \bar{\tau}} Q$ can be rewritten to*

$$\frac{\frac{\frac{\overline{\quad} a}{P_1 \xrightarrow{\pi} P'_1} \text{STR-CONG}}{\mathcal{F}_1[P_1] \xrightarrow{\pi} \mathcal{F}_1[P'_1]} \quad \frac{\frac{\overline{\quad} \bar{a}}{P_2 \xrightarrow{\bar{\pi}} P'_2} \text{STR-CONG}}{\mathcal{F}_2[P_2] \xrightarrow{\bar{\pi}} \mathcal{F}_2[P'_2]} \text{STR-CONG}}{\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2] \xrightarrow{\tau \parallel \bar{\tau}} \mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]} \text{STR-PAR} \text{STR-CONG}$$

$$\frac{\mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\tau \parallel \bar{\tau}} \mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]}{\mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\tau \parallel \bar{\tau}} \mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]}}$$

where a and $\bar{a}$ are one of ACT-LINK_1 , ACT-LINK_2 , ACT-SEND , ACT-RECV , ACT-CLOSE , ACT-WAIT , ACT-SELECT_1 , ACT-SELECT_2 , ACT-OFFER_1 , or ACT-OFFER_2 .

Proof. By induction on the derivation of the transition $P \xrightarrow{\pi} Q$. The case for STR-PAR follows by Lemma 3.90. The case for STR-CONG follows by induction, composing the evaluation context introduced by STR-CONG with the evaluation context in the induction hypothesis.

□

Lemma 3.92. Any derivation of a transition $P \xrightarrow{\tau} Q$ can be rewritten to

$$\begin{array}{c}
 \frac{\frac{\frac{}{P_1 \xrightarrow{\pi} P'_1} a}{\mathcal{F}_1[P_1] \xrightarrow{\pi} \mathcal{F}_1[P'_1]} \text{STR-CONG} \quad \frac{\frac{\frac{}{P_2 \xrightarrow{\bar{\pi}} P'_2} \bar{a}}{\mathcal{F}_2[P_2] \xrightarrow{\bar{\pi}} \mathcal{F}_2[P'_2]} \text{STR-CONG}}{\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2] \xrightarrow{\pi \parallel \bar{\pi}} \mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]} \text{STR-PAR}}{\mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\pi \parallel \bar{\pi}} \mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]} \text{STR-CONG}}{(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\tau} \mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]} t} \text{STR-CONG} \\
 \mathcal{E}_1[(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]]] \xrightarrow{\tau} \mathcal{E}_1[\mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]]
 \end{array}$$

where t is one of TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂, and a and $\bar{a}$ are one of ACT-LINK₁, ACT-LINK₂, ACT-SEND, ACT-RECV, ACT-CLOSE, ACT-WAIT, ACT-SELECT₁, ACT-SELECT₂, ACT-OFFER₁, or ACT-OFFER₂.

Proof. By induction on the structure of the transition. The cases for TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂ follow by Lemma 3.91. The case for STR-CONG follows by induction, composing the evaluation context introduced by STR-CONG with the evaluation context in the induction hypothesis.

□

The equivalences between action transitions and parallel action transitions and structural congruence follow as corollaries from Lemmas 3.90 and 3.91 and Corollary 3.26, and their converses follow from Lemma 3.89.

$$\begin{array}{lcl}
 P \xrightarrow{\pi} Q & \iff & P \equiv \pi.P' \parallel R \\
 P \xrightarrow{\pi \parallel \bar{\pi}} Q & \iff & P \equiv \pi.P_1 \parallel \bar{\pi}.P_2 \parallel R
 \end{array}$$

Harmony follows by relating the normal-form derivations of τ -transition (Lemma 3.92) and reduction (Lemma 3.86).

Proposition 3.93 (Harmony). If $P \vdash \mathcal{G}$, then $P \xrightarrow{\tau} \equiv Q \iff P \longrightarrow Q$.

Proof. There are two cases:

- Case ($\Rightarrow$).

By Lemma 3.92, rewrite the derivation of the transition $P \xrightarrow{\tau} Q$ to

$$\begin{array}{c}
 \frac{\frac{\frac{}{P_1 \xrightarrow{\pi} P'_1} a}{\mathcal{F}_1[P_1] \xrightarrow{\pi} \mathcal{F}_1[P'_1]} \text{STR-CONG} \quad \frac{\frac{\frac{}{P_2 \xrightarrow{\bar{\pi}} P'_2} \bar{a}}{\mathcal{F}_2[P_2] \xrightarrow{\bar{\pi}} \mathcal{F}_2[P'_2]} \text{STR-CONG}}{\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2] \xrightarrow{\pi \bar{\pi}} \mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]} \text{STR-PAR}}{\mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\pi \bar{\pi}} \mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]} \text{STR-CONG}} \\
 \frac{\mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\pi \bar{\pi}} \mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]}{(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\tau} \mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]} t \\
 \frac{(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]] \xrightarrow{\tau} \mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]}{\mathcal{E}_1[(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]]] \xrightarrow{\tau} \mathcal{E}_1[\mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]]} \text{STR-CONG}
 \end{array}$$

where t is one of TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂, and a and $\bar{a}$ are one of ACT-LINK₁, ACT-LINK₂, ACT-SEND, ACT-RECV, ACT-CLOSE, ACT-WAIT, ACT-SELECT₁, ACT-SELECT₂, ACT-OFFER₁, or ACT-OFFER₂. By repeated application of Lemmas 3.24 and 3.25:

$$\mathcal{E}_1[(\nu x \bar{x}) \mathcal{E}_3[\mathcal{F}_1[P_1] \parallel \mathcal{F}_2[P_2]]] \equiv \mathcal{E}_1[\mathcal{E}_3[\mathcal{F}_1[\mathcal{F}_2[(\nu x \bar{x})(P_1 \parallel P_2)]]]]$$

By case analysis on t and inversion on a and $\bar{a}$, the structure of the messages π and $\bar{\pi}$, and the processes P_1 , P'_1 , P_2 , and P'_2 , and $\mathcal{E}_2$:

$$\mathcal{E}_1[\mathcal{E}_3[\mathcal{F}_1[\mathcal{F}_2[(\nu x \bar{x})(P_1 \parallel P_2)]]]] \longrightarrow \mathcal{E}_1[\mathcal{E}_3[\mathcal{F}_1[\mathcal{F}_2[\mathcal{E}_2[P'_1 \parallel P'_2]]]]]$$

By repeated application of Lemmas 3.24 and 3.25:

$$\mathcal{E}_1[\mathcal{E}_3[\mathcal{F}_1[\mathcal{F}_2[\mathcal{E}_2[P'_1 \parallel P'_2]]]]] \equiv \mathcal{E}_1[\mathcal{E}_2[\mathcal{E}_3[\mathcal{F}_1[P'_1] \parallel \mathcal{F}_2[P'_2]]]]$$

The result follows by composing these results with E-EQUIV.

- Case ($\Leftarrow$).

By Lemma 3.86, rewrite the derivation of the reduction $P \longrightarrow Q$ to

$$\frac{P \equiv (\nu x \bar{x})(P_1 \parallel P_2) \quad \frac{\frac{(\nu x \bar{x})(P_1 \parallel P_2) \longrightarrow R}{\mathcal{E}[(\nu x \bar{x})(P_1 \parallel P_2)] \longrightarrow \mathcal{E}[R]} r}{\mathcal{E}[(\nu x \bar{x})(P_1 \parallel P_2)] \longrightarrow \mathcal{E}[R]} \quad \mathcal{E}[R] \equiv Q}{P \longrightarrow Q}$$

where r is E-LINK, E-SEND, E-CLOSE, E-SELECT₁, or E-SELECT₂. By case analysis on r and inversion on the structure of P_1 , P_2 , and R , $(\nu x \bar{x})(P_1 \parallel P_2) \xrightarrow{\tau} R$. By STR-CONG, $\mathcal{E}[(\nu x \bar{x})(P_1 \parallel P_2)] \xrightarrow{\tau} \mathcal{E}[R]$. By Lemma 3.89 and $P \equiv (\nu x \bar{x})(P_1 \parallel P_2)$, $P \xrightarrow{\tau} \mathcal{E}[R]$. By transitivity, $P \xrightarrow{\tau} \equiv Q$.

□

3.3 Discussion

This section proceeds as follows:

- In § 3.3.1, I introduce an alternative semantics for the absurd offer based on exception handling.
- In § 3.3.2, I introduce synchronous semantics for link.
- In § 3.3.3, I introduce a variant of HCP that decomposes the binary and n -ary offer into the singleton offer and summation, which allows us to type guarded summation and enables us to factor actions out into their own syntactic category.
- In § 3.3.4, I discuss channel names and endpoint names.
- In § 3.3.4, I discuss the relation between tensor and par and send and receive.
- In § 3.3.5, I discuss the relation between HCP and deep inference logic.
- In § 3.3.6, I discuss the relation between HCP and hypersequent calculus.
- In § 3.3.7, I discuss the relation between the empty hypersequent and hypersequent composition in HCP and the MIX_0 and MIX rules.

3.3.1 Zap: the Exceptionally Absurd Offer

In this section, we revisit the semantics for the absurd offer. In § 2.1.6, we discussed the ‘inert’ semantics for the absurd offer. It does nothing, just sits around. If we think of the absurd offer as an exception handler, the inert semantics make sense. It is waiting for an exception that—by the correctness of the type system—will never arrive. However, the absurd offer is not *quite* an exception handler. An exception handler handles an exceptional case—either the computation succeeds, and we compute a value, or, in some exceptional circumstances, something goes wrong, and the exception handler is invoked. The absurd offer deals with the *certainty* that something will go wrong. If you are certain, why bother waiting? Under this view, we refer to the absurd offer as *zap* or *zapper*. It is not inert, and does not wait. It *zaps* channels, kills processes, and shuts the whole thing down.

Wadler’s CP handles zap with a commuting conversion, defined on typing derivations:

$$\frac{a \triangleright \{\} \vdash_c \Gamma, a : \mathbb{T}, x : \bar{A} \quad P \vdash_c \Delta, \bar{x} : \bar{A}}{(vx\bar{x})(a \triangleright \{\} \parallel P) \vdash_c \Gamma, \Delta, a : \mathbb{T}} \longrightarrow_c \frac{}{a \triangleright \{\} \vdash_c \Gamma, \Delta, a : \mathbb{T}}$$

As mentioned in § 2.1.6, this reduction rule does not satisfy type erasure, since the process $x \triangleright \{\}$ does not record which sessions it is allowed to kill. In HCP and CP, as presented in this thesis, the absurd offer does record these sessions as part of the process syntax. Hence, we could adapt and repair Wadler’s reduction rule. To save on ink and eye strain, let us write N, x to mean $N \cup \{x\}$ and N, M to mean $N \cup M$.

$$\frac{\bar{x} \in \text{fn}(P) \quad M \triangleq \text{fn}(P) \setminus \{\bar{x}\}}{(vx\bar{x})(w \not\vdash N, x \parallel P) \longrightarrow w \not\vdash N, M} \text{E-ZAP-NOW}$$

The above rule is correct, in the sense that it matches Wadler’s semantics, but it is a bit over-eager. It permits $w \dot{\downarrow} N$ to kill any process that holds any of the names in N , regardless of whether those processes are currently attempting to communicate on those channels. Let us consider what requirements this places on an implementation.

- If we wanted to implement this as part of our channel implementation, it would require every process to always be listening on all of their endpoints.
- We could implement this using another mechanism entirely. For instance, using POSIX Threads, we could implement this by decorating each channel with the thread ID of the process on the other end of the channel, and implement kill using `pthread_kill` or `pthread_cancel`.

In either case, the adoption of E-ZAP-NOW significantly complicates an implementation in ways that are not obvious from the description of the formal system.

We can address the over-eagerness by restricting the reduction rule to only apply when the other process is ready:

$$\frac{\text{ready}(P, \bar{x}) \quad M \triangleq \text{fn}(P) \setminus \{\bar{x}\}}{(\nu x \bar{x})(w \dot{\downarrow} N, x \parallel P) \longrightarrow w \dot{\downarrow} N, M} \text{E-ZAP-WHEN-READY}$$

What does this semantics require of an implementation? Any process only needs to monitor the endpoints it is *ready* on, which is certainly an improvement. What is a zipper ready on? Consider the two processes

$$(a) \ (\nu x \bar{x})(a \dot{\downarrow} N, x \parallel \bar{x} \dot{\downarrow} M) \quad (b) \ (\nu x \bar{x})(a \dot{\downarrow} N, x \parallel b \dot{\downarrow} M, \bar{x})$$

For both of these processes to reduce by E-ZAP-WHEN-READY, a zipper $x \dot{\downarrow} N$ must be ready on all its endpoints— x and all the endpoints in N . This reveals an awkwardness: the endpoint x is not special. A zipper may introduce multiple endpoints of type $\bar{\cdot}$. From the perspective of the logic, there is no way to tell these apart. The only thing that matters is that there is at least one, and the principal purpose of the syntax $x \dot{\downarrow} N$ is to ensure that. Unfortunately, this makes reduction with E-ZAP-WHEN-READY non-confluent. The second process has two reductions, each of which singles out a different endpoint as special.

$$\begin{array}{ccc} & (\nu x \bar{x})(a \dot{\downarrow} N, x \parallel b \dot{\downarrow} M, \bar{x}) & \\ \swarrow & & \searrow \\ a \dot{\downarrow} N, M, b & & b \dot{\downarrow} N, M, a \end{array}$$

Note that Wadler’s semantics has the same problem, but, in his case, the non-confluence for the absurd offer is part of the general non-confluence caused by the commuting conversions (see § 2.3).

The solution is simple: treat all endpoints in a zipper the same. We alter the syntax for zippers to $\downarrow N$ and alter the typing and reduction rules to reflect the new syntax.

$$P, Q, R ::= \dots \mid x \downarrow N \mid \downarrow N$$

$$\frac{N = \text{fn}(\Gamma)}{\downarrow N, x \vdash \Gamma, x : \top} \text{T-ZAP} \quad \frac{\text{ready}(P, \bar{x}) \quad M \triangleq \text{fn}(P) \setminus \{\bar{x}\}}{(\nu x \bar{x})(\downarrow N, x \parallel P) \longrightarrow \downarrow N, M} \text{E-ZAP}$$

The critical pair is resolved. The two processes are definitionally equal:

$$\begin{array}{c} (\nu x \bar{x})(\downarrow N, a, x \parallel \downarrow M, b, \bar{x}) \\ \swarrow \quad \searrow \\ \downarrow N, M, a, b \quad \quad \quad \downarrow N, M, b, a \end{array}$$

Alternatively, if one wishes to avoid definitional set equality, the permutation of the endpoints in N could be arranged via the structural congruence. In this case, there is no need to add rules for contraction, since the typing rule guarantees the uniqueness of the endpoints in N .

To reflect the fact that a zipper must be ready on all its endpoints, we alter the definition of *ready*.

Definition 3.94 (Ready). A process P is ready to act on x , written $\text{ready}(P, x)$, if it is of one of the forms:

$$\begin{array}{l} x \leftrightarrow y \quad x[y].P \quad x[].P \quad x \triangleleft \text{inl}.P \quad x \triangleleft \text{inr}.P \quad \downarrow N, x \\ y \leftrightarrow x \quad x(y).P \quad x().P \quad x \triangleright \{\text{inl}: P; \text{inr}: Q\} \end{array}$$

A process P is ready if it is ready to act on some endpoint.

The addition of T-ZAP and E-ZAP preserves all the metatheoretical properties of HCP presented in this chapter:

- The proof of *preservation* requires an additional case for E-ZAP (see Proposition 3.95).
- The proof of *progress* requires no changes. All relevant changes are contained in the reduction lemma (see Lemma 3.96).
- The definitions of canonical form and the dependency and connection graphs are unchanged, as are the corresponding proofs of the adequacy of canonical forms, right-branching form, and disentanglement. All relevant changes are contained in the definition of *ready* (see Definition 3.94).
- The calculus continues to be in correspondence with CP, with the caveat that similar changes must be made to CP's syntax and semantics for zippers.
- Finally, we make similar changes to the label-transition system, such that harmony continues to hold.

We presented the updated proofs for preservation and progress and the updated label-transition system.

Proposition 3.95 (Preservation). *If $P \vdash \mathcal{G}$ and $P \longrightarrow Q$, then $Q \vdash \mathcal{G}$.*

Proof. By induction on the derivation of the reduction. The cases for E-LINK, E-SEND, E-CLOSE, E-SELECT₁, E-SELECT₂, E-EQUIV, and E-CONG, are as in Proposition 3.30.

- Case E-ZAP.

The reduction rule guarantees that P is ready on $\bar{x}$. Therefore, P must be typed under a single hyper-environment. Hence, the typing derivation is of the form:

$$\frac{\frac{\frac{N = \text{fn}(\Gamma)}{\not\downarrow N, x \vdash \Gamma, x : A} \text{T-ZAP} \quad P \vdash \Delta, \bar{x} : \bar{A}}{\not\downarrow N, x \parallel P \vdash \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}} \text{T-PAR}}{\frac{}{(vx\bar{x})(\not\downarrow N, x \parallel P) \vdash \Gamma, \Delta} \text{T-NEW}}$$

The result is well-typed by the following typing derivation:

$$\frac{N, M = \text{fn}(\Gamma, \Delta)}{\not\downarrow N, M \vdash \Gamma, \Delta} \text{T-ZAP}$$

The side condition $N, M = \text{fn}(\Gamma, \Delta)$ is satisfied. This follows from the side conditions of typing and reduction, and Proposition 3.18:

$$N = \text{fn}(\Gamma) \quad \text{and} \quad M \triangleq \text{fn}(P) \setminus \{\bar{x}\} = \text{fn}(\Delta)$$

The other condition is that there must be *some* endpoint $y : \top$ in Γ, Δ . By typing, there must be some endpoint $y : \top$ in $\Gamma, x : A$. There are two cases:

1. If $x \neq y$, then there exists some endpoint $y : \top \in \Gamma$.
2. If $x = y$, then $y : \top \notin \Gamma$. By duality, the endpoint $\bar{x}$ must have type **1**. As there are no introduction forms for **1**, the endpoint of type **1** must have been introduced by T-ZAP. Hence, there must be some endpoint $z : \top \in \Delta$.

The proof for subcase (2) relies on the consistency of HCLL, which we have not explicitly proven, but which follows from disentanglement and the consistency of CLL.

□

Lemma 3.96 (Reduction). *If $(vx\bar{x})(P \parallel Q) \vdash \mathcal{G}$, and P and Q are ready to act on x and $\bar{x}$, respectively, there exists some R such that $(vx\bar{x})(P \parallel Q) \longrightarrow R$.*

Proof. By inversion on the derivation of $(vx\bar{x})(P \parallel Q) \vdash \mathcal{G}$.

The ten cases that correspond to E-LINK, E-SEND, E-CLOSE, E-SELECT₁, E-SELECT₂, and their symmetric variants are as in Lemma 3.34.

There are two new cases, where either P or Q is a zapper, which correspond to E-ZAP and its symmetric variant under E-EQUIV with SC-PARCOMM.

□

Proposition 3.97 (Progress). *If $P \vdash \mathcal{G}$, then either P is in canonical form, or there exists some Q such that $P \rightarrow Q$.*

Proof. The proof is unchanged from Proposition 3.35, except that the appeal to Lemma 3.34 is replaced by an appeal to Lemma 3.96.

□

Let us extend the label-transition system and the proof of harmony. We extend the action labels with two new actions.

$$\pi ::= \dots \mid x \downarrow N \mid x \uparrow N$$

The action $x \downarrow N$ denotes sending a kill signal, and receiving the unused endpoints N from the killed process. The action $x \uparrow N$ denotes the dual, receiving a kill signal, and sending the unused endpoints N .

We add two corresponding action transition rules and one corresponding communication rule.

Action Rules	Communication Rules
$\begin{array}{c} \downarrow N, x \xrightarrow{x \downarrow M} \downarrow N, M \\ P \xrightarrow{x \uparrow N} 0 \\ \text{if ready}(P, x) \text{ and } \neg \text{zapper}(P) \\ \text{where } N \triangleq \text{fn}(P) \setminus \{x\} \end{array}$	$\begin{array}{c} \text{ACT-KILL} \\ \text{ACT-DIE} \\ \text{TAU-KILL-DIE} \\ \frac{P \xrightarrow{x \downarrow N \parallel \bar{x} \uparrow N} P'}{(\nu x \bar{x})P \xrightarrow{\tau} P'} \end{array}$

Proposition 3.98 (Harmony). *If $P \vdash \mathcal{G}$, then $P \xrightarrow{\tau} Q \iff P \rightarrow Q$.*

Proof. The proof proceeds as Proposition 3.93. In case $(\Rightarrow)$, there is an additional case for TAU-KILL-DIE with ACT-KILL and ACT-DIE. In case $(\Leftarrow)$, there is an additional case for E-ZAP.

□

A Zapper Is Made of Other, Smaller Zappers

A zapper must be ready to act on all its endpoints in parallel. Hence, we should think of the zapper $\downarrow x_1, \dots, x_n$ as a representation of the parallel composition of individual zappers for each endpoint, i.e. $\downarrow x_1 \parallel \dots \parallel \downarrow x_n$. Under this view, the permutation of the endpoints in a zapper follows

from the permutation of parallel processes under structural congruence, which is rather pleasing. This is exactly how zappers are represented in Fowler’s Exceptional GV [EGV, Fowler et al., 2019]. However, this representation is not well-typed in HCP:

$$(a) \frac{}{\textcolor{red}{\cancel{x}} a, b \vdash a : \textcolor{blue}{T}, b : A} \quad \frac{\frac{}{\textcolor{red}{\cancel{x}} a \vdash a : \textcolor{blue}{T}} \quad \frac{}{\textcolor{red}{\cancel{x}} b \vdash b : A}}{\textcolor{red}{\cancel{x}} a \parallel \textcolor{red}{\cancel{x}} b \vdash a : \textcolor{blue}{T} \parallel b : A} (b)$$

To type (b), we would have to change the local constraint that at least one endpoint introduced by the zapper must have type $\textcolor{blue}{T}$ to the *non-local* constraint that at least one zapper must introduce an endpoint with type $\textcolor{blue}{T}$. Since EGV is not interested in preserving consistency and allows exceptions to occur in any context, it does not require this restriction.

To assign (a) and (b) the same type, we would have to admit MIX , which breaks the connection structure of HCP and the correspondence with CLL. Since EGV is not interested in preserving that correspondence for its runtime terms, it admits MIX in its runtime type system.

While the representation $\textcolor{red}{\cancel{x}} x_1 \parallel \dots \parallel \textcolor{red}{\cancel{x}} x_n$ is not a good fit for HCP, it provides an intuition for the implementation of zappers. A zapper is not a single monolithic process that takes ownership of all the unused resources of the processes it kills, and the killed process should not transfer ownership of all their unused endpoints to such a monolithic zapper. Zappers are small processes responsible for cleaning up individual resources. Once a zapper has successfully delivered a kill signal, it may terminate. The killed process, upon receiving a kill signal, simply spawns off one zapper process for each of its unused resources, and then terminates.

Could You to Tell Me When You Are Done?

In this section, we have described an exceptional semantics for the absurd offer in which all cancelled endpoints are treated equally. In this final portion of the section, I hope to convince you that there is no well-behaved interpretation that treats one endpoint specially, as implied by the syntax used in Chapter 2 and Chapter 3, and that this is a direct consequence of the fact that the logical rule corresponding to T-ABSURD can introduce multiple propositions $\textcolor{blue}{T}$ which it cannot distinguish.

Let us consider an alternative interpretation for the zapper $x \textcolor{red}{\cancel{x}} N$ where the endpoint x is special:

- When some process P receives a kill signal on x , it becomes the zapper $x \textcolor{red}{\cancel{x}} N$, where N contains its unused endpoints, i.e. $N = \text{fn}(P) \setminus \{x\}$.
- A zapper $x \textcolor{red}{\cancel{x}} N$ does the following for each endpoint in N in parallel:
 - It sends a kill signal.

- It waits to receive a notification to confirm that the process on the other side has finished cleaning up and has terminated.
- It closes the channel, and removes the endpoint from N .
- When N is empty, the zipper notifies the process on the other side of x that it has finished cleaning up, and terminates.

Such an interpretation is useful. For instance, if we implement processes as OS threads, then the main thread must not terminate before all the child threads have finished cleaning up.

Under this interpretation, a zipper $x \downarrow N$ is ready on all the endpoints in N in parallel, but is blocked on x until N is empty. Consider the following reduction rules:

$$\begin{array}{ll}
 (vx\bar{x})(w \downarrow N, x \parallel \bar{x} \downarrow \emptyset) \longrightarrow w \downarrow N & \text{E-ZAP-NOTIFY} \\
 (vx\bar{x})(w \downarrow N, x \parallel P) \longrightarrow (vx\bar{x})(w \downarrow N, x \parallel \bar{x} \downarrow M) & \text{E-ZAP-KILL} \\
 & \text{if ready}(P, \bar{x}) \text{ and } \neg \text{zipper}(P) \\
 & \text{where } M \triangleq \text{fn}(P) \setminus \{\bar{x}\}
 \end{array}$$

These rules preserve typing. The rule E-ZAP-KILL rebinds x and $\bar{x}$ with types $x : \mathbf{0}$ and $\bar{x} : \mathbf{T}$, but this is usual for session types. However, they are only well-behaved under certain circumstances. Consider the two processes

$$(a) \ (vx\bar{x})(a \downarrow N, x \parallel \bar{x} \downarrow M) \quad (vx\bar{x})(a \downarrow N, x \parallel b \downarrow M, \bar{x}) \ (b)$$

If a process up to structural congruence does not contain (b), then progress, termination, and confluence follow by an argument from the structure of the connection graph. To simplify matters, let us consider a fully-connected process with a single zipper

$$\mathcal{C}[a \downarrow N, P_1, \dots, P_n]$$

where each P_i is ready on some bound endpoint x_i . The process reduces to a single zipper in $2n$ steps.

1. By n applications of E-ZAP-KILL, following the structure of the connection tree outward from the zipper $a \downarrow N$, we convert each process P_i to a zipper $x_i \downarrow N_i$ where $N_i = \text{fn}(P_i) \setminus \{x_i\}$.
2. If P_i is a leaf of the connection graph, it is of the form $x_i \downarrow \emptyset$. By n application of E-ZAP-NOTIFY, following the structure of the connection tree inwards from the leaves, we convert the entire process to $a \downarrow M$, where M contains all the free endpoints from N and each N_i .

If some P_i is blocked on a free endpoint, reduction becomes blocked as usual. If the process contains multiple zippers, but the connection structure for these zippers is always as in (a), then that merely reduces the number of uses of E-ZAP-KILL that are needed. If the process is not

fully-connected, any component with at least zipper reduces to become a single zipper, and the remaining components reduce as usual.

Under this interpretation, $\mathbf{0}$ and $\mathbf{T}$ are assigned dual behaviours—an endpoint of type $\mathbf{0}$ is used to send a kill signal and then wait for confirmation, and an endpoint of type $\mathbf{T}$ is used, dually, to receive a kill signal, and then send confirmation and terminate—and therein lies the problem. The typing derivation for process (b)

$$(\nu x \bar{x})(a \dot{\downarrow} N, x \parallel b \dot{\downarrow} M, \bar{x})$$

could assign $x : \mathbf{0}$ and $\bar{x} : \mathbf{T}$, or vice versa, but by E-ZAP-KILL, they will be treated the same, not dually. The type system does not have sufficient structure to ensure the desired dual behaviour.

3.3.2 Synchronous Links and Lazy Forwarders

In this section, we revisit the semantics of the link. In § 2.1.2, we mentioned that the semantics of link does not explicitly forward messages, but rather treats links as a suspended α -renaming. Recall that the reduction rule for link is as follows:

$$(\nu x \bar{x})(x \leftrightarrow y \parallel P) \longrightarrow P\{\bar{x}/y\}$$

Let us consider what requirements this places on an implementation. The process P is not required to be ready to act on the endpoint $\bar{x}$, which means that link cannot be implemented in terms of synchronous message passing. In essence, link is asynchronous. For instance, if every channel has an associated message buffer, link could be implemented by asynchronously sending a redirection notice.

There is a tension between the asynchronous semantics of link, and the synchronous semantics of all other actions in CP. This is apparent in the definition of canonical form and its adequacy, which have separate cases for link and for all other actions. There are two options to resolve this tension: we either make link synchronous or all other actions asynchronous. CP is easily adapted to be asynchronous [see, e.g. the treatment of synchronous and asynchronous GV in Fowler, 2019]. In this section, we will consider our options for a synchronous link.

Blocking Link

The easiest approach to making link synchronous is to require that the other process is ready on the relevant endpoint.

$$\frac{\text{ready}(P, \bar{x})}{(\nu x \bar{x})(x \leftrightarrow y \parallel P) \longrightarrow P\{\bar{x}/y\}} \text{E-LINK-READY}$$

The new rule allows strictly fewer reductions, and as such does not affect the proof of preservation. However, it simplifies the definitions of canonical form and the proof of progress, and strengthens the adequacy of canonical forms.

The definition of canonical form no longer requires condition (1).

Definition 3.99 (Canonical Form). A process P is in canonical form, written $\text{canonical}(P)$, if P is of the form $\mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$) and (for $1 \leq i, j \leq n$) no P_i and P_j are ready to act on dual endpoints $\{x, \bar{x}\} \in \text{dn}(\mathcal{C}[\cdot])$.

The proof of progress for HCP (Proposition 3.35) has two cases: either condition (1) fails, or condition (2) fails. The proof of progress for HCP with E-LINK-READY no longer requires the first case, as link reduction is fully covered by the second case.

Proposition 3.100 (Progress). If $P \vdash \mathcal{G}$, then either P is in canonical form, or there exists some Q such that $P \rightarrow Q$.

The adequacy of canonical forms is strengthened by the change. In HCP with E-LINK-READY, “blocked on free endpoints” fully characterises canonical forms.

Corollary 3.101. If $P \vdash \top$, then $\text{canonical}(P) \iff \text{blocking}(P) \subseteq \text{fn}(P)$.

Identity Expansion

We could implement a synchronous link using the process calculus equivalent of *identity expansion*—the procedure that rewrites proofs in the logic to remove non-atomic uses of the axiom. Under this view, link is a macro which statically computes the link process from the type, using the following expansions:

$$\begin{aligned}
 & \overset{A \otimes B}{x \leftrightarrow y}, \overset{A \otimes B}{y \leftrightarrow x} \triangleq x(z). y[w]. (z \overset{A}{\leftrightarrow} w \parallel x \overset{B}{\leftrightarrow} y) \\
 & \overset{\perp}{x \leftrightarrow y}, \overset{1}{y \leftrightarrow x} \triangleq x(). y[]. 0 \\
 & \overset{A \& B}{x \leftrightarrow y}, \overset{A \oplus B}{y \leftrightarrow x} \triangleq x \triangleright \{ \text{inl}: y \triangleleft \text{inl}. x \overset{A}{\leftrightarrow} y; \text{inr}: y \triangleleft \text{inr}. x \overset{B}{\leftrightarrow} y \} \\
 & \overset{\top}{x \leftrightarrow y}, \overset{0}{y \leftrightarrow x} \triangleq x \downarrow y
 \end{aligned}$$

(The type written over the arrow is the type of the left endpoint.)

Identity expansion works when the type is statically known, but it breaks when we add polymorphism, and repairing it requires us to compute the link process dynamically and to keep session types around at runtime.

From the perspective of an implementation, identity expansion also inflates the size of the program, from a single link instruction to a number of instructions equal to the size of the type.

Action Rules

$x \leftrightarrow y \xrightarrow{x()} y[], 0$	FWD-CLOSE
$x \leftrightarrow y \xrightarrow{x(z)} y[w]. (z \leftrightarrow w \parallel x \leftrightarrow y)$	FWD-SEND
$x \leftrightarrow y \xrightarrow{x \triangleright \text{inl}} y \triangleleft \text{inl}. x \leftrightarrow y$	FWD-SELECT ₁
$x \leftrightarrow y \xrightarrow{x \triangleright \text{inr}} y \triangleleft \text{inr}. x \leftrightarrow y$	FWD-SELECT ₂
$x \leftrightarrow y \xrightarrow{y()} x[], 0$	BWD-CLOSE
$x \leftrightarrow y \xrightarrow{y(z)} x[w]. (z \leftrightarrow w \parallel x \leftrightarrow y)$	BWD-SEND
$x \leftrightarrow y \xrightarrow{y \triangleright \text{inl}} x \triangleleft \text{inl}. x \leftrightarrow y$	BWD-SELECT ₁
$x \leftrightarrow y \xrightarrow{y \triangleright \text{inr}} x \triangleleft \text{inr}. x \leftrightarrow y$	BWD-SELECT ₂

Figure 3.5: Label-Transition System for Lazy Forwarders

The Lazy Forwarder

We could implement a synchronous link that performs identity expansion, not based on the type, but based on the messages the link receives. In response to a send action, the link unfolds into a receive action followed by a send action, and the process reduces by E-SEND.

$$\begin{aligned}
 (\nu x \bar{x})(x[z]. P \parallel \bar{x} \leftrightarrow y) &\triangleq (\nu x \bar{x})(x[z]. P \parallel \bar{x}(\bar{z}). y[w]. (\bar{z} \leftrightarrow w \parallel \bar{x} \leftrightarrow y)) \\
 &\longrightarrow (\nu x \bar{x})(\nu z \bar{z})(P \parallel y[w]. (\bar{z} \leftrightarrow w \parallel \bar{x} \leftrightarrow y))
 \end{aligned}$$

The operational semantics of the lazy forwarder are easier to understand from the label-transition system, since the reduction rules fold this two-step behaviour into single rule. The label-transition rules for lazy forwarders are in Figure 3.5.

The lazy forwarder has the same semantics as identity expansion, but works in the presence of polymorphism. Lazy forwarders are synchronous, rather than asynchronous, but they behave differently from blocking links. They are directed, as they only forward received messages. Hence, the following process is stuck:

$$(\nu x \bar{x})(x \leftrightarrow y \parallel \bar{x}(z). P)$$

This is a direct consequence of the correspondence between the process calculus and the sequent calculus for CLL, since the expansion that sends and then receives—i.e. $y[w]. x(z). (z \leftrightarrow w \parallel x \leftrightarrow y)$ —is not typeable. However, with delayed actions [Kokke et al., 2019a], the process can send and then receive, regardless of the order of the two actions. Hence, under delayed actions lazy forwarders and blocking link may behave the same.

Lazy forwarders are type preserving by the correctness of identity expansion, and satisfy progress as well as the adequacy of canonical

forms. However, the canonical forms are quite different, due to the directionality of lazy forwarders.

Firstly, we must alter the definition of ready, since lazy forwarders are only ready on the endpoint on which they receive, which can be inferred from the types of its endpoints.

Definition 3.102 (Ready). A process P is ready to act on x , written $\text{ready}(P, x)$, if it is of one of the forms:

$$\begin{array}{ccccccc}
 \begin{array}{c} A \otimes B \\ x \leftrightarrow y \end{array} & \begin{array}{c} \perp \\ x \leftrightarrow y \end{array} & \begin{array}{c} A \& B \\ x \leftrightarrow y \end{array} & \begin{array}{c} T \\ x \leftrightarrow y \end{array} & x[y].P & x[].P & x \triangleleft \text{inl}.P \quad x \triangleleft \text{inr}.P \quad x \not\leq N \\
 \begin{array}{c} A \otimes B \\ y \leftrightarrow x \end{array} & \begin{array}{c} 1 \\ y \leftrightarrow x \end{array} & \begin{array}{c} A \oplus B \\ y \leftrightarrow x \end{array} & \begin{array}{c} 0 \\ y \leftrightarrow x \end{array} & x(y).P & x().P & x \triangleright \{\text{inl}: P; \text{inr}: Q\}
 \end{array}$$

A process P is ready if it ready to act on some endpoint.

The definition of canonical form and the proof of progress simplify in the same manner as for blocking links. However, while the text defining canonical form is the same, the actual canonical forms are quite different, due to the different definitions of *ready*. For instance, the previously mentioned stuck process is in canonical form:

$$(\nu x \bar{x})(\overset{A \otimes B}{x \leftrightarrow y} \parallel \bar{x}(z).P)$$

While type annotations are unnecessary at runtime, determining whether or not a lazy forwarder is ready requires case analysis on its type to determine its direction. For the process above, we can infer that the endpoint x has type $A \otimes B$, and hence the link is only ready on y .

The asymmetry of lazy forwarders allows us to simplify the definition of the dependency graph, by removing the special case for link. Links no longer generate undirected edges, but directed arcs which align with their direction:

$$\text{Dep}\left(\overset{A \otimes B}{x \leftrightarrow y}\right) = \overrightarrow{xy} \quad \text{Dep}\left(\overset{A \oplus B}{x \leftrightarrow y}\right) = \overrightarrow{yx}$$

The adequacy of canonical forms is strengthened in the same manner as with blocking links, i.e. “blocked on free endpoints” fully characterises canonical forms.

Lazy forwarders continue to work in the presence of polymorphism, since the direction of the link is determined dynamically, in response to communication.

3.3.3 Variant Types and Guarded Summation

In this section, we generalise select and offer to variant types, as mentioned in § 2.1.5, then decompose the n -ary offer into guarded summation, which lets us factor out actions into their own syntactic

sort and which enormously simplifies the reduction and label-transition semantics, as well as the associated metatheory. The decomposition guarantees that summations are guarded using a limited form of focusing, which is a technique from proof theory to control the structure of proofs without reducing the expressivity of the logic [see Andreoli, 1992].

The Variant With Variants

To generalise select and offer to variant types, we remove the restriction that labels must always be drawn from $\{\text{inl}, \text{inr}\}$ [following Dardha and Gay, 2018]. Let L range over finite sets of labels, and let l and k range over individual labels (not to be confused with ℓ , which ranges over the entirely unrelated labels used in the label-transition system). We write L, l to mean $L \cup \{l\}$. Labels occur both in processes and in types, and as such appear both in blue and red. They are coloured accordingly, and when they appear by themselves, they are coloured according to whether or not they can be erased. Regardless of colour, labels with the same name refer to the same label.

We replace the binary select and offer actions and the corresponding types with n -ary labelled variants:

$$\begin{array}{ll}
 P, Q, R ::= & \dots \\
 & | \cancel{x \triangleleft \text{inl}. P} \mid \cancel{x \triangleleft \text{inr}. P} \\
 & | \cancel{x \triangleright \{\text{inl}: P; \text{inr}: Q\}} \\
 & | x \triangleleft l. P \\
 & | x \triangleright \{l: P\}_{l \in L} \\
 \\
 A, B ::= & \dots \\
 & | A \oplus B \\
 & | \cancel{A \& B} \\
 & | \bigoplus \{l: A_l\}_{l \in L} \\
 & | \big\& \{l: A_l\}_{l \in L}
 \end{array}$$

The variant offers $x \triangleright \{l: P_l\}_{l \in L}$ offers a set of alternatives labelled by the labels from some finite set L , and the variant selection $x \triangleleft l. P$ selects the alternative labelled by l . The variant selection and offer types, $\bigoplus \{l: A_l\}_{l \in L}$ and $\big\& \{l: A_l\}_{l \in L}$, respectively, associate each label $l \in L$ with the type of the alternative.

The syntax $\{l: P_l\}_{l \in L}$ and $\{l: A_l\}_{l \in L}$ denote sets of labelled alternatives. Formally, these are functions from labels to processes and session types, respectively. Note that the occurrences of l and l between the curly braces are not free, but bound by the expressions in the subscripts.

For concrete processes and types, we write the set of alternatives in full, e.g. as $x \triangleright \{l_1: P_1; \dots; l_n: P_n\}$. In these cases, we omit the subscript, since the set of labels can be inferred.

The choice to treat alternatives as sets means that they can be reordered. For instance, $x \triangleright \{\text{inl}: P; \text{inr}: Q\}$ and $x \triangleright \{\text{inr}: Q; \text{inl}: P\}$ are definitionally equal. Alternatively, we could treat the sets of labels as ordered, and treat alternatives as ordered sequences.⁴

⁴Dardha and Gay [2018] use a different construction, where both the label and the

The typing rules and reduction semantics for variant selection and offer are standard generalisations of the binary versions.

$$\begin{array}{c}
\text{T-SELECT} \\
\frac{P \vdash \Gamma, x : A_i \quad i \in L}{x \triangleleft i. P \vdash \Gamma, x : \bigoplus \{i : A_i\}_{i \in L}} \\
\\
\text{T-OFFER} \\
\frac{(P_i \vdash \Gamma, x : A_i)_{i \in L}}{x \triangleright \{i : P_i\}_{i \in L} \vdash \Gamma, x : \& \{i : A_i\}_{i \in L}} \\
\\
(\nu x \bar{x})(x \triangleleft i. P \parallel \bar{x} \triangleright \{i : Q_i\}_{i \in L}) \longrightarrow (\nu x \bar{x})(P \parallel Q_i) \quad \text{E-SELECT}
\end{array}$$

In the premise of T-OFFER, the notation “ $(\dots)_{i \in L}$ ” means that one typing derivation is required for each label and alternative. Should we want to support type inference, then every variant selection on a free endpoint requires one type annotation for each label not selected, though this is no different from the binary left and right selection actions.

There is an awkwardness between variant types and the absurd offer. The absurd offer should be the nullary variant, which would let us replace $\mathbf{0}$ by $\bigoplus \{\}$, $\top$ by $\& \{\}$, and the absurd offer by $x \triangleright \{\}$. As discussed in § 3.1.6 and § 2.1.6, this works in HCP with the inert semantics and in CP with both the inert semantics and Wadler’s commuting conversions, though, in both cases, it complicates the statement of linearity. As discussed in § 2.1.6 and § 3.3.1, zappers require that the names of the discarded endpoints are kept at runtime. Morally, zappers are the nullary offer, and the expected propositions hold, e.g. $\& \{i_1 : A; i_2 : \top\}$ is equivalent to A . However, zappers require more information than would follow from the nullary case of the generalisation. To combine variant types with zappers, we must rule out empty variants and retain the additive units.

The metatheory of HCP generalises easily to variant types, and all metatheoretical properties of HCP are preserved. In some cases the proofs simplify, as the variant syntax lets us treat selections uniformly. However, we will not belabour this point, as the generalisation to variant types is not novel, and is principally intended as an introduction to the discussion of guarded summation.

Focusing and Guarded Summation

The π -calculus usually defines the syntax of actions separately from the syntax of processes, and the two are combined by prefixing, i.e. $\pi.P$. This leads to a much simpler theory, as we can treat all actions uniformly. As mentioned in § 2.1.5 and § 3.2.8, this cannot easily be done in CP and HCP, as the offer is a single monolithic construct that combines all labels and alternative processes.

In the π -calculus, alternatives are combined by *summation*, e.g. the process $P + Q$ acts either as P or as Q . To ensure that the choice is not

alternative are indexed by some index i drawn from some set I , e.g. $\{i : P_i\}_{i \in I}$. They do not specify whether these denote a pair of functions from indices to labels and alternatives, or whether these denote ordered sequences of pairs.

arbitrary, each alternative must be *guarded*, which means that it must be ready to act on some distinct endpoint, and the first alternative to receive a message is selected. For instance, when evaluating the process

$$x[].0 \parallel (x().P + y().Q)$$

the communication on x selects P and discards Q . Such a process calculus, where the selection is made by the choice of endpoint, is difficult to type in formalisms like CP, where endpoints are linear and channels must be bound by a name restriction. (Summation naturally corresponds to a structural *with*, but the introduction of additive hypersequents complicates name restriction, which must account for the fact that each alternative may use some channel at a different type.)

In this section, we introduce an intermediate system, where alternatives are combined by summation, and each alternative must be guarded by an offer action on the same endpoint, but with a different label. Actions are defined as usual, e.g. as in § 3.2.8, with the exception that selection and offer permit arbitrary labels. (We postpone the discussion of link.)

$$\pi ::= x[y] \mid x(y) \mid x[] \mid x() \mid x \triangleleft l \mid x \triangleright l$$

Processes are defined by name restriction, parallel composition, the terminated process, *prefixing*, and *summation*.

$$P, Q, R ::= (\nu x \bar{x})P \mid P \parallel Q \mid 0 \mid P + Q \mid \pi.P$$

The process $x \triangleright l.P$ introduces a unary offer, and offers can be combined by summation. The typing rules for the offer and summation use a limited form of *focusing* [see Andreoli, 1992]. Focusing is a technique from proof theory that can be used to heavily restrict the form of proofs without affecting expressivity, and is often used for more efficient proof search. In our case, focusing is only used to enforce guardedness. We add a new typing judgement, $P \vdash \Gamma \uparrow x : A$, which remembers what endpoint the process is ready to make an offer on—or, if we focus all actions, ready to act on.⁵

$$\begin{array}{c} \text{T-OFFER}_1 \\ \frac{P \vdash \Gamma, x : A}{x \triangleright l.P \vdash \Gamma \uparrow x : \&\{l : A\}} \end{array} \quad \begin{array}{c} \text{T-ABSURD} \\ \frac{N = \text{fn}(\Gamma)}{x \not\triangleleft N \vdash \Gamma \uparrow x : \&\{\}} \end{array} \quad \begin{array}{c} \text{T-FOCUS} \\ \frac{P \vdash \Gamma \uparrow x : A}{P \vdash \Gamma, x : A} \end{array}$$

$$\begin{array}{c} \text{T-SUM} \\ \frac{P \vdash \Gamma \uparrow x : \&\{l : A_l\}_{l \in L} \quad Q \vdash \Gamma \uparrow x : \&\{l : A_l\}_{l \in L'}}{P + Q \vdash \Gamma \uparrow x : \&\{l : A_l\}_{l \in L \cup L'}} \end{array}$$

The rule T-OFFER₁ types the unary offer. It is the unary case of the rule for the variant offer, T-OFFER, except that it remembers what endpoint the

⁵The upward arrow in “ $P \vdash \Gamma \uparrow x : A$ ” comes from Andreoli’s sequent focused on an asynchronous proposition, which has a dual focused on a synchronous proposition that uses a downward arrow.

offer is made on. The rule T-FOCUS allows us to forget this information at any point, though, once forgotten, we cannot recover it. The rule T-SUM types guarded summation, which requires that each alternative is ready to make an offer on the same endpoint.

We extend the structural congruence with the following rules, such that summations are associative and commutative, and the absurd offer is the unit for summation:

$$\begin{array}{ll}
 P + x \not\downarrow N & \equiv P & \text{SC-SUMABSURD} \\
 P + Q & \equiv Q + P & \text{SC-SUMCOMM} \\
 P + (Q + R) & \equiv (P + Q) + R & \text{SC-SUMASSOC}
 \end{array}$$

The reduction rule for selection chooses the alternative with the matching label, and discards the other alternatives.

$$(\nu x \bar{x})(x \triangleleft l. P \parallel (\bar{x} \triangleright l. Q + R)) \longrightarrow (\nu x \bar{x})(P \parallel Q) \quad \text{E-SELECT}$$

The metatheory of HCP generalises easily to the variant with guarded summation. The proofs of preservation for the structural congruence and reduction must be updated by adding cases for the new rules. The proofs for progress, deadlock freedom, and the adequacy of canonical forms are easily updated. Alternatively, these properties follow from the operational correspondence between HCP with binary choice and HCP with guarded summation. We present the updated proofs of preservation and a sketch for the operational correspondence.

The rules SC-SUMABSURD, SC-SUMCOMM, SC-SUMASSOC, and E-SELECT preserve types.

Lemma 3.103. *If $P \equiv Q$, then $P \vdash \mathcal{G}$ if and only if $Q \vdash \mathcal{G}$.*

Proof. By induction on the derivation of the equivalence $P \equiv Q$.

The cases for reflexivity, symmetry, transitivity, and applications of SC-LINKCOMM, SC-PARNIL, SC-PARCOMM, SC-PARASSOC, SC-NEWCOMM, and SC-SCOPEEXT are as in Lemma 3.27. The case for congruence closure follows, similarly to Lemma 3.27, by induction and the injectivity of the type derivation rules.

- Case SC-SUMABSURD.

$$\begin{array}{c}
 \frac{P \vdash \Gamma \uparrow x : \{A_i\}_{i \in I} \quad \frac{N = \text{fn}(\Gamma) \quad x \not\downarrow N \vdash \Gamma \uparrow x : \{\}}{P \parallel x \not\downarrow N \vdash \Gamma \uparrow x : \{A_i\}_{i \in I}}}{P \vdash \Gamma \uparrow x : \{A_i\}_{i \in I}} \\
 \text{III}
 \end{array}$$

- Case SC-SUMCOMM.

$$\frac{P \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L} \quad Q \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L'}}{P + Q \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L \cup L'}}$$

III

$$\frac{Q \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L'} \quad P \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L}}{Q + P \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L \cup L'}}$$

- Case SC-SUMASSOC.

$$\frac{P_1 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_1} \quad \frac{P_2 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_2} \quad P_3 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_3}}{P_2 + P_3 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_2 \cup L_3}}}{P_1 + (P_2 + P_3) \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_1 \cup L_2 \cup L_3}}$$

III

$$\frac{\frac{P_1 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_1} \quad P_2 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_2}}{P_1 + P_2 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_1 \cup L_2}} \quad P_3 \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_3}}{P_1 + (P_2 + P_3) \vdash \Gamma \uparrow x : \&\{l:A_i\}_{i \in L_1 \cup L_2 \cup L_3}}$$

□

Proposition 3.104 (Preservation). *If $P \vdash \mathcal{G}$ and $P \longrightarrow Q$, then $Q \vdash \mathcal{G}$.*

Proof. By induction on the derivation of the reduction. The cases for E-SEND, E-CLOSE, E-EQUIV, and E-CONG, are as in Proposition 3.30. The case for E-LINK, if included, follows by similarly to Proposition 3.30.

- Case E-SELECT.

By inversion, the typing derivation for the summation $\bar{x} \triangleright l. Q + R$ is of the form

$$\frac{\frac{Q \vdash \Delta, \bar{x} : \bar{A}_k}{\bar{x} \triangleright k. Q \vdash \Delta \uparrow \bar{x} : \&\{k:\bar{A}_k\}} \quad R \vdash \Delta \uparrow \bar{x} : \&\{l:\bar{A}_i\}_{i \in L}}{\bar{x} \triangleright k. Q + R \vdash \Delta \uparrow \bar{x} : \&\{l:\bar{A}_i\}_{i \in L, k}}}{\bar{x} \triangleright k. Q + R \vdash \Delta, \bar{x} : \&\{l:\bar{A}_i\}_{i \in L, k}}$$

By inversion, the typing derivation for the left-hand side of the reduction is of the form

$$\frac{\frac{P \vdash \Gamma, x : A_k}{x \triangleleft k. P \vdash \Gamma, x : \oplus\{l:A_i\}_{i \in L, k}} \quad \bar{x} \triangleright k. Q + R \vdash \Delta, \bar{x} : \&\{l:\bar{A}_i\}_{i \in L, k}}{x \triangleleft k. P \parallel (\bar{x} \triangleright k. Q + R) \vdash \Gamma, x : \oplus\{l:A_i\}_{i \in L, k} \parallel \Delta, \bar{x} : \&\{l:\bar{A}_i\}_{i \in L, k}}}{(\nu x \bar{x})(x \triangleleft k. P \parallel (\bar{x} \triangleright k. Q + R)) \vdash \Gamma, \Delta}$$

The result follows as

$$\frac{\frac{P \vdash \Gamma, x : A_k \quad Q \vdash \Delta, \bar{x} : \bar{A}_k}{P \parallel Q \vdash \Gamma, x : A_k \parallel \Delta, \bar{x} : \bar{A}_k}}{(\nu x \bar{x})(P \parallel Q) \vdash \Gamma, \Delta}$$

□

The operational correspondence follows disentanglement, by translating each summation to a sequence of binary choices. The translation on processes proceeds as follows:

1. Take the maximum summation context, which is defined by analogy to the maximum configuration context.
2. Eliminate any superfluous absurd offers, which ensures that any use of SC-SUMABSURD holds reflexively under the translation.
3. Normalise the structure of the summation to right-branching form.
4. Fix an arbitrary order on all labels and reorder the unary offers accordingly, which ensures that any use of SC-SUMCOMM and SC-SUMASSOC holds reflexively under the translation.
5. Translate the summation as a series of binary offers.

The translation on session types proceeds similarly, and the order on labels ensures that the translation of the session types matches the translation of the processes.

The translation on processes preserves types—up to the translation on types—and preserves structural congruence and reduction. However, the correspondence between the reduction semantics is not one-to-one. A process that reduces in one step with E-SELECT requires a number of steps with E-SELECT₁ and E-SELECT₂ that is worst-case linear in the number of alternatives. The worst-case linear increase is a consequence of the right-branching form, and we can improve to a logarithmic increase by translating summations as balanced binary trees.

What About Zappers? The exceptional semantics for the absurd offer, as discussed in § 3.3.1, are compatible with guarded summation. The rule E-ZAP-KILL works as written, and it is important *not* to permit any other actions in the summation.

What About Link? There are no issues with link in HCP with guarded summation, but there is a choice on how to implement link. We can either add link as a process or add it as an action:

- *As a process.* From the perspective of the logic, link should be a process constructor, since link, like name restriction and parallel composition, corresponds to a structural rule of the logic, whereas the actions correspond to logical rules. This prevents us from treating link together with the other actions. In HCP, where link is asynchronous and the other actions are synchronous, we are already required to treat link separately. Hence, for HCP, adding link as a process constructor is a good fit.
- *As an action.* On the other hand, when using a synchronous semantics for link, such as those presented in § 3.3.2, we can treat it

together with the other actions. Hence, for HCP with a synchronous link and for HCP with asynchronous actions, adding link as an action might be a good fit as well, even if it forces us to write “ $x \leftrightarrow y.0$ ”.

Processes, Summations, and an Absurd Unit Sangiorgi and Walker [2003] define processes and summations as separate syntactic sorts

$$\begin{aligned} P, Q &::= S \mid P \parallel Q \mid (\nu x)P \\ S, R &::= 0 \mid \pi.P \mid S + R \end{aligned}$$

which guarantees that summations are guarded *syntactically*. There is no great need for such separation in HCP, as guardedness is guaranteed by the type system. However, examining Sangiorgi and Walker’s definition reveals two interesting things.

Firstly, the rule T-FOCUS corresponds to the syntactic embedding of summations into processes. If we focus the existing typing rules for actions, the typing judgements divide neatly into separate typing judgements for processes and summations

$$\begin{aligned} P &\vdash \mathcal{G} && \text{process typing} \\ S &\vdash \Gamma \uparrow x : A && \text{summation typing} \end{aligned}$$

with the rules for prefixing and the rule T-FOCUS moving back and forth between summation typing and process typing.

Secondly, Sangiorgi and Walker [2003] define 0 as a *summation*. Since summations are processes, it does double duty as both the empty summation and the empty process, i.e. as both the unit for summation and parallel composition. In HCP with guarded summation, the absurd offer $x \not\downarrow N$ is the unit for summation, whereas the terminated process 0 is the unit for parallel composition.

3.3.4 Channel or Endpoint Names?

Should names refer to communication channels or to their endpoints? I already revealed my hand in Chapter 2, where CP’s names refer to channel endpoints. However, for CP, the choice is not hugely significant. Case in point, names refer to channels in Wadler’s CP.

Under the restrictive view, channel names are a natural choice. Endpoint names are incompatible with the restrictive view, as there is no mechanism to tell us which endpoints are connected to the same channel. We could introduce such a mechanism. We could use co-names, where the endpoint name for sending is computed from the endpoint name for receiving by the overbar function, e.g. if x is the endpoint for receiving, then $\bar{x}$ is the endpoint for sending. We could generalise, and compute

the endpoints for any number of participants in a multiparty session from one special endpoint name. However, such a mechanism invariably requires one special endpoint name from which to compute the others, and is there really any difference between such a special endpoint name and a channel name?

Under the creative view, we have a legitimate choice between channel and endpoint names. Let us consider our options with a few questions.

How to Connect Two Unconnected Parallel Processes?

Consider the following example:

$$a\langle c \rangle. P \parallel b(y). Q$$

With channel names, we must rename one of the channels so that their names coincide, then use a name restriction to bind that name—e.g. rename b to a , then use the name restriction (va) .

$$a\langle c \rangle. P \parallel b(y). Q \quad \text{to} \quad a\langle c \rangle. P \parallel a(y). Q \quad \text{to} \quad (va)(a\langle c \rangle. P \parallel a(y). Q)$$

With endpoint names, we only need to use the name restriction (vab) .

$$a\langle c \rangle. P \parallel b(y). Q \quad \text{to} \quad (vab)(a\langle c \rangle. P \parallel b(y). Q)$$

With channel names, we need renaming to connect the two processes—an operation in the meta-language, which mutates the process to boot—but with endpoint names, name restriction internalises the operation of connection into the object language.

What Is in A Name?

With channel names, some form of role annotation is necessary to ensure that the various uses of a channel name are coherent.

As discussed, $L\pi$ annotates session types with their role—whether the corresponding channel is used to send, receive, neither, or both. The typing rule for parallel composition checks that the various roles are coherent—i.e. at most one send and one receive. The typing rule for name restriction checks that all roles are fulfilled—i.e. there is exactly one send and one receive.

In Lindley and Morris' GV, similar role annotations are used. The typing rule for parallel composition checks that the various roles for *one* channel are coherent, and requires all other channels names are unique. It combines one positive and one negative use of one channel name into one locked use of that channel—i.e. for exactly one channel name x , $x : S$ and $x : \bar{S}$ are combined into $x : S^\#$. The typing rule for name restriction checks that all roles are fulfilled—i.e. that the channel is locked.

An early version of HCP [Kokke et al., 2019b, with errata] used channel names and did not require an explicit role annotation. However, the typing rule for name restriction does require that each name occur at most twice and with dual types.

With endpoint names, such role annotations are unnecessary, as each endpoint name is associated with a unique role. The type system checks that all endpoints are used with a coherent set of roles when checking the corresponding v -binder.

Do Tensor and Par Correspond to Send and Receive?

Reader familiar with CP may be confused at the assertion that tensor captures some notion of *disjointness*, rather than sending. Wadler [2012] interprets tensor and par as sending and receiving, respectively.

$$\frac{P \vdash_c \Gamma, y : A \quad Q \vdash_c \Delta, x : B}{x[y].(P \parallel Q) \vdash_c \Gamma, \Delta, x : A \otimes B} \text{T-SEND} \quad \frac{P \vdash_c \Gamma, y : A, x : B}{x(y).P \vdash_c \Gamma, x : A \wp B} \text{T-RECV}$$

Carbone et al. [2016, p. 4-5, *May one invert output and input?*] argue that the inverse interpretation, where tensor is interpreted as receive and par is interpreted as send, is nonsense. The argument is that a process should change its behaviour based on the information it receives, not the information it sends. That is a fair argument, but it misses a crucial aspect of Wadler’s interpretations for tensor and par. The send and receive actions, $x[y]$ and $x(y)$, implement a restricted form of delegation, but do not transmit information. They are about plumbing, not what is in the pipes.⁶ Hence, the argument does not apply.

The crucial part of the interpretation of tensor/par is not the send/receive actions, but the *disjointness* of $P \parallel Q$ and the *jointness* of P . Given a channel with endpoints $x : A \otimes B$ and $\bar{x} : A \wp B$, the process that holds $\bar{x}$ determines the order in which the sub-sessions A and B are resolved. The process that holds x must guarantee that either order works, and the manner in which CP guarantees this is by forcing the sub-sessions to be handled by entirely disjoint processes.

3.3.5 Deep Inference and Display Calculus

Hypersequents add *stratified* structural connectives. HCLL’s hyper-environments may contain environments, but its environments may not contain hyper-environments. (In a sense, HCLL’s structures are restricted to conjunctive normal form.) As such, we can view HCLL as a stratified

⁶I’m paraphrasing Robert Atkey, who frequently refers to the function of multiplicative as plumbing or throat-clearing. For a proper quote: “Multiplicatives correspond to matters of communication topology, while the additives will correspond to actual data transfer” [Atkey, 2017].

variant of the formulations of CLL as deep inference calculi [e.g. Flat BV, Guglielmi, 2007, p. 51; or LS, Straßburger, 2002] or display calculi [e.g. Belnap, 1989], where every logical connective has a corresponding structural connective without any such stratification.

3.3.6 Hypersequent Calculus

The hypersequents used in HCLL are similar to those used by Avron [1991]. Both are stratified extensions to the structure of the logic. However, their interpretations are rather different. HCLL's hypersequents are linear, multiplicative, and conjunctive, whereas Avron's hypersequents are unrestricted and “usually disjunctive” [Avron, 1996, p. 4].

3.3.7 Hypersequents, Mix, and Mix₀

The rules for hyper-environments resemble those for Mix and Mix₀.

$$\frac{\vdash \mathcal{G} \quad \vdash \mathcal{H}}{\vdash \mathcal{G} \parallel \mathcal{H}} (\parallel) \quad \frac{}{\vdash \circ} (\circ) \quad \frac{\vdash \Gamma \quad \vdash \Delta}{\vdash \Gamma, \Delta} \text{Mix} \quad \frac{}{\vdash \emptyset} \text{Mix}_0$$

However, HCLL is a conservative extension of CLL, but Mix and Mix₀ alter the logic: Mix is logically equivalent to $A \otimes B \multimap A \wp B$, and Mix₀ is logically equivalent to $\mathbf{1} \multimap \perp$. The correspondence between branching and tensor is illuminating: Mix converts *branching* (corresponding to $\otimes$) into a *comma* (corresponding to $\wp$), and Mix₀ converts *having no premises* (corresponding to $\mathbf{1}$) into the empty environment (corresponding to $\perp$). On the contrary, ($\parallel$) and ($\circ$) convert structure of the proof tree to structural connectives with the same logical interpretation.

3.4 Conclusion

In this chapter, I introduced Hypersequent CP with its typing rules, reduction semantics, label-transition semantics, and their metatheory. Hypersequent CP is a variant of CP that uses hypersequents to tighten the correspondence with the π -calculus. I proved *preservation* (Proposition 3.30) and *progress* (Proposition 3.35). I proved that CP's processes are deadlock-free (Corollary 3.42), and that its canonical forms are adequate (Corollary 3.50). I proved that HCP's connection graphs are forests (Proposition 3.52), and that any process can be disentangled into a collection of CP processes (Proposition 3.61). I proved harmony between HCP's reduction and label-transition semantics (Proposition 3.93). I defined a pair of inverse translations from CP to HCP (Definition 3.73) and from HCP to CP (Definition 3.83), and proved that they preserve types (Proposition 3.75 and Proposition 3.84) and give rise to a sound and

complete operational correspondence (Proposition 3.77; Proposition 3.87). Finally, I discussed the relation between HCP and similar logical systems, and introduced three variants of HCP with an exceptional semantics for the absurd offer, with synchronous semantics for links, and with guarded summation.

In the future, it would be interesting to extend HCP with the extensions for CP described in previous work. Some of these extensions are already described in the literature. Montesi and Peressotti [2021] describe a variant of HCP extended with the second-order quantifiers and exponentials, which are interpreted as polymorphism and replication. Qian et al. [2021] describe a variant of HCP extended with the exponentials and DiLL’s co-exponentials, which are interpreted as replication and client-server interaction. Other extensions, such as Lindley and Morris’ fixed points, have not yet been adapted to HCP.

Furthermore, it would be interesting to investigate a variant of Hypersequent CP with logical and structural connectives that capture sequential composition, such as “before” from Retoré’s Pomset logic or “seq” from Guglielmi’s BV. Any system with sequential composition necessarily deviates from the π -calculus, which only has trivial sequential composition—prefixing a process with an action. However, general sequential composition is important for programming languages. The paper that introduced session types, Honda [1993], describes a type system for a process calculus with general sequential composition. Honda’s type system is unsound, in the sense that well-typed processes may deadlock, and additional syntactic constraints are required to guarantee deadlock freedom. In hindsight, this is unsurprising—Tiu [2006] proved that Guglielmi’s “seq” cannot be captured by a standard sequent calculus, and this result is assumed to extend to Retoré’s “before” and other systems that capture sequential composition. However, the inference rules of Guglielmi’s BV and the proof nets of Retoré’s Pomset logic do not easily lend themselves to an interpretation as processes. The decorated sequent calculus for Pomset logic, introduced by Slavnov [2019], is an interesting candidate for further study.

3.5 Omitted Proofs

Lemma 3.40. *If P is ready, then $\text{Dep}(P)$ is essentially acyclic and connected.*

Proof. By case analysis on P .

- Case P is of the form $x \leftrightarrow y$.

$$\begin{aligned} V_{\text{Dep}(P)} &= \{x, y\} \\ E_{\text{Dep}(P)} &= \{xy\} \\ A_{\text{Dep}(P)} &= \emptyset \end{aligned}$$

Hence, $\text{Dep}(\mathbf{P})$ is essentially acyclic, as there are no arcs, and connected, as the two vertices are connected by an edge.

- Case $\mathbf{P}$ is of the form $x[y]. \mathbf{P}', x(y). \mathbf{P}', x[] . \mathbf{P}', x(). \mathbf{P}', x \triangleleft \text{inl}. \mathbf{P}', x \triangleleft \text{inr}. \mathbf{P}'$, or $x \triangleright \{\text{inl}: \mathbf{P}_1; \text{inr}: \mathbf{P}_2\}$.

$$V_{\text{Dep}(\mathbf{P})} = \text{fn}(\mathbf{P})$$

$$E_{\text{Dep}(\mathbf{P})} = \emptyset$$

$$A_{\text{Dep}(\mathbf{P})} = \{\vec{x}\vec{y} \mid y \in \text{fn}(\mathbf{P}) \wedge x \neq y\}$$

Hence, $\text{Dep}(\mathbf{P})$ is essentially acyclic, as all arcs are out of x , and connected, as every other vertex is connected to x .

□

Proposition 3.41. *If $\mathbf{P} \vdash \mathcal{G}$, then the dependency graph $\text{Dep}(\mathbf{P})$ is essentially acyclic, and there is an isomorphism f between the typing environments in $\mathcal{G}$ and the components of $\text{Dep}(\mathbf{P})$ that preserves fn , i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. The process $\mathbf{P}$ is of the form $\mathcal{C}^n[\mathbf{T}_1, \dots, \mathbf{T}_n]$ (for some $n \geq 0$). By induction on $\mathcal{C}[\cdot]$ and inversion $\mathbf{P}$ and the derivation of $\mathbf{P} \vdash \mathcal{G}$.

There are four cases:

- Case $\mathcal{C}[\cdot]$ is of the form $\square$.

Let G be $\text{Dep}(\mathbf{P})$. As $\mathcal{C}[\cdot]$ is maximal, $\mathbf{P}$ is ready. By Lemma 3.40, G is essentially acyclic and connected. As $\mathbf{P}$ is ready, $\mathcal{G}$ is of the form Γ . Let f be the function $\{\Gamma \mapsto G\}$. By Proposition 3.18, $\text{fn}(\Gamma) = \text{fn}(\mathbf{P})$. By definition, $V_G = \text{fn}(\mathbf{P})$. Hence, $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.

- Case $\mathcal{C}[\cdot]$ is of the form 0 .

By definition, $\text{Dep}(\mathbf{P})$ is the null graph. By inversion on the derivation $\mathbf{P} \vdash \mathcal{G}$, $\mathcal{G}$ is of the form $\circ$, i.e. there are no typing environments in $\mathcal{G}$. Let f be the empty function $\emptyset$, which vacuously preserves fn .

- Case $\mathcal{C}[\cdot]$ is of the form $(\nu x \bar{x}) \mathcal{C}'[\cdot]$.

By inversion, $\mathbf{P}$ is of the form $(\nu x \bar{x}) \mathbf{P}'$ such that $(\nu x \bar{x}) \mathbf{P}' \vdash \mathcal{G}' \parallel \Gamma, \Delta$ and $\mathbf{P}' \vdash \mathcal{G}' \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}$.

Let G be $\text{Dep}(\mathbf{P})$ and G' be $\text{Dep}(\mathbf{P}')$.

By definition, $V_G = V_{G'}$, $E_G = E_{G'} \cup \{x\bar{x}\}$, and $A_G = A_{G'}$.

By induction, G' is essentially acyclic, and there is an isomorphism f' between the components of G' and the typing environments in $\mathcal{G}' \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}$ that preserves fn .

Let C_1 be $f'(\Gamma, x : A)$ and C_2 be $f'(\Delta, \bar{x} : \bar{A})$.

By definition, C_1 and C_2 are disjoint and essentially acyclic.

Let C be the graph formed by connecting C_1 and C_2 with the edge $x\bar{x}$.

By definition, C is connected and a component of G .

By Lemma A.2, C and G are essentially acyclic.

Let f be the function $\{\Gamma, \Delta \mapsto C\} \cup f' \triangleleft \{\Gamma, x : A, \Delta, \bar{x} : \bar{A}\}$ (where $\triangleleft$ is domain subtraction, see Definition A.4).

The function f is an isomorphism that preserves fn, by definition (for Γ, Δ) and by induction (for the typing environments in $\mathcal{G}$).

- Case $\mathcal{C}[\cdot]$ is of the form $\mathcal{C}_1[\cdot] \parallel \mathcal{C}_2[\cdot]$.

By inversion, P is of the form $P_1 \parallel P_2$ such that $P_1 \parallel P_2 \vdash \mathcal{G}_1 \parallel \mathcal{G}_2$, $P_1 \vdash \mathcal{G}_1$, and $P_2 \vdash \mathcal{G}_2$.

Let G be $\text{Dep}(P)$, G_1 be $\text{Dep}(P_1)$, and G_2 be $\text{Dep}(P_2)$.

By definition, $V_G = V_{G_1} \cup V_{G_2}$, $E_G = E_{G_1} \cup E_{G_2}$, and $A_G = A_{G_1} \cup A_{G_2}$.

By induction, G_1 and G_2 are essentially acyclic, and there are isomorphisms f_1 and f_2 between the typing environments in $\mathcal{G}_1$ and $\mathcal{G}_2$, respectively, the components of G_1 and G_2 , respectively, that preserve fn.

By T-PAR, G_1 and G_2 are disjoint. Hence, G is essentially acyclic.

Let f be the function $f_1 \cup f_2$.

As union preserves the relevant properties of f_1 and f_2 , f is an isomorphism between the components of G and the typing environments in $\mathcal{G}$ that preserves fn.

□

Proposition 3.52. *If $P \vdash \mathcal{G}$, then $\text{Con}(P)$ is a forest, and there is an isomorphism f between the typing environments in $\mathcal{G}$ and the trees of $\text{Con}(P)$ that preserves fn, i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. The process $P = \mathcal{C}^n[T_1, \dots, T_n]$ (for some $n \geq 0$).

By induction on $\mathcal{C}[\cdot]$ and inversion P and the derivation of $P \vdash \mathcal{G}$.

There are four cases:

- Case $\mathcal{C}[\cdot]$ is of the form $\square$.

Let G be $\text{Con}(P)$. As $\mathcal{C}[\cdot]$ is maximal, P is ready. By definition, G is the singleton graph, which is a tree. As P is ready, $\mathcal{G}$ is of the form Γ . Let f be the function $\{\Gamma \mapsto G\}$. By Proposition 3.18, $\text{fn}(\Gamma) = \text{fn}(P)$. By definition, $V_G = \{P\}$. Hence, $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.

- Case $\mathcal{C}[\cdot]$ is of the form 0 .

By definition, G is the null graph, which is a forest. By inversion on the derivation $P \vdash \mathcal{G}$, $\mathcal{G}$ is of the form $\circ$, i.e. there are no typing environments in $\mathcal{G}$. Let f be the empty function $\emptyset$, which vacuously preserves fn.

- Case $\mathcal{C}[\cdot]$ is of the form $(\nu x \bar{x})\mathcal{C}'[\cdot]$.

By inversion, P is of the form $(\nu x \bar{x})P'$ such that $(\nu x \bar{x})P' \vdash \mathcal{G}' \parallel \Gamma, \Delta$ and $P' \vdash \mathcal{G}' \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}$.

Let G be $\text{Con}(P)$ and G' be $\text{Con}(P')$.

By Proposition 3.18 and Lemma 3.22, there exist unique $T_i, T_j \in V_{G'}$ such that $x \in \text{fn}(T_i)$ and $\bar{x} \in \text{fn}(T_j)$

By definition, $V_G = V_{G'}$ and $\ell_G = \ell_{G'} \cup \{T_i T_j \mapsto (x, \bar{x})\}$.

By induction, G' is a forest, and there is an isomorphism f' between the components of G' and the typing environments in $\mathcal{G}' \parallel \Gamma, x : A \parallel \Delta, \bar{x} : \bar{A}$ that preserves fn.

Let T_1 be $f'(\Gamma, x : A)$ and T_2 be $f'(\Delta, \bar{x} : \bar{A})$.

By definition, T_1 and T_2 are disjoint trees.

Let T be the graph formed by connecting T_1 and T_2 with the edge $P_i P_j$. By Lemma A.1, T is a tree. Hence, G is a forest.

Let f be the function $\{\Gamma, \Delta \mapsto C\} \cup f' \triangleleft \{\Gamma, x : A, \Delta, \bar{x} : \bar{A}\}$ (where $\triangleleft$ is domain subtraction, see Definition A.4).

The function f is an isomorphism that preserves fn, by definition (for Γ, Δ) and by induction (for the typing environments in $\mathcal{G}'$).

- Case $\mathcal{C}[\cdot]$ is of the form $\mathcal{C}_1[\cdot] \parallel \mathcal{C}_2[\cdot]$.

By inversion, P is of the form $P_1 \parallel P_2$ such that $P_1 \parallel P_2 \vdash \mathcal{G}_1 \parallel \mathcal{G}_2$, $P_1 \vdash \mathcal{G}_1$, and $P_2 \vdash \mathcal{G}_2$.

Let G be $\text{Con}(P)$, G_1 be $\text{Con}(P_1)$, and G_2 be $\text{Con}(P_2)$.

By definition, $V_G = V_{G_1} \cup V_{G_2}$ and $\ell_G = \ell_{G_1} \cup \ell_{G_2}$.

By induction, G_1 and G_2 are forests, and there are isomorphisms f_1 and f_2 between the typing environments in $\mathcal{G}_1$ and $\mathcal{G}_2$, respectively, the trees of G_1 and G_2 , respectively, that preserve fn.

By T-PAR, G_1 and G_2 are disjoint. Hence, G is a forest.

Let f be the function $f_1 \cup f_2$.

As union preserves the relevant properties of f_1 and f_2 , f is an isomorphism between the components of G and the typing environments in $\mathcal{G}$ that preserves fn.

□

Lemma 3.56. *If $P \vdash \mathcal{G} \parallel \Gamma$, then there exist processes Q and R such that $Q \vdash \mathcal{G}$, $R \vdash \Gamma$, $P \stackrel{\text{is}}{=} Q \parallel R$, and R is in right-branching tree form.*

Proof. By Proposition 3.52, $\text{Con}(P)$ is a forest and some component T is a tree such that $\text{fn}(\Gamma) = \text{fn}(V_T)$. By induction, guarded by the size of T .

- Case $|V_T| = 1$.

There is some thread R such that $V_T = \{R\}$, and some evaluation context $\mathcal{E}$ such that $P = \mathcal{E}[R]$. As T is a component of $\text{Con}(P)$, $\text{fn}(R) \cap \text{bn}(\mathcal{E}) = \emptyset$. By Corollary 3.26, $P \stackrel{\text{is}}{=} \mathcal{E}[0] \parallel R$. Let Q be $\mathcal{E}[0]$. Recall that R is in right-branching form, as it is ready.

- Case $|V_T| > 1$.

Let T_i be any leaf of T . As T_i is a leaf, exactly one free endpoint in T_i is bound in P . Let us name that endpoint x and its dual $\bar{x}$. There exist some $\mathcal{E}_1$, $\mathcal{E}_2$, $\mathcal{F}_i$, $\mathcal{F}_j$, and P_j such that

$$\begin{aligned} P &\stackrel{\text{is}}{=} \mathcal{E}_1[(vx\bar{x})(\mathcal{E}_2[\mathcal{F}_i[T_i] \parallel \mathcal{F}_j[P_j]])] && \langle \text{by Corollary 3.23} \rangle \\ &\stackrel{\text{is}}{=} (vx\bar{x})(\mathcal{E}_1[\mathcal{E}_2[\mathcal{F}_i[T_i] \parallel \mathcal{F}_j[P_j]])] && \langle \text{by Lemma 3.24} \rangle \\ &\stackrel{\text{is}}{=} (vx\bar{x})(\mathcal{E}_1[\mathcal{E}_2[\mathcal{F}_i[T_i \parallel \mathcal{F}_j[P_j]]]]) && \langle \text{by Lemma 3.25} \rangle \\ &\stackrel{\text{is}}{=} (vx\bar{x})(T_i \parallel \mathcal{E}_1[\mathcal{E}_2[\mathcal{F}_i[\mathcal{F}_j[P_j]]]]) && \langle \text{by Lemma 3.25} \rangle \end{aligned}$$

Let P' be $\mathcal{E}_1[\mathcal{E}_2[\mathcal{F}_i[\mathcal{F}_j[P_j]]]]$. By Lemma 3.27, P' is well-typed.

Let T' be the tree formed by deleting the vertex T_i and its only edge from T . Clearly, $|V_T| = |V_{T'}| + 1$. Every thread in P is preserved in P' , except T_i . Therefore, the connection graph $\text{Con}(P')$ is the forest formed by replacing T with T' .

By induction with P' , there exist processes Q and R' such that $P' \stackrel{\text{is}}{=} Q \parallel R'$ and R' is in right-branching tree form. Therefore,

$$\begin{aligned} P &\stackrel{\text{is}}{=} (vx\bar{x})(T_i \parallel \mathcal{E}_1[\mathcal{E}_2[\mathcal{F}_i[\mathcal{F}_j[P_j]]]]) && \langle \text{see above} \rangle \\ &\stackrel{\text{is}}{=} (vx\bar{x})(T_i \parallel (Q \parallel R')) && \langle \text{by induction} \rangle \\ &\stackrel{\text{is}}{=} Q \parallel (vx\bar{x})(T_i \parallel R') && \langle \text{by Lemma 3.25} \rangle \end{aligned}$$

Let R be $(vx\bar{x})(T_i \parallel R')$. As R' is in right-branching tree form, R is in right-branching tree form.

It remains to show that the processes Q and R have the correct types.

For some Γ_i , Γ_j , and A , $T_i \vdash \Gamma_i, x : A$ and $P_j \vdash \Gamma_j, \bar{x} : \bar{A}$.

Since T_i is a leaf, the evaluation contexts $\mathcal{E}_1$, $\mathcal{E}_2$, and $\mathcal{F}_i$ must preserve the contents of its typing environment Γ_i . In essence, we can assign $\mathcal{E}_1$, $\mathcal{E}_2$, and $\mathcal{F}_i$ type schemas, rather than types. For some $\mathcal{G}'$, $\mathcal{G}'_i$, $\mathcal{G}'_j$, Γ'' ,

Γ'' , and Γ' the evaluation contexts $\mathcal{E}_1$, $\mathcal{E}_2$, and $\mathcal{F}_i$ have the following type schemas, and $\mathcal{F}_j$ has the following type.

$$\begin{array}{ll}
\forall \Delta. \mathcal{E}_1 \vdash \mathcal{G}' \parallel \Gamma'', \Delta & \rightarrow \mathcal{G} \parallel \Gamma', \Delta \quad \text{or} \\
\forall \Delta. \mathcal{E}_1 \vdash \mathcal{G}' \parallel \Gamma'', \bar{x} : \bar{A} \parallel \Delta & \rightarrow \mathcal{G} \parallel \Gamma', \bar{x} : \bar{A} \parallel \Delta \\
\forall \Delta. \mathcal{E}_2 \vdash \mathcal{G}'_i \parallel \mathcal{G}'_j \parallel \Gamma''', \bar{x} : \bar{A} \parallel \Delta & \rightarrow \mathcal{G}' \parallel \Gamma'', \bar{x} : \bar{A} \parallel \Delta \\
\forall \Delta. \mathcal{F}_i \vdash \Delta & \rightarrow \mathcal{G}'_i \parallel \Delta \\
\mathcal{F}_j \vdash \Gamma_j, \bar{x} : \bar{A} & \rightarrow \mathcal{G}'_j \parallel \Gamma''', \bar{x} : \bar{A}
\end{array}$$

Hence, $\Gamma = \Gamma', \Gamma_i$ and $P' \vdash \mathcal{G} \parallel \Gamma', \bar{x} : \bar{A} \parallel \Gamma_i, x : A$.

By induction, $Q \vdash \mathcal{G}$ and $R' \vdash \Gamma', \bar{x} : \bar{A}$. Hence, $R \vdash \Gamma$.

□

Proposition 3.57. *If $P \vdash \mathcal{G}$, then there exists a process Q , such that $P \cong Q$, and Q is in right-branching forest form, and there is an isomorphism f between the processes in right-branching tree form in Q and the typing environments in $\mathcal{G}$ that preserves fn, i.e. $\text{fn}(\Gamma) = \text{fn}(f(\Gamma))$.*

Proof. By induction on the hyper-environment $\mathcal{G}$.

- Case $\mathcal{G}$ is of the form $\circ$.

By Lemma 3.55, $P \cong 0$. Let Q be 0 . The result follows.

- Case $\mathcal{G}$ is of the form $\mathcal{G} \parallel \Gamma$ (reusing $\mathcal{G}$).

By Lemma 3.56, there exist processes P' and R such that $P' \vdash \mathcal{G}$, $R \vdash \Gamma$, $P \cong P' \parallel R$, and R is in right-branching tree form.

By induction, there exists a process Q' such that $P' \cong Q'$ and Q' is in right-branching forest form. There are two cases:

- If $Q' = 0$, let Q be R . By inversion, $\mathcal{G} = \circ$. The result follows, as, by SC-PARNIL, $P \cong R$, and R is in right-branching forest form.
- Otherwise, let Q be $Q' \parallel R$. The result follows, as $P \cong Q$ and Q is in right-branching forest form.

□

Lemma 3.89. *If $P \xRightarrow{\ell} Q$, then $P \xrightarrow{\ell} Q$.*

Proof. Unfolding the compositions of the structural congruence and the transition relations, the goal is as follows:

$$P \equiv P' \wedge P' \xrightarrow{\ell} Q \implies P \xrightarrow{\ell} Q' \wedge Q' \equiv Q$$

I refer to the subgoals $P \xrightarrow{\ell} Q'$ and $Q' \equiv Q$ as results (1) and (2), respectively.

By induction on the derivation of the structural congruence $P \equiv Q$ and inversion on the transition $P \xrightarrow{\ell} P'$. The case for reflexivity follows immediately. The cases for symmetric use of the rules of structural congruence follow by an analogy to the cases below. The case for transitivity follows by induction. The case for congruence follows by induction on the transition.

- Case SC-LINKCOMM.

The structural congruence is of the form:

$$x \leftrightarrow y \equiv y \leftrightarrow x$$

By inversion on the transition:

- Subcase ACT-LINK₁. Result (1) follows by ACT-LINK₂.
- Subcase ACT-LINK₂. Result (1) follows by ACT-LINK₁.

Result (2) follows by reflexivity.

- Case SC-PARNIL.

The structural congruence is of the form (reusing P):

$$P \parallel 0 \equiv P$$

Result (1) follows by STR-CONG with $\square \parallel 0$.

Result (2) follows by SC-PARNIL.

- Case SC-PARCOMM.

The structural congruence is of the form (reusing P and Q):

$$P \parallel Q \equiv Q \parallel P$$

By inversion on the transition:

- Subcase STR-PAR.
Result (1) follows as $\pi \parallel \bar{\pi}$ is unordered.
- Subcase STR-CONG with $\mathcal{E} \parallel Q$.
Result (1) follows by STR-CONG with $Q \parallel \mathcal{E}$.
- Subcase STR-CONG with $P \parallel \mathcal{E}$.
Result (1) follows by STR-CONG with $\mathcal{E} \parallel P$.

Result (2) follows by SC-PARCOMM.

- Case SC-PARASSOC.

The structural congruence is of the form (reusing P and Q):

$$P \parallel (Q \parallel R) \equiv (P \parallel Q) \parallel R$$

By inversion on the transition:

- Subcase STR-PAR.

Result (1) follows as $\pi \parallel \bar{\pi}$ is unordered.

- Subcase STR-CONG with $P \parallel \square$ then STR-PAR.

Result (1) follows as $\pi \parallel \bar{\pi}$ is unordered.

- Subcase STR-CONG with $\mathcal{E} \parallel (Q \parallel R)$.

Result (1) follows by STR-CONG with $(\mathcal{E} \parallel Q) \parallel R$.

- Subcase STR-CONG with $P \parallel (\mathcal{E} \parallel R)$.

Result (1) follows by STR-CONG with $(P \parallel \mathcal{E}) \parallel R$.

- Subcase STR-CONG with $P \parallel (Q \parallel \mathcal{E})$.

Result (1) follows by STR-CONG with $(P \parallel Q) \parallel \mathcal{E}$.

Result (2) follows by SC-PARASSOC.

- Case SC-NEWCOMM.

The structural congruence is of the form (reusing P):

$$(\nu x \bar{x})P \equiv (\nu \bar{x}x)P$$

By inversion on the transition. In the cases for TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, TAU-SELECT-OFFER₂, result (1) follows as $\pi \parallel \bar{\pi}$ is unordered.

In the cases for TAU-LINK and TAU-CLOSE-WAIT, result (2) follows by reflexivity. In the cases for TAU-SEND-RECV, TAU-SELECT-OFFER₁, TAU-SELECT-OFFER₂, result (2) follows by SC-NEWCOMM.

- Case SC-NEWASSOC.

The structural congruence is of the form (reusing P):

$$(\nu x \bar{x})(\nu y \bar{y})P \equiv (\nu y \bar{y})(\nu x \bar{x})P$$

By inversion on the transition:

- Subcases TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂ then STR-CONG with $(\nu y \bar{y})\square$.

Result (1) follows by STR-CONG with $(\nu y \bar{y})\square$ then TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂.

In the cases for TAU-LINK and TAU-CLOSE-WAIT, result (2) follows by reflexivity. In the cases for TAU-SEND-RECV, TAU-SELECT-OFFER₁, TAU-SELECT-OFFER₂, result (2) follows by SC-NEWASSOC.

- Subcases STR-CONG with $(\nu x\bar{x})\square$ then TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂.

Result (1) follows by TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂ then STR-CONG with $(\nu x\bar{x})\square$.

In the cases for TAU-LINK and TAU-CLOSE-WAIT, result (2) follows by reflexivity. In the cases for TAU-SEND-RECV, TAU-SELECT-OFFER₁, TAU-SELECT-OFFER₂, result (2) follows by SC-NEWASSOC.

- Subcase STR-CONG with $(\nu x\bar{x})(\nu y\bar{y})\square$.

Result (1) follows by STR-CONG with $(\nu y\bar{y})(\nu x\bar{x})\square$.

Result (2) follows by SC-NEWASSOC.

- Case SC-SCOPEEXT.

The structural congruence is of the form (reusing P and Q):

$$(\nu x\bar{x})(P \parallel Q) \equiv P \parallel (\nu x\bar{x})Q$$

By inversion on the transition:

- Subcases TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂ then STR-PAR or STR-CONG with $\mathcal{E} \parallel Q$.

Impossible. The side condition for SC-SCOPEEXT requires $x, \bar{x} \notin \text{fn}(P)$, but either $x \in \text{fn}(P)$ or $\bar{x} \in \text{fn}(P)$ must hold.

- Subcases TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂ then STR-CONG with $P \parallel \mathcal{E}$. Result (1) follows by STR-CONG with $P \parallel \square$, then TAU-LINK, TAU-SEND-RECV, TAU-CLOSE-WAIT, TAU-SELECT-OFFER₁, or TAU-SELECT-OFFER₂, then STR-CONG with $\mathcal{E}$.

In the cases for TAU-LINK and TAU-CLOSE-WAIT, result (2) follows by reflexivity. In the cases for TAU-SEND-RECV, TAU-SELECT-OFFER₁, TAU-SELECT-OFFER₂, result (2) follows by SC-NEWASSOC.

- Subcases STR-CONG with $(\nu x\bar{x})\square$ then STR-PAR.

Result (1) follows by STR-PAR then STR-CONG with $(\nu x\bar{x})\square$.

Result (2) follows by SC-SCOPEEXT.

- Subcases STR-CONG with $(\nu x\bar{x})(\mathcal{E} \parallel Q)$

Result (1) follows by STR-CONG with $\mathcal{E} \parallel (\nu x\bar{x})Q$

Result (2) follows by SC-SCOPEEXT.

- Subcases STR-CONG with $(\nu x\bar{x})(P \parallel \mathcal{E})$.

Result (1) follows by STR-CONG with $P \parallel (vx\bar{x})\mathcal{E}$.

Result (2) follows by SC-SCOPEEXT.

□

Chapter 4

Hypersequent Good Variation

This chapter presents Hypersequent GV (HGV), concurrent λ -calculus with session-typed concurrency primitives, which has a tight correspondence to Hypersequent CP.

Hypersequent GV was first introduced by Fowler et al. [2021], as a variant of GV that uses hypersequents for its runtime type system. While CP has not undergone many significant changes since its publication by Wadler [2012], GV's history is much more storied:

- Good Variation was first published under that name by Wadler [2012], who adapted the LAST calculus of Gay and Vasconcelos [2010] to fit a correspondence with CP. However, Wadler's GV has neither an operational semantics—except implicitly, by its translation to CP—nor an account of polymorphism or replication.

This work has an error in its translation from GV to CP, which does not preserve typing in the case for session termination (Theorem 3).

- Lindley and Morris [2014]¹ extend GV with operations for forwarding, polymorphism, and servers and clients, rectify the error in Wadler's translation from GV to CP that breaks typing, and describe the first translation from CP to GV.

This work does not give an operational semantics for GV, and its translation from GV to CP does not preserve reduction.

- Lindley and Morris [2015] give an operational semantics for GV, drop polymorphism and servers and clients, refactor GV so that the concurrency primitives are constant functions rather than term constructors, and switch to an encoding of the choice operators following Dardha et al. [2017].

¹Lindley and Morris [2014] name their system “Harmonious GV” as the operations in the their version of GV restore the balance between and CP. Later publications in the same line of work refer to their system as GV.

This work has an error in its operational semantics, which does not satisfy the diamond property in the case for link (Theorem 13), and in its translation from GV to CP, which does not preserve the reductions for products and sums² (Lemma 21 and Theorems 22 and 23).

- Lindley and Morris [2016b] present μCP and μGV , which adds fixed points to both GV and CP. Orthogonally, this work also rectifies the error in the semantics for link, and adds the direct unit “ $\top$ ” and the direct product “ $\&$ ” [Girard and Lafont, 1987].

This work has the same error in its translation from GV to CP as Lindley and Morris [2015], which does not preserve the reductions for products and sums (Theorem 16).

- Lindley and Morris [2017] present FST, which adds record & variant types with row polymorphism, an account of subkinding for linear and unrestricted types, asynchrony, and access points, and drops link, weak explicit substitution, and the correspondence with CP. This work does not present any of the metatheory of FST.
- Fowler et al. [2019] and Fowler [2019] present EGV, which adds exceptions and handlers to GV. This work also drops link, weak explicit substitution, and the correspondence with CP.
- Fowler et al. [2021] present Hypersequent GV, which rectifies a design flaw that complicates all previous metatheory of GV by dropping lock typing in favour of hypersequents. Orthogonally, this work also rectifies the error in the translation from GV to CP, and presents the first such translation from that both preserves and reflects reduction. (The extended version of this paper is presented in Chapter 4.)

In one sense, Hypersequent GV relates to GV as Hypersequent CP relates to CP. It uses hypersequents to separate name restriction from parallel composition in its configuration typing. In doing so, it removes lock typing and makes the structural congruence type preserving, which significantly simplifies the metatheory of GV.

In another sense, Hypersequent GV is simply GV. Recall that GV separates its language into a static fragment and a runtime fragment, and the static fragment—the user-facing programming language—is unchanged from

²Lindley and Morris [2015] use weak explicit substitution in an effort to obtain a translation from GV to CP that preserves reduction, which works in principle, as the translation preserves β -reduction. Unfortunately, their implementation of weak explicit substitution only suspends substitution on λ -abstraction. Hence, the translation does not preserve reductions that involve the other binders, i.e. those for products and sums. Consequently, Lemma 21, Theorem 22, and Theorem 23 fail to hold. Presumably, this is easily solved by either suspending substitution on all binders, or by encoding the eliminators for products and sums as constants and making λ -abstraction the only binder.

GV. Hypersequent CP generalises CP’s process structure from trees to forests, and likewise Hypersequent GV generalises GV’s configuration structures from trees to forests. However, reduction from a single term—a single user program—only ever reaches tree-structured configurations. That does not mean that we could not use the framework of GV to study configurations arising from multiple user programs—we could, although, since they are disconnected, this is unlikely to provide us any novel insights.

The bulk of this chapter consists of the paper *Separating Sessions Smoothly* by Fowler et al. [2023], hereafter referred to as Paper I. References made from the main body of this thesis into Paper I will be prefixed by an “I”, e.g. “Theorem I.3.20”. This chapter proceeds as follows:

- In § 4.1, we provide a legend and an errata for Paper I.
- In § 4.2, we present Hypersequent GV and its metatheory, together with translations to and from GV, and an operational correspondence with Hypersequent CP. This section consists entirely of Paper I, and proceeds as follows:
 - In § I.2, we discuss the complications in GV’s metatheory that arise from the use of lock typing.
 - In § I.3, we present HGV.
 - In § I.3.1, we present the metatheory for HGV.

Notably, we prove *preservation* (Theorem I.3.3), the *tree-structure* of connections in configurations (Theorem I.3.14), *global progress* (Theorem I.3.20), the *diamond property* (Theorem I.3.21), and *termination* (Theorem I.3.22).

- In § I.4, we present the relation between HGV and GV.

Notably, we prove that any GV configuration is typeable in HGV (Theorem I.4.3), and that any HGV configuration can be rewritten by structural congruence to obtain a configuration that is typeable in GV (Corollary I.4.7).

- In § I.5, we present the relation between HGV and HCP.

Notably, we define fine-grain call-by-value HGV (HGV*, § I.5, under “HGV*”), define a naive translation from HGV to HGV* (§ I.5, Definition I.5.8) and define a translation from HGV* to HCP (§ I.5, Figure I.8), and prove that the latter translation preserves types (§ I.5, Lemma I.5.9) and is a sound and complete operational correspondence (§ I.5, Theorem I.5.11). Finally, we define a translation from HCP to HGV, by composing existing translations (§ I.5, under “Translating HCP to HGV”).

- In § I.6, we present extensions of HGV.

- In § I.7, we discuss the possibility of using hyper-environments in term typing.
- In § I.8, we discuss the related work.

4.1 Legend and Errata

The conventions and terminology in Paper I are different from those used in the rest of this thesis.

- The terms and types of *both* Hypersequent CP and Hypersequent GV are printed in **red** and **blue**, respectively, and are rendered in a font with serif. The keywords in HGV’s syntax are bolded.
- The names for the rules of structural congruence, reduction, and the type system differ slightly from those used for HCP in Chapter 3.
- Instead of including a single rule for congruence under evaluation contexts, the paper includes separate congruence rules for name restriction and parallel composition (E-RES and E-PAR).
- Instead of the *connection graph*, the paper uses the “abstract process structure” (§ I.3.1, Definition I.3.9), which is derived from a hyper-environment and co-name set, as opposed to the configuration.
- Instead of *ready thread*, the paper uses the term “blocked thread” (§ I.3.1, Definition I.3.16.)
- The duality for HCP is written as $A^\perp$ rather than $\bar{A}$ (see Figure I.6).
- The lts for HCP (see Figure I.7) combines a number of action rules into ACT-PREF, renames ACT-OFFER₁ and ACT-OFFER₂ to ACT-OFF-INL and ACT-OFF-INR, respectively, makes the distinction between α -transition and β -transition explicit in the rules for τ -transition, rather than a post-facto restriction, replaces the congruence rule STR-CONG with individual congruence rules for name restriction and parallel composition (STR-RES, STR-PAR₁, and STR-PAR₂), and renames STR-PAR to STR-SYN.

To the best of my knowledge, there are no significant errors in Paper I, but there are a small number of typesetting errors:

- The definition of internal and external choice (§ I.3, under “Choice”) define $S \oplus S'$ and $S \& S'$ in terms of S_1 and S_2 , rather than S and S' . The correct definitions are:

$$S \oplus S' \triangleq !1.!(\bar{S} + \bar{S}').end_l \quad S \& S' \triangleq ?1.?(S + S').end_r$$

- The definition of an *abstract process structure* (§ I.3.1, Definition I.3.9) uses the function `envs` without prior definition. The function maps hyper-environments to sets of environments, and may be defined as $envs(\Gamma_1 \parallel \dots \parallel \Gamma_n) \triangleq \{\Gamma_1, \dots, \Gamma_n\}$.

- The definition of a *ground configuration* (§ I.3.1, Definition I.3.19) states that a configuration C is a ground configuration if $\cdot \vdash C : T$, amongst other conditions. This is syntactically ill-formed, as T is not a configuration type, and should be $\cdot \vdash C : \bullet T$.
- One of TG-CONNECT_1 and TG-CONNECT_2 can be elided from GV's configuration typing rules (Figure I.5) without compromising type preservation of reduction, since reduction is only defined on bound endpoints and one can always dualise the type of the bound endpoints. Having both comes at a cost, as parallel compositions no longer have unique typing derivations. Fortunately, the problem is easily addressed by eliding either TG-CONNECT_1 or TG-CONNECT_2 .

The rules TG-CONNECT_1 and TG-CONNECT_2 were mistakenly reproduced from Fowler [2019]. In Fowler's GV, both rules are needed, since reduction is allowed outside of the scope of a name restriction, and therefore, it is not possible to dualise the type of the locked channel without breaking preservation. However, the version of GV presented in this chapter only permits reduction under a name restriction.

- The proof of *global progress* (§ I.3.1, Theorem I.3.20) states that any ground configuration that cannot be reduced can be written as:

$$(\nu x_1 y_1)(\circ \mathcal{A}_1 \parallel \dots \parallel (\nu x_n y_n)(\circ \mathcal{A}_n \parallel \bullet V) \dots)$$

This is syntactically ill formed, as $\mathcal{A}$ denotes an auxiliary thread, which already includes the thread flag $\circ$. The correct statement is:

$$(\nu x_1 y_1)(\bullet \mathcal{A}_1 \parallel \dots \parallel (\nu x_n y_n)(\bullet \mathcal{A}_n \parallel \bullet V) \dots)$$

- The statement of the *diamond property* (§ I.3.1, Theorem I.3.21) is:

$$\text{If } \mathcal{G} \vdash \mathcal{C} : T, \mathcal{C} \longrightarrow \mathcal{D}, \mathcal{C} \longrightarrow \mathcal{D}', \text{ then } \mathcal{D} \equiv \mathcal{D}'.$$

This does not hold—and, arguably, is not the diamond property—since it does not permit the result to perform any reductions. Hence, if the reductions $\mathcal{C} \longrightarrow \mathcal{D}$ and $\mathcal{C} \longrightarrow \mathcal{D}'$ eliminate different redexes, there is no leeway to eliminate the redex targeted by $\mathcal{C} \longrightarrow \mathcal{D}$ in $\mathcal{D}'$ and vice versa. The correct statement is:

$$\text{If } \mathcal{G} \vdash \mathcal{C} : T \text{ and } \mathcal{C} \longrightarrow \mathcal{D}_1 \text{ and } \mathcal{C} \longrightarrow \mathcal{D}_2, \text{ then either } \mathcal{D}_1 \equiv \mathcal{D}_2 \\ \text{or there exists some } \mathcal{D}_3 \text{ such that } \mathcal{D}_1 \longrightarrow \mathcal{D}_3 \text{ and } \mathcal{D}_2 \longrightarrow \mathcal{D}_3.$$

- The definition of *flattening* (§ I.4, Definition I.4.1) translates the empty hyper-environment to the empty hyper-environment, i.e. $\downarrow \emptyset = \emptyset$.

This is incorrect, as the codomain of flattening is HGV *typing environments*, and the empty typing environment is denoted by $\cdot$. The correct definition is $\downarrow \emptyset = \cdot$.

4.2 Paper I: Separating Sessions Smoothly

This section contains the paper with the same title, written in collaboration with Simon Fowler, Ornela Dardha, Sam Lindley, and J. Garrett Morris, which was originally published in the journal Logical Methods in Computer Science, Volume 19, Issue 3, 2023, and is an extended journal version of the paper with the same title and authors originally published in the proceedings for the 32nd International Conference on Concurrency Theory (CONCUR 2021) as part of the Leibniz International Proceedings in Informatics (LIPIcs) series.

J. Garrett Morris acted as my second Ph.D. supervisor from 2016 to 2017, and Sam Lindley acted as my second Ph.D. supervisor from 2018 to 2024.

The work presented in the paper was conceived of by all the authors. I co-developed Hypersequent GV and was primarily responsible for the correspondence between Hypersequent GV and Hypersequent CP.

SEPARATING SESSIONS SMOOTHLY

SIMON FOWLER ^a, WEN KOKKE ^b, ORNELA DARDHA ^a, SAM LINDLEY ^c,
AND J. GARRETT MORRIS ^d

^a University of Glasgow, UK

e-mail address: simon.fowler@glasgow.ac.uk, ornela.dardha@glasgow.ac.uk

^b University of Strathclyde, UK

e-mail address: wen.kokke@strath.ac.uk

^c The University of Edinburgh, UK

e-mail address: sam.lindley@ed.ac.uk

^d The University of Iowa, USA

e-mail address: garrett-morris@uiowa.edu

ABSTRACT. This paper introduces Hypersequent GV (HGV), a modular and extensible core calculus for functional programming with session types that enjoys deadlock freedom, confluence, and strong normalisation. HGV exploits hyper-environments, which are collections of type environments, to ensure that structural congruence is type preserving. As a consequence we obtain an operational correspondence between HGV and HCP—a process calculus based on hypersequents and in a propositions-as-types correspondence with classical linear logic (CLL). Our translations from HGV to HCP and vice-versa both preserve and reflect reduction. HGV scales smoothly to support Girard’s Mix rule, a crucial ingredient for channel forwarding and exceptions.

1. INTRODUCTION

Session types [Hon93, THK94, HVK98] are types used to model and verify communication protocols in concurrent and distributed systems: just as data types rule out dividing an integer by a string, session types rule out sending an unexpected message. Session types originated in process calculi, but there is a gap between process calculi, which model the evolving state of concurrent systems, and the descriptions of these systems in mainstream programming languages. This paper addresses two foundations for session types: (1) a session-typed concurrent lambda calculus called GV [LM15], intended to be a modular and extensible basis for functional programming languages with session types; and, (2) a session-typed process calculus called CP [Wad14], with a propositions-as-types correspondence to classical linear logic (CLL) [Gir87].

Processes in CP correspond exactly to proofs in CLL and deadlock freedom follows from cut-elimination for CLL. However, while CP is strongly tied to CLL, at the same time it departs from the π -calculus. Independent π -calculus features can only appear in combination in CP: CP combines name restriction with parallel composition $((\nu x)(P \parallel Q))$, corresponding

to CLL’s cut rule, and combines sending (of bound names only) with parallel composition ($x[y].(P \parallel Q)$), corresponding to CLL’s tensor rule. This results in a proliferation of process constructors and prevents the use of standard techniques from concurrency theory, such as labelled-transition semantics and bisimulation, since the expected transitions give rise to ill-typed terms. For example, we cannot write the expected transition rule for output, $x[y].(P \parallel Q) \xrightarrow{x[y]} P \parallel Q$, since $P \parallel Q$ is not a valid CP process. A similar issue arises when attempting to design a synchronisation transition rule for bound output (see [KMP19b] for a detailed discussion). Inspired by Carbone *et al.* [CMS18] who use hypersequents [Avr91] to give a logical grounding to choreographic programming languages [Mon13], Hypersequent CP (HCP) [KMP19a, KMP19b, MP18] restores the independence of these features by factoring out parallel composition into a standalone construct while retaining the close correspondence with CLL proofs. HCP typing reasons about collections of processes using collections of type environments (or *hyper-environments*).

GV extends linear λ -calculus with constants for session-typed communication. Following Gay and Vasconcelos [GV10], Lindley and Morris [LM15] describe GV’s semantics by combining a reduction relation on single terms, following standard λ -calculus rules, and a reduction relation on concurrent configurations of terms, following standard π -calculus rules. They give a semantic characterisation of deadlocked processes, an extrinsic [Rey00] type system for configurations, and show that well-typed configurations are deadlock-free. There is, however, a large fly in this otherwise smooth ointment: GV’s process equivalence does not preserve typing. As a result, it is not enough for Lindley and Morris to show progress and preservation for well-typed configurations; instead, they must show progress and preservation for *all* configurations *equivalent* to well-typed configurations. This not only complicates the metatheory of GV, but the burden is inherited by any effort to build on GV’s account of concurrency [FLMD19].

In this paper, we show that using hyper-environments in the typing of configurations enables a metatheory for GV that, compared to that of Lindley and Morris, is simpler, is more general, and as a result is easier to use and easier to extend. Hypersequent GV (HGV) repairs the treatment of process equivalence—equivalent configurations are equivalently typeable—and avoids the need for formal gimmickry connecting name restriction and parallel composition. HGV admits standard semantic techniques for concurrent programs: we use bisimulation to show that our translations both preserve *and* reflect reduction, whereas Lindley and Morris resort to weak explicit substitutions [LM99] and only show that their translations between GV and CP preserve reduction. HGV is also more easily extensible: we outline three examples, including showing that HGV naturally extends to disconnected sets of communication processes, without any change to the proof of deadlock freedom, and that it serves as a simpler foundation for existing work on exceptions in GV [FLMD19].

Contributions. The paper contributes the following:

- Section 3 introduces Hypersequent GV (HGV), a modular and extensible core calculus for functional programming with session types which uses hyper-environments to ensure that structural congruence is type preserving.
- Section 4 shows that every well-typed GV configuration is also a well-typed HGV configuration, and every tree-structured HGV configuration is equivalent to a well-typed GV configuration.
- Section 5 gives an operational correspondences between HGV and HCP via translations in both directions that preserve and reflect reduction.

- Section 6 demonstrates the extensibility of HGV through: (1) unconnected processes, (2) a simplified treatment of forwarding, and (3) an improved foundation for exceptions.

Section 2 reviews GV and its metatheory, Section 7 discusses why it is difficult to apply hyper-environments to term typing, Section 8 discusses related work, and Section 9 concludes and discusses future work.

This paper is an improved and extended version of a paper published at CONCUR 2021 [FKD⁺21]. Additional highlights include:

- a more detailed account of process structures;
- a more detailed account of extensions;
- a more detailed account of the metatheory for HCP; and
- a modified formulation of HCP's labelled transition system and the translation of **fork** in Section 5 fixing errors in the operational correspondence result from the CONCUR 2021 paper.

Proofs of all of the technical results are included in the paper.

2. THE EQUIVALENCE EMBROGLIO

GV programs are deadlock free, which GV ensures by restricting process structures to trees. A *process structure* is an undirected graph where nodes represent processes and edges represent channels shared between the connected nodes. Session-typed programs with an acyclic process structure are deadlock-free by construction. We illustrate this with a session-typed vending machine example written in GV.

Example 2.1. Consider the session type of a vending machine below, which sells chocolate bars and lollipops. If the vending machine is free, the customer can press ① to receive a chocolate bar or ② to receive a lollipop. If the vending machine is busy, the session ends.

$$\text{VendingMachine} \triangleq \oplus \left\{ \begin{array}{l} \text{Free} : \& \{ \textcircled{1} : !\text{ChocolateBar}.\text{end}_! , \textcircled{2} : !\text{Lollipop}.\text{end}_! \} \\ \text{Busy} : \text{end}_! \end{array} \right\}$$

The customer's session type is *dual*: where the vending machine sends a **ChocolateBar**, the customer receives a **ChocolateBar**, and so forth. Figure 1 shows the vending machine and customer as a GV program with its process structure.

GV establishes the restriction to tree-structured processes by restricting the primitive for spawning processes. In GV, **fork** has type $(S \multimap \text{end}_!) \multimap \bar{S}$. It takes a closure of type $S \multimap \text{end}_!$ as an argument, creates a channel with endpoints of dual types S and $\bar{S}$, spawns the closure as a new process by supplying one of the endpoints as an argument, and then returns the other endpoint. In essence, **fork** is a branching operation on the process structure: it creates a new node connected to the current node by a single edge. Linearity guarantees that the tree structure is preserved, even in the presence of higher-order channels.

Lindley and Morris [LM15] introduce a semantics for GV, which evaluates programs embedded in process configurations, consisting of embedded programs, flagged as main ($\bullet$) or child ($\circ$) threads, ν -binders to create new channels, and parallel compositions:

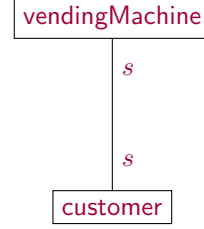
$$\mathcal{C}, \mathcal{D} ::= \bullet M \mid \circ M \mid (\nu x)\mathcal{C} \mid (\mathcal{C} \parallel \mathcal{D})$$

They introduce these process configurations together with a standard structural congruence, which allows, amongst other things, the reordering of processes using commutativity ($\mathcal{C} \parallel \mathcal{C}' \equiv \mathcal{C}' \parallel \mathcal{C}$), associativity ($\mathcal{C} \parallel (\mathcal{C}' \parallel \mathcal{C}'') \equiv (\mathcal{C} \parallel \mathcal{C}') \parallel \mathcal{C}''$), and scope extrusion

```

let vendingMachine =  $\lambda s.$ 
  let  $s = \text{select Free } s$  in
    let  $s = \text{offer } s \left\{ \begin{array}{l} \textcircled{1} \mapsto \text{send chocolateBar} \\ \textcircled{2} \mapsto \text{send lollipop} \end{array} \right\}$ 
    close  $s$ 
in let customer =  $\lambda s.$ 
  offer  $s \left\{ \begin{array}{l} \text{Free} \mapsto \text{let } s = \text{select } \textcircled{1} \text{ } s \text{ in} \\ \quad \text{let } (cb, s) = \text{recv } s \text{ in} \\ \quad \text{wait } s; \text{eat } cb \\ \text{Busy} \mapsto \text{wait } s; \text{hungry} \end{array} \right\}$ 
in let  $s = \text{fork } (\lambda s. \text{vendingMachine } s)$ 
in customer  $s$ 

```



(A) Vending machine and customer as a GV program. (B) Process structure of Figure 1a.

FIGURE 1. Example program with acyclic process structure.

$(\mathcal{C} \parallel (\nu x)\mathcal{C}') \equiv (\nu x)(\mathcal{C} \parallel \mathcal{C}')$ if $x \notin \text{fv}(\mathcal{C})$). They guarantee acyclicity by defining an extrinsic type system for configurations. In particular, the type system requires that in every parallel composition $\mathcal{C} \parallel \mathcal{D}$, configurations $\mathcal{C}$ and $\mathcal{D}$ must have exactly one channel in common, and that in a name restriction $(\nu x)\mathcal{C}$, channel x cannot be used until it is shared across a parallel composition.

These restrictions are sufficient to guarantee deadlock freedom. Unfortunately, they are *not* preserved by process equivalence. As Lindley and Morris write, (noting that their name restrictions bind *channels* rather than endpoint pairs, and their (νxy) abbreviates $(\nu x)(\nu y)$):

Alas, our notion of typing is not preserved by configuration equivalence. For example, assume that $\Gamma \vdash (\nu xy)(C_1 \parallel (C_2 \parallel C_3))$, where $x \in \text{fv}(C_1)$, $y \in \text{fv}(C_2)$, and $x, y \in \text{fv}(C_3)$. We have that $C_1 \parallel (C_2 \parallel C_3) \equiv (C_1 \parallel C_2) \parallel C_3$, but $\Gamma \not\vdash (\nu xy)((C_1 \parallel C_2) \parallel C_3)$, as both x and y must be shared between the processes $C_1 \parallel C_2$ and C_3 .

As a result, standard notions of progress and preservation are not enough to guarantee deadlock freedom, as reduction sequences could include equivalence steps from well-typed to non-well-typed terms. Instead, they must prove a stronger result:

Theorem 3 (Lindley and Morris [LM15]). *If $\Gamma \vdash \mathcal{C}$, $\mathcal{C} \equiv \mathcal{C}'$, and $\mathcal{C}' \longrightarrow \mathcal{D}'$, then there exists $\mathcal{D}$ such that $\mathcal{D} \equiv \mathcal{D}'$ and $\Gamma \vdash \mathcal{D}$.*

This is not a one-time cost: languages based on GV must either also give up on type preservation for structural congruence [FLMD19] or admit deadlocks [ITT⁺19, TV20].

Note that CP only avoids the same issue through its combined $(\nu x)(P \parallel Q)$ term; attempts to split the term into a separate name restriction and parallel composition would also lose typability of equivalence.

3. HYPERSEQUENT GV

We present Hypersequent GV (HGV), a linear λ -calculus extended with session types and primitives for session-typed communication. HGV shares its syntax and static typing with GV, but uses hyper-environments for runtime typing to simplify and generalise its semantics.

Typing rules for terms

Typing rules for terms				$\boxed{\Gamma \vdash M : T}$
$\frac{}{x : T \vdash x : T}$	$\frac{}{\cdot \vdash K : T}$	$\frac{\text{TM-LAM}}{\Gamma, x : T \vdash M : U}$	$\frac{\text{TM-APP}}{\Gamma \vdash M : T \multimap U \quad \Delta \vdash N : T}$	
$\frac{}{\cdot \vdash () : \mathbf{1}}$	$\frac{\text{TM-LETUNIT}}{\Gamma, \Delta \vdash \text{let } () = M \text{ in } N : T}$	$\frac{\text{TM-PAIR}}{\Gamma, \Delta \vdash (M, N) : T \times U}$		
$\frac{\text{TM-LETPAIR}}{\Gamma, \Delta \vdash \text{let } (x, y) = M \text{ in } N : U}$	$\frac{\text{TM-ABSURD}}{\Gamma \vdash \text{absurd } M : T}$	$\frac{\text{TM-INL}}{\Gamma \vdash \text{inl } M : T + U}$		
$\frac{\text{TM-INR}}{\Gamma \vdash \text{inr } M : T + U}$	$\frac{\text{TM-CASESUM}}{\Gamma, \Delta \vdash \text{case } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} : U}$			

Type schemas for communication primitives

$\boxed{K : T}$		
$\text{link} : S \times \bar{S} \multimap \text{end}_!$	$\text{send} : T \times !T.S \multimap S$	$\text{wait} : \text{end}_? \multimap \mathbf{1}$
$\text{fork} : (S \multimap \text{end}_!) \multimap \bar{S}$	$\text{recv} : ?T.S \multimap T \times S$	

Duality

$\boxed{\bar{S}}$			
$\overline{!T.S} = ?T.\bar{S}$	$\overline{?T.S} = !T.\bar{S}$	$\overline{\text{end}_!} = \text{end}_?$	$\overline{\text{end}_?} = \text{end}_!$

FIGURE 2. HGV, duality and typing rules for terms.

Types, terms, and static typing. Types (T, U) comprise a unit type ($\mathbf{1}$), an empty type ($\mathbf{0}$), product types ($T \times U$), sum types ($T + U$), linear function types ($T \multimap U$), and session types (S).

$$T, U ::= \mathbf{1} \mid \mathbf{0} \mid T \times U \mid T + U \mid T \multimap U \mid S \quad S ::= !T.S \mid ?T.S \mid \text{end}_! \mid \text{end}_?$$

Session types (S) comprise output ($!T.S$: send a value of type T , then behave like S), input ($?T.S$: receive a value of type T , then behave like S), and dual end types ($\text{end}_!$ and $\text{end}_?$). The dual endpoints restrict process structure to *trees* [Wad14]; conflating them loosens this restriction to *forests* [ALM16]. We let Γ, Δ range over type environments.

The terms and typing rules are given in Figure 2. The linear λ -calculus rules are standard; communication primitives K are given as constants. Each communication primitive K has a type schema: **link** takes a pair of compatible endpoints and forwards all messages between them; **fork** takes a function, which is passed one endpoint (of type S) of a fresh channel yielding a new child thread, and returns the other endpoint (of type $\bar{S}$); **send** takes a pair of a value and an endpoint, sends the value over the endpoint, and returns an updated endpoint; **recv** takes an endpoint, receives a value over the endpoint, and returns the pair of the received value and an updated endpoint; and **wait** synchronises on a terminated endpoint of type $\text{end}_?$. Output is dual to input, and $\text{end}_!$ is dual to $\text{end}_?$. Duality is involutive, i.e., $\bar{\bar{S}} = S$.

We write $M; N$ for **let** $() = M$ **in** N , **let** $x = M$ **in** N for $(\lambda x.N) M$, $\lambda().M$ for $\lambda z.z; M$, and $\lambda(x, y).M$ for $\lambda z.\text{let } (x, y) = z \text{ in } M$. We write $K : T$ for $\cdot \vdash K : T$ in typing derivations.

Typing rules for configurations

$$\boxed{\mathcal{G} \vdash \mathcal{C} : R}$$

$$\frac{\text{TC-NEW} \quad \mathcal{G} \parallel \Gamma, x : S \parallel \Delta, y : \bar{S} \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma, \Delta \vdash (\nu xy) \mathcal{C} : R}$$

$$\frac{\text{TC-PAR} \quad \mathcal{G} \vdash \mathcal{C} : R \quad \mathcal{H} \vdash \mathcal{D} : R'}{\mathcal{G} \parallel \mathcal{H} \vdash \mathcal{C} \parallel \mathcal{D} : R \sqcap R'}$$

$$\frac{\text{TC-MAIN} \quad \Gamma \vdash M : T}{\Gamma \vdash \bullet M : \bullet T}$$

$$\frac{\text{TC-CHILD} \quad \Gamma \vdash M : \text{end}_!}{\Gamma \vdash \circ M : \circ}$$

$$\frac{\text{TC-LINK}}{x : S, y : \bar{S}, z : \text{end}_? \vdash x \xleftrightarrow{z} y : \circ}$$

Configuration types

$$R ::= \circ \mid \bullet T$$

Configuration type combination

$$\bullet T \sqcap \circ = \bullet T \quad \circ \sqcap \bullet T = \bullet T \quad \circ \sqcap \circ = \circ$$

$$\boxed{R \sqcap R'}$$

FIGURE 3. HGV, typing rules for configurations.

Remark 3.1. We include **link** because it is convenient for the correspondence with CP, which interprets CLL's axiom as forwarding. We *can* encode **link** in GV via a type directed translation akin to CLL's *identity expansion*.

Configurations and runtime typing. Process configurations $(\mathcal{C}, \mathcal{D}, \mathcal{E})$ comprise child threads $(\circ M)$, the main thread $(\bullet M)$, link threads $(x \xleftrightarrow{z} y)$, name restrictions $((\nu xy) \mathcal{C})$, and parallel compositions $(\mathcal{C} \parallel \mathcal{D})$. We refer to a configuration of the form $\circ M$ or $x \xleftrightarrow{z} y$ as an *auxiliary thread*, and a configuration of the form $\bullet M$ as a *main thread*. We let $\mathcal{A}$ range over auxiliary threads and $\mathcal{T}$ range over all threads (auxiliary or main).

$$\phi ::= \bullet \mid \circ \quad \mathcal{C}, \mathcal{D}, \mathcal{E} ::= \phi M \mid x \xleftrightarrow{z} y \mid \mathcal{C} \parallel \mathcal{D} \mid (\nu xy) \mathcal{C}$$

The configuration language is reminiscent of π -calculus processes, but has some non-standard features. Name restriction uses double binders [Vas12] in which one name is bound to each endpoint of the channel. Link threads [LM16] handle forwarding. A link thread $x \xleftrightarrow{z} y$ waits for the thread connected to z to terminate before forwarding all messages between x and y .

Configuration typing departs from GV [LM15], exploiting *hypersequents* [Avr91] to recover modularity and extensibility. Inspired by HCP [MP18, KMP19b, KMP19a], configurations are typed under a *hyper-environment*, an unordered collection of disjoint type environments. We let $\mathcal{G}, \mathcal{H}$ range over hyper-environments, writing $\emptyset$ for the empty hyper-environment, $\mathcal{G} \parallel \Gamma$ for disjoint extension of $\mathcal{G}$ with type environment Γ , and $\mathcal{G} \parallel \mathcal{H}$ for disjoint concatenation of $\mathcal{G}$ and $\mathcal{H}$.

The typing rules for configurations are given in Figure 3. Rules TC-NEW and TC-PAR are key to deadlock freedom: TC-NEW joins two disjoint configurations with a new channel, and merges their type environments; TC-PAR combines two disjoint configurations, and registers their disjointness by separating their type environments in the hyper-environment. Rules TC-MAIN, TC-CHILD, and TC-LINK type main, child, and link threads, respectively; all three require a singleton hyper-environment. A configuration has type $\circ$ if it has no main thread, and $\bullet T$ if it has a main thread of type T . The configuration type combination operator ensures that a well-typed configuration has at most one main thread.

Operational semantics. Figure 4 gives the operational semantics for HGV, presented as a deterministic reduction relation on terms and a nondeterministic reduction relation on configurations. HGV values (U, V, W) , evaluation contexts (E) , and term reduction rules

Values and evaluation contexts

Values	U, V, W	$::=$	$K \mid \lambda x.M \mid () \mid (V, W) \mid \mathbf{inl} V \mid \mathbf{inr} V$
Evaluation contexts	E	$::=$	$\square$ $\mid E M \mid V E$ $\mid \mathbf{let} () = E \mathbf{in} M$ $\mid (E, M) \mid (V, E) \mid \mathbf{let} (x, y) = E \mathbf{in} M$ $\mid \mathbf{inl} E \mid \mathbf{inr} E \mid \mathbf{case} E \{\mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N\}$
Thread contexts	F	$::=$	ϕE

Term reduction

$$M \longrightarrow_M N$$

E-LAM	$(\lambda x.M) V$	$\longrightarrow_M$	$M\{V/x\}$
E-UNIT	$\mathbf{let} () = () \mathbf{in} M$	$\longrightarrow_M$	M
E-PAIR	$\mathbf{let} (x, y) = (V, W) \mathbf{in} M$	$\longrightarrow_M$	$M\{V/x, W/y\}$
E-INL	$\mathbf{case} \mathbf{inl} V \{\mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N\}$	$\longrightarrow_M$	$M\{V/x\}$
E-INR	$\mathbf{case} \mathbf{inr} V \{\mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N\}$	$\longrightarrow_M$	$N\{V/y\}$
E-LIFT	$E[M]$	$\longrightarrow_M$	$E[N]$, if $M \longrightarrow_M N$

Structural congruence

$$\mathcal{C} \equiv \mathcal{D}$$

SC-PARASSOC	$\mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E}) \equiv (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}$	SC-PARCOMM	$\mathcal{C} \parallel \mathcal{D} \equiv \mathcal{D} \parallel \mathcal{C}$
SC-NEWCOMM	$(\nu xy)(\nu zw)\mathcal{C} \equiv (\nu zw)(\nu xy)\mathcal{C}$	SC-NEWSWAP	$(\nu xy)\mathcal{C} \equiv (\nu yx)\mathcal{C}$
SC-SCOPEEXT	$(\nu xy)(\mathcal{C} \parallel \mathcal{D}) \equiv \mathcal{C} \parallel (\nu xy)\mathcal{D}$, if $x, y \notin \text{fv}(\mathcal{C})$	SC-LINKCOMM	$x \xleftrightarrow{\mathcal{C}} y \equiv y \xleftrightarrow{\mathcal{C}} x$

Configuration reduction

$$\mathcal{C} \longrightarrow \mathcal{D}$$

E-REIFY-FORK	$F[\mathbf{fork} V]$	$\longrightarrow$	$(\nu xx')(F[x] \parallel \circ (V x'))$, where x, x' fresh
E-REIFY-LINK	$F[\mathbf{link} (x, y)]$	$\longrightarrow$	$(\nu zz')(x \xleftrightarrow{\mathcal{C}} y \parallel F[z'])$, where z, z' fresh
E-COMM-LINK	$(\nu zz')(\nu xx')(x \xleftrightarrow{\mathcal{C}} y \parallel \circ z' \parallel \phi M)$	$\longrightarrow$	$\phi (M\{y/x'\})$
E-COMM-SEND	$(\nu xy)(F[\mathbf{send} (V, x)] \parallel F'[\mathbf{recv} y])$	$\longrightarrow$	$(\nu xy)(F[x] \parallel F'[(V, y)])$
E-COMM-CLOSE	$(\nu xy)(\circ y \parallel F[\mathbf{wait} x])$	$\longrightarrow$	$F[()]$

E-RES	$\frac{\mathcal{C} \longrightarrow \mathcal{C}'}{(\nu xy)\mathcal{C} \longrightarrow (\nu xy)\mathcal{C}'}$	E-PAR	$\frac{\mathcal{C} \longrightarrow \mathcal{C}'}{\mathcal{C} \parallel \mathcal{D} \longrightarrow \mathcal{C}' \parallel \mathcal{D}}$	E-EQUIV	$\frac{\mathcal{C} \equiv \mathcal{C}' \quad \mathcal{C}' \longrightarrow \mathcal{D}' \quad \mathcal{D}' \equiv \mathcal{D}}{\mathcal{C} \longrightarrow \mathcal{D}}$	E-LIFT-M	$\frac{M \longrightarrow_M M'}{F[M] \longrightarrow F[M']}$
-------	---	-------	---	---------	---	----------	---

FIGURE 4. HGV, operational semantics.

$(\longrightarrow_M)$ define a standard call-by-value, left-to-right evaluation strategy. A closed term either reduces to a value or is blocked on a communication action.

Thread contexts (F) extend evaluation contexts to threads. The structural congruence rules are standard apart from SC-LINKCOMM, which ensures links are undirected, and SC-NEWSWAP, which swaps names in double binders.

The configuration reduction relation gives a semantics for HGV's communication and concurrency constructs. The first two rules, E-REIFY-FORK and E-REIFY-LINK, create child and link threads, respectively. The next three rules, E-COMM-LINK, E-COMM-SEND, and E-COMM-CLOSE perform communication actions. The final four rules enable reduction under name restriction and parallel composition, rewriting by structural congruence, and term reduction in threads. Two rules handle links: E-REIFY-LINK creates a new *link thread* $x \xleftrightarrow{\mathcal{C}} y$ which blocks on z of type $\mathbf{end}_?$, one endpoint of a fresh channel. The other endpoint, z' of

type **end_!**, is placed in the evaluation context of the parent thread. When z' terminates a child thread, E-COMM-LINK performs forwarding by substitution.

Remark 3.2. Note that E-COMM-LINK does not fire if z' is returned by a main thread. In closed configurations, typing ensures that such a configuration cannot arise: intuitively, a main thread can only obtain endpoints by **fork** or by receiving an endpoint.

Endpoints generated to communicate with forked threads (i.e., those passed to a child thread) will always have a session type terminating with **end_?**, and a child thread cannot transmit an endpoint ending in **end_!**, since the endpoint must be returned. Consequently, there is no way for a main thread to obtain endpoints with dual session types as required by the type of **link**. The case for open configurations is accounted for by our open progress result (see Section 3.1).

Choice. HGV does not include constructs for internal and external choice (for example, as shown in the vending machine example in Section 1). Internal and external choice are instead encoded with sum types and session delegation [Kob03, DGS17]. Prior encodings of choice in GV [LM15] are asynchronous. Instead, to encode synchronous choice we add a ‘dummy’ synchronisation before exchanging the value of sum type, as follows:

$$\begin{array}{ll}
 S \oplus S' \triangleq !1.!(\overline{S_1} + \overline{S_2}).\mathbf{end}_! & \mathbf{select} \ell \triangleq \lambda x. \left(\mathbf{let} \ x = \mathbf{send} \ (\ell, x) \ \mathbf{in} \right. \\
 S \& S' \triangleq ?1.?(S_1 + S_2).\mathbf{end}_? & \left. \mathbf{fork} \ (\lambda y. \mathbf{send} \ (\ell \ y, x)) \right) \\
 \oplus\{\} \triangleq !1.!0.\mathbf{end}_! & \mathbf{offer} \ L \ \{\mathbf{inl} \ x \mapsto M; \mathbf{inr} \ y \mapsto N\} \\
 \&\{\} \triangleq ?1.?0.\mathbf{end}_? & \triangleq \mathbf{let} \ (\ell, z) = \mathbf{recv} \ L \ \mathbf{in} \ \mathbf{let} \ (w, z) = \mathbf{recv} \ z \\
 & \mathbf{in} \ \mathbf{wait} \ z; \mathbf{case} \ w \ \{\mathbf{inl} \ x \mapsto M; \mathbf{inr} \ y \mapsto N\} \\
 & \mathbf{offer} \ L \ \{\} \triangleq \mathbf{let} \ (\ell, c) = \mathbf{recv} \ L \ \mathbf{in} \ \mathbf{let} \ (z, c) = \mathbf{recv} \ c \\
 & \mathbf{in} \ \mathbf{wait} \ c; \mathbf{absurd} \ z
 \end{array}$$

3.1. Metatheory. HGV enjoys type preservation, deadlock freedom, confluence, and strong normalisation.

Preservation. Hyper-environments enable type preservation under structural congruence, which significantly simplifies the metatheory compared to GV.

Theorem 3.3 (Preservation).

- (1) If $\mathcal{G} \vdash C : R$ and $C \equiv D$, then $\mathcal{G} \vdash D : R$.
- (2) If $\mathcal{G} \vdash C : R$ and $C \longrightarrow D$, then $\mathcal{G} \vdash D : R$.

Proof. By induction on the derivations of $C \equiv D$ and $C \longrightarrow D$. See Appendix A. □

Before moving onto progress, we must introduce some technical machinery to allow us to reason about the structure of HGV programs.

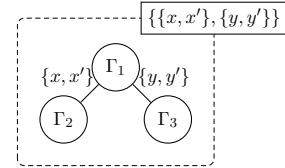
Abstract process structures. Unlike in GV, in HGV we cannot rely on the fact that exactly one channel is split over each parallel composition. Instead, we introduce the notion of an *abstract process structure* (APS). Abstract process structures are a crucial ingredient in showing that HGV configurations can be written in *tree canonical form*, which helps both with establishing progress results and also the correspondence between HGV and GV.

We begin by establishing the intuition behind the notion of an APS, and then describe the formal definitions. An APS is a graph defined over a hyper-environment $\mathcal{G}$ and a set of undirected pairs of co-names (a *co-name set*) $\mathcal{N}$ drawn from the names in $\mathcal{G}$.

The nodes of an APS are the type environments in $\mathcal{G}$. Each edge is labelled by a distinct co-name pair $\{x_1, x_2\} \in \mathcal{N}$, such that $x_1 : S \in \Gamma_1$ and $x_2 : \bar{S} \in \Gamma_2$.

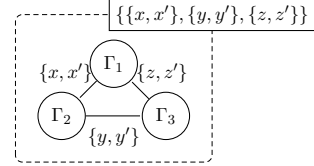
Example 3.4.

Let $\mathcal{G} = \Gamma_1 \parallel \Gamma_2 \parallel \Gamma_3$, where $\Gamma_1 = x : S_1, y : S_2$, $\Gamma_2 = x' : \bar{S}_1, z : T$, and $\Gamma_3 = y' : \bar{S}_2$, and suppose $\mathcal{N} = \{\{x, x'\}, \{y, y'\}\}$. The APS for $\mathcal{G}$ and $\mathcal{N}$ is illustrated to the right.



Example 3.5.

Let $\mathcal{G} = \Gamma_1 \parallel \Gamma_2 \parallel \Gamma_3$, where $\Gamma_1 = x : S_1, z' : \bar{S}_3$, and $\Gamma_2 = x' : \bar{S}_1, y : S_2$, and $\Gamma_3 = y' : \bar{S}_2, z : S_3$, and suppose $\mathcal{N} = \{\{x, x'\}, \{y, y'\}, \{z, z'\}\}$. The APS for $\mathcal{G}$ and $\mathcal{N}$ is illustrated to the right.



Let us now discuss the formal definition of an APS. We begin by recalling the definition of an *undirected edge-labelled multigraph*: an undirected graph that allows multiple edges between vertices.

Definition 3.6 (Undirected Multigraph). An *undirected multigraph* G is a 3-tuple $(\mathcal{V}, \mathcal{E}, r)$ where:

- (1) $\mathcal{V}$ is a set of vertices
- (2) $\mathcal{E}$ is a set of edge names
- (3) r is a function $r : \mathcal{E} \mapsto \{\{v, w\} : v, w \in \mathcal{V}\}$ from edge names to an unordered pair of vertices

Denote the size of a set as $|\cdot|$. A *path* is a sequence of edges connecting two vertices. A multigraph $G = (\mathcal{V}, \mathcal{E}, r)$ is *connected* if $|\mathcal{V}| = 1$, or if for every pair of vertices $v, w \in \mathcal{V}$ there is a path between v and w . A multigraph is *acyclic* if no path forms a cycle. A *leaf* is a vertex connected to the remainder of a graph by a single edge.

Definition 3.7 (Leaf). Given an undirected multigraph $(\mathcal{V}, \mathcal{E}, r)$, a vertex $v \in \mathcal{V}$ is a *leaf* if there exists a single $e \in \mathcal{E}$ such that $v \in r(e)$.

In an undirected tree containing at least two vertices, there must be at least two leaves.

Lemma 3.8. *If $G = (\mathcal{V}, \mathcal{E}, r)$ is an undirected tree where $|V| \geq 2$, then there exist at least two leaves in $\mathcal{V}$.*

Proof. For G to be an undirected tree where $|V| \geq 2$ and have fewer than two leaves, then there would need to be a cycle, contradicting acyclicity. $\square$

With the graph preliminaries in place, we are now ready to introduce the formal definition of an APS.

Definition 3.9 (Abstract process structure). The *abstract process structure* of a hyper-environment $\mathcal{H}$ with respect to a co-name set $\mathcal{N} = \{\{x_1, y_1\}, \dots, \{x_n, y_n\}\}$ is an undirected multigraph $(\mathcal{V}, \mathcal{E}, r)$ defined as follows:

- (1) $\mathcal{V} = \text{envs}(\mathcal{H})$
- (2) $\mathcal{E} = \mathcal{N}$
- (3) $r = (\{x, y\} \mapsto \{\Gamma_1, \Gamma_2\})$ for each $\{x, y\} \in \mathcal{N}$ such that $\Gamma_1 \in \text{envs}(\mathcal{H}), \Gamma_2 \in \text{envs}(\mathcal{H}), x \in \text{fv}(\Gamma_1), y \in \text{fv}(\Gamma_2)$

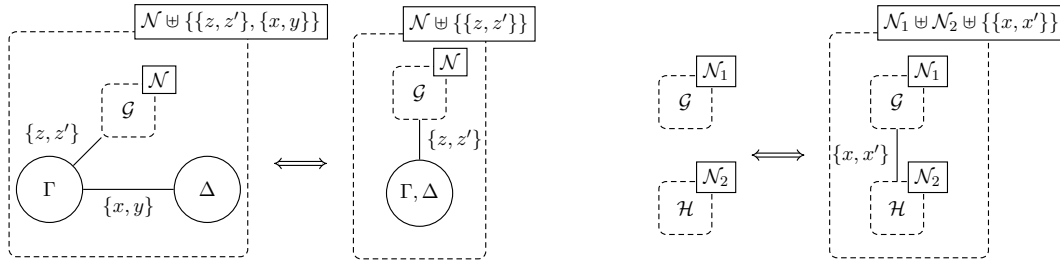
Example 3.10. The formal definition of the APS described in Example 3.4 is defined as:

- $\mathcal{V} = \{\Gamma_1, \Gamma_2, \Gamma_3\}$
- $\mathcal{E} = \{\{x, x'\}, \{y, y'\}\}$
- $r(\{x, x'\}) \mapsto \{\Gamma_1, \Gamma_2\}$
- $r(\{y, y'\}) \mapsto \{\Gamma_1, \Gamma_3\}$

Whereas Example 3.4 is a tree, Example 3.5 contains a cycle. Only configurations typeable under a hyper-environment with a *tree structure* can be written in tree canonical form.

Definition 3.11 (Tree structure). A hyper-environment $\mathcal{H}$ with co-name set $\mathcal{N}$ has a *tree structure*, written $\text{Tree}(\mathcal{H}, \mathcal{N})$, if its APS is connected and acyclic.

An HGV program $\bullet M$ has a single type environment, so is tree-structured; the same goes for child and link threads. A key feature of HGV is a subformula principle, which states that all hyper-environments arising in the derivation of an HGV program are tree-structured. It follows that a configuration resulting from the reduction of an HGV program is also tree structured. Read bottom-up, TC-NEW and TC-PAR preserve tree structure, which is illustrated by the following two pictures.



The following lemma states this intuition formally. By analogy to Kleene equality, we write $\mathcal{P} \stackrel{\hat{=}}{\iff} \mathcal{Q}$, to mean that either $\mathcal{P}$ or $\mathcal{Q}$ is undefined, or $\mathcal{P} \iff \mathcal{Q}$.

Lemma 3.12 (Tree structure).

- $\text{Tree}((\mathcal{H} \parallel \Gamma_1, x_1 : S \parallel \Gamma_2, x_2 : \bar{S}), \mathcal{N} \uplus \{\{x_1, x_2\}\}) \stackrel{\hat{=}}{\iff} \text{Tree}((\mathcal{H} \parallel \Gamma_1, \Gamma_2), \mathcal{N})$
- $\text{Tree}((\mathcal{H}_1 \parallel \Gamma_1, x_1 : S), \mathcal{N}_1) \wedge \text{Tree}((\mathcal{H}_2 \parallel \Gamma_2, x_2 : \bar{S}), \mathcal{N}_2) \stackrel{\hat{=}}{\iff} \text{Tree}((\mathcal{H}_1 \parallel \Gamma_1, x_1 : S \parallel \mathcal{H}_2 \parallel \Gamma_2, x_2 : \bar{S}), \mathcal{N}_1 \uplus \mathcal{N}_2 \uplus \{\{x_1, x_2\}\})$

Proof. By the definition of $\stackrel{\hat{=}}{\iff}$, we need only consider the cases where both sides of the bi-implication are defined. Both results follow from the observation that adding an edge between two trees results in a tree, and removing an edge from a tree partitions the tree into two subtrees. $\square$

Tree canonical form. We now define a canonical form for configurations that captures the tree structure of an APS. Tree canonical form enables a succinct statement of *open progress* (Lemma 3.17) and a means for embedding HGV in GV (Proposition 4.5).

Definition 3.13 (Tree canonical form). A configuration $\mathcal{C}$ is in *tree canonical form* if it can be written: $(\nu x_1 y_1)(\mathcal{A}_1 \parallel \cdots \parallel (\nu x_n y_n)(\mathcal{A}_n \parallel \phi N) \cdots)$ where $x_i \in \text{fv}(\mathcal{A}_i)$ for $1 \leq i \leq n$.

Every well-typed HGV configuration typeable under a single type environment can be written in tree canonical form.

Theorem 3.14 (Well-typed configurations in tree canonical forms). *If $\Gamma \vdash \mathcal{C} : R$, then there exists some $\mathcal{D}$ such that $\mathcal{C} \equiv \mathcal{D}$ and $\mathcal{D}$ is in tree canonical form.*

Proof. By induction on the number of ν -binders in $\mathcal{C}$. In the case that $n = 0$, it must be the case that $\Gamma \vdash \phi M : R$ for some thread M , since parallel composition is only typeable under a hyper-environment containing two or more type environments. Therefore, $\mathcal{C}$ is in tree canonical form by definition.

In the case that $n \geq 1$, by Theorem 3.3, we can rewrite the configuration as:

$$(\nu x_1 y_1) \cdots (\nu x_n y_n)(\circ M_1 \parallel \cdots \parallel \circ M_n \parallel \phi N)$$

Fix $\mathcal{N} = \{\{x_i, y_i\} \mid 1 \leq i \leq n\}$. By definition, Γ has a tree structure with respect to an empty co-name set. By repeated applications of TC-NEW, there exists some $\mathcal{G}$ such that $\mathcal{G} \vdash \circ M_1 \parallel \cdots \parallel \circ M_n \parallel \phi N : T$; by Lemma 3.12 (clause 1, right-to-left), $\mathcal{G}$ has a tree structure.

Construct the APS for $\mathcal{G}$ using names $\mathcal{N}$; by Lemma 3.8, there exist $\Gamma_1, \Gamma_2 \in \text{envs}(\mathcal{H})$ such that Γ_1 and Γ_2 are leaves of the tree and therefore by the definition of the APS contain precisely one ν -bound name. By TC-PAR, there must exist two threads $\mathcal{C}_1, \mathcal{C}_2$ such that $\Gamma_1 \vdash \mathcal{C}_1 : R_1$ and $\Gamma_2 \vdash \mathcal{C}_2 : R_2$. By runtime type combination, at least one of R_1, R_2 must be $\circ$; without loss of generality assume this is R_1 . Suppose (again without loss of generality) that the ν -bound name contained in Γ_1 is x_1 and $L_1 = M_1$.

Let $\mathcal{D} = (\nu x_2 y_2) \cdots (\nu x_n y_n)(\circ M_2 \parallel \cdots \parallel \circ M_n \parallel \phi N)$. By Theorem 3.3 and the fact that x_1 is the only ν -bound variable in M_1 , we have that $\mathcal{C} \equiv (\nu x_1 y_1)(\circ M_1 \parallel \mathcal{D})$. By the induction hypothesis, there exists some $\mathcal{D}'$ such that $\mathcal{D} \equiv \mathcal{D}'$ and $\mathcal{D}'$ is in canonical form. By construction we have that $\mathcal{C} \equiv (\nu x_1 y_1)(\circ M_1 \parallel \mathcal{D}')$, which is in tree canonical form as required. $\square$

As hyper-environments capture parallelism, a configuration $\mathcal{C}$ typeable under hyper-environment $\Gamma_1 \parallel \cdots \parallel \Gamma_n$ is equivalent to n independent parallel processes.

Proposition 3.15 (Independence). *If $\Gamma_1 \parallel \cdots \parallel \Gamma_n \vdash \mathcal{C} : R$, then there exist $R_1, \dots, R_n$ and $\mathcal{D}_1, \dots, \mathcal{D}_n$ such that $R = R_1 \sqcap \cdots \sqcap R_n$ and $\mathcal{C} \equiv \mathcal{D}_1 \parallel \cdots \parallel \mathcal{D}_n$ and $\Gamma_i \vdash \mathcal{D}_i : R_i$ for each i .*

Proof. By induction on the derivation of $\Gamma_1 \parallel \cdots \parallel \Gamma_n \vdash \mathcal{C} : R$. The cases for TC-MAIN, TC-CHILD, and TC-LINK follow immediately. The cases for TC-NEW and TC-PAR follow from the IH and structural congruence rules. $\square$

It follows from Theorem 3.14 and Proposition 3.15 that any well-typed HGV configuration can be written as a forest of independent configurations in tree canonical form.

Progress and Deadlock Freedom. With tree canonical forms defined, we can now state a progress result. A thread is *blocked* on an endpoint x if it is ready to perform a communication action on x .

Definition 3.16 (Blocked thread). We say that thread $\mathcal{T}$ is *blocked on variable z* , written $\text{blocked}(\mathcal{T}, z)$, if either: $\mathcal{T} = \circ z$; $\mathcal{T} = x \xrightarrow{z} y$, for some x, y ; or $\mathcal{T} = F[N]$ for some F , where N is **send** (V, z) , **recv** z , or **wait** z .

We let Ψ range over type environments containing only session-typed variables, *i.e.*, $\Psi ::= \cdot \mid \Psi, x : S$, which lets us reason about configurations that are closed except for runtime names. Using Lemma 3.17 we obtain *open progress* for configurations with free runtime names.

Lemma 3.17 (Open Progress). *Suppose $\Psi \vdash \mathcal{C} : T$ where*

$$\mathcal{C} = (\nu x_1 y_1)(\mathcal{A}_1 \parallel \cdots \parallel (\nu x_n y_n)(\mathcal{A}_n \parallel \phi N) \cdots)$$

is in tree canonical form. Either $\mathcal{C} \longrightarrow \mathcal{D}$ for some $\mathcal{D}$, or:

- (1) *For each $\mathcal{A}_i$ ($1 \leq i \leq n$), $\text{blocked}(\mathcal{A}_i, z)$ for some $z \in \{x_i\} \cup \{y_j \mid 1 \leq j < i\} \cup \text{fv}(\Psi)$*
- (2) *Either N is a value or $\text{blocked}(\phi N, z)$ for some $z \in \{y_i \mid 1 \leq i \leq n\} \cup \text{fv}(\Psi)$*

Proof. Open progress follows as a direct corollary of a slightly more verbose property which holds on HGV processes, proved by induction on the derivation of an inductive definition of tree canonical forms. See Appendix A for details. $\square$

Closed configurations enjoy a stronger result: if a closed configuration cannot reduce, then each auxiliary thread must either be a value, or be blocked on its neighbouring endpoint.

Lemma 3.18 (Closed Progress). *Suppose $\Psi \vdash \mathcal{C} : R$ where*

$$\mathcal{C} = (\nu x_1 y_1)(\mathcal{A}_1 \parallel \cdots \parallel (\nu x_n y_n)(\mathcal{A}_n \parallel \phi N) \cdots)$$

is in tree canonical form. Either $\mathcal{C} \longrightarrow \mathcal{D}$ for some $\mathcal{D}$, or:

- (1) *For each $\mathcal{A}_j$ for $1 \leq j \leq n$, $\text{blocked}(\mathcal{A}_j, x_j)$*
- (2) *N is a value*

Proof. Since the environment is closed, by Lemma 3.17, for each $\mathcal{A}_j$ it must be that $\text{blocked}(\mathcal{A}_j, z)$ for some $z \in \{y_i \mid i \in 1..j-1\} \cup \{x_j\}$.

Note that if two names x, y are co-names, and one thread is blocked on x , and another is blocked on y , then due to typing the names must be dual and reduction can occur.

Consider $\mathcal{A}_1$. Since the environment is closed, $\mathcal{A}_1$ must be blocked on x_1 . Next, consider $\mathcal{A}_2$; the thread cannot be blocked on y_1 as reduction would occur. By the definition of tree canonical forms, $\mathcal{A}_2$ must contain x_2 and by the typing rules cannot contain y_2 , so the thread must be blocked on x_2 . The argument extends to the remainder of the configuration. $\square$

Finally, for *ground configurations*, where the main thread does not return a runtime name or capture a runtime name in a closure, we obtain a yet tighter result, *global progress*, which implies deadlock freedom [CDM14].

Definition 3.19 (Ground configuration). A configuration $\mathcal{C}$ is a *ground configuration* if $\cdot \vdash \mathcal{C} : T$, $\mathcal{C}$ is in canonical form, and T does not contain session types or function types.

Our main progress result states that a ground configuration can reduce, or is a value.

Typing rules for configurations

$$\boxed{\Gamma \vdash_{\text{GV}} \mathcal{C} : T}$$

$$\begin{array}{c}
\text{TG-NEW} \\
\frac{\Gamma, \langle x, y \rangle : S^\sharp \vdash_{\text{GV}} \mathcal{C} : R}{\Gamma \vdash_{\text{GV}} (\nu xy)\mathcal{C} : R}
\end{array}
\quad
\begin{array}{c}
\text{TG-CONNECT}_1 \\
\frac{\Gamma_1, x : S \vdash_{\text{GV}} \mathcal{C} : R \quad \Gamma_2, y : \bar{S} \vdash_{\text{GV}} \mathcal{D} : R'}{\Gamma_1, \Gamma_2, \langle x, y \rangle : S^\sharp \vdash_{\text{GV}} \mathcal{C} \parallel \mathcal{D} : R \sqcap R'}
\end{array}
\quad
\begin{array}{c}
\text{TG-CONNECT}_2 \\
\frac{\Gamma_1, y : \bar{S} \vdash_{\text{GV}} \mathcal{C} : R \quad \Gamma_2, x : S \vdash_{\text{GV}} \mathcal{D} : R'}{\Gamma_1, \Gamma_2, \langle x, y \rangle : S^\sharp \vdash_{\text{GV}} \mathcal{C} \parallel \mathcal{D} : R \sqcap R'}
\end{array}$$

$$\begin{array}{c}
\text{TG-CHILD} \\
\frac{\Gamma \vdash_{\text{GV}} M : \text{end}_!}{\Gamma \vdash_{\text{GV}} \circ M : \circ}
\end{array}
\quad
\begin{array}{c}
\text{TG-MAIN} \\
\frac{\Gamma \vdash_{\text{GV}} M : T}{\Gamma \vdash_{\text{GV}} \bullet M : \bullet T}
\end{array}
\quad
\begin{array}{c}
\text{TG-LINK} \\
\frac{}{x : S, y : \bar{S}, z : \text{end}_? \vdash_{\text{GV}} x \overset{z}{\leftrightarrow} y : \circ}
\end{array}$$

FIGURE 5. GV, typing rules for configurations.

Theorem 3.20 (Global progress). *Suppose $\mathcal{C}$ is a ground configuration. Either there exists some $\mathcal{D}$ such that $\mathcal{C} \longrightarrow \mathcal{D}$, or $\mathcal{C} = \bullet V$ for some value V .*

Proof. By Lemma 3.18, either $\mathcal{C}$ can reduce, or $\mathcal{C}$ can be written:

$$(\nu x_1 y_1)(\circ \mathcal{A}_1 \parallel \cdots \parallel (\nu x_n y_n)(\circ \mathcal{A}_n \parallel \bullet V) \cdots)$$

where $\text{blocked}(\mathcal{A}_i, x_i)$ for each $\{x_i \mid i \in 1..n\}$.

Since $\mathcal{C}$ is ground, $\text{fv}(V) = \emptyset$. By definition, tree canonical form ensures that no cycles are present amongst threads, so no auxiliary thread can be blocked. It follows that if $\mathcal{C} \not\longrightarrow$, then there cannot be any auxiliary threads and thus $\mathcal{C} = \bullet V$ for some value V . $\square$

Determinism and Strong Normalisation. HGV enjoys a strong form of determinism known as the diamond property, and due to linearity it enjoys strong normalisation. Unlike with preservation and progress, the addition of hypersequents does not substantially change the arguments from [LM15].

Theorem 3.21 (Diamond property). *If $\mathcal{G} \vdash \mathcal{C} : T$, $\mathcal{C} \longrightarrow \mathcal{D}$, and $\mathcal{C} \longrightarrow \mathcal{D}'$, then $\mathcal{D} \equiv \mathcal{D}'$.*

Proof. Similar to that of GV [LM15, Fow19]: $\longrightarrow_{\text{M}}$ is deterministic, and due to linearity, any overlapping reductions are separate and may be performed in either order. $\square$

Theorem 3.22 (Termination). *If $\mathcal{G} \vdash \mathcal{C} : T$, there are no infinite sequences $\mathcal{C} \longrightarrow \cdots$.*

Proof. As with GV [LM15, Fow19], due to linearity, HGV has an elementary strong normalisation proof. Let the size of a configuration be the sum of the sizes of all abstract syntax trees of all terms contained in threads. The size of a configuration is invariant under $\equiv$ and strictly decreases under $\longrightarrow$, so no infinite reduction sequences can exist. $\square$

4. RELATION BETWEEN HGV AND GV

In this section, we show that well-typed GV configurations are well-typed HGV configurations, and well-typed HGV configurations with tree structure are well-typed GV configurations.

GV. HGV and GV share a common term language and reduction semantics, so only differ in their runtime typing rules. Figure 5 gives the runtime typing rules for GV. We adapt the rules to use a double-binder formulation to concentrate on the essence of the relationship with HGV, but it is trivial to translate GV with single binders into GV with double binders.

GV uses a pseudo-type $S^\sharp$ to type channels. Unlike endpoints, channels cannot appear in terms. Read bottom-up, rule TG-NEW types a name restriction $(\nu xy)\mathcal{C}$, adding $\langle x, y \rangle : S^\sharp$ to the type environment, which along with TG-CONNECT₁ and TG-CONNECT₂ ensures that a session channel of type S will be split into endpoints x and y over a parallel composition. In turn, this enforces a tree process structure. The remaining typing rules are as in HGV.

A simple embedding of GV into HGV. The simplest embedding of GV in HGV relies on the observation from Section 2 that each parallel composition splits a single channel. Let $\mathcal{C} \parallel_{\langle x, y \rangle} \mathcal{D}$ denote two configurations $\mathcal{C}$ and $\mathcal{D}$ connected by a channel with endpoints x, y . We can write an arbitrary closed GV configuration in the form:

$$\mathcal{C}_1 \parallel_{\langle x_1, y_1 \rangle} \cdots \parallel_{\langle x_{n-2}, y_{n-2} \rangle} \mathcal{C}_{n-1} \parallel_{\langle x_{n-1}, y_{n-1} \rangle} \mathcal{C}_n$$

where each $\mathcal{C}$ does not contain a further parallel composition, and any main thread is in $\mathcal{C}_n$. We can then subsequently embed the configuration in HGV as:

$$(\nu x_1 y_1)(\mathcal{C}_1 \parallel \cdots \parallel (\nu x_{n-2} y_{n-2})(\mathcal{C}_{n-2} \parallel (\nu x_{n-1} y_{n-1})(\mathcal{C}_{n-1} \parallel \mathcal{C}_n)) \cdots)$$

which is well-typed by construction. As a corollary, every well-typed, closed GV configuration is equivalent to a well-typed, closed HGV configuration.

A structure-preserving embedding of GV into HGV. Though the simple embedding of GV into HGV is sound, it does not respect the *intention* of GV. In fact, we can provide a stronger result: every well-typed open GV configuration is exactly a well-typed HGV configuration.

Definition 4.1 (Flattening). Flattening, written $\downarrow$, converts GV type environments and HGV hyper-environments into HGV environments.

$$\begin{array}{lll} \downarrow \cdot & = & \cdot \\ \downarrow (\Gamma, \langle x, x' \rangle : S^\sharp) & = & \downarrow \Gamma, x : S, x' : \bar{S} \\ \downarrow (\Gamma, x : T) & = & \downarrow \Gamma, x : T \end{array} \quad \begin{array}{lll} \downarrow \emptyset & = & \emptyset \\ \downarrow (\mathcal{G} \parallel \Gamma) & = & \downarrow \mathcal{G}, \Gamma \end{array}$$

Definition 4.2 (Splitting). Splitting converts GV type environments into hyper-environments. Given channels $\{\langle x_i, x'_i \rangle : S_i^\sharp\}_{i \in 1..n}$ in Γ , a hyper-environment $\mathcal{G}$ is a *splitting* of Γ if $\downarrow \mathcal{G} = \downarrow \Gamma$ and $\exists \Gamma_1, \dots, \Gamma_{n+1}$ such that $\mathcal{G} = \Gamma_1 \parallel \cdots \parallel \Gamma_{n+1}$, and $\text{Tree}(\mathcal{G}, \{\{x_1, x'_1\}, \dots, \{x_n, x'_n\}\})$.

A well-typed GV configuration is typeable in HGV under a splitting of its type environment.

Theorem 4.3 (Typeability of GV configurations in HGV). *If $\Gamma \vdash_{\text{GV}} \mathcal{C} : R$, then there exists some $\mathcal{G}$ such that $\mathcal{G}$ is a splitting of Γ and $\mathcal{G} \vdash \mathcal{C} : R$.*

Proof. By induction on the derivation of $\Gamma \vdash_{\text{GV}} \mathcal{C} : T$ (see Appendix B). □

Example 4.4. Consider a configuration where a child thread pings the main thread:

$$(\nu xy)(\circ (\text{send } (\text{ping}, x)) \parallel \bullet (\text{let } ((), y) = \text{recv } y \text{ in wait } y))$$

We can write a GV typing derivation as follows:

$$\frac{\frac{x : !1.\text{end!}, \text{ping} : 1 \vdash_{\text{GV}} \circ (\text{send } (\text{ping}, x)) : \circ \quad y : ?1.\text{end?} \vdash_{\text{GV}} \bullet (\text{let } ((), y) = \text{recv } y \text{ in wait } y) : \bullet 1}{\langle x, y \rangle : !1.\text{end!}^\sharp, \text{ping} : 1 \vdash_{\text{GV}} (\nu xy)(\circ (\text{send } (\text{ping}, x)) \parallel \bullet (\text{let } ((), y) = \text{recv } y \text{ in wait } y)) : 1}}{\text{ping} : 1 \vdash_{\text{GV}} (\nu xy)(\circ (\text{send } (\text{ping}, x)) \parallel \bullet (\text{let } ((), y) = \text{recv } y \text{ in wait } y)) : 1}$$

The corresponding HGV derivation is:

$$\frac{\frac{x : !1.\text{end}_!, \text{ping} : 1 \vdash \circ(\text{send}(\text{ping}, x)) : \circ \quad y : ?1.\text{end}_? \vdash \bullet(\text{let}(\langle \rangle, y) = \text{recv } y \text{ in wait } y) : \bullet 1}{x : !1.\text{end}_!, \text{ping} : 1 \parallel y : ?1.\text{end}_? \vdash (\nu xy)(\circ(\text{send}(\text{ping}, x)) \parallel \bullet(\text{let}(\langle \rangle, y) = \text{recv } y \text{ in wait } y)) : \bullet 1}}{\text{ping} : 1 \vdash (\nu xy)(\circ(\text{send}(\text{ping}, x)) \parallel \bullet(\text{let}(\langle \rangle, y) = \text{recv } y \text{ in wait } y)) : \bullet 1}$$

Note that $x : !1.\text{end}_!, \text{ping} : 1 \parallel y : ?1.\text{end}_?$ is a splitting of $\langle x, y \rangle : (!1.\text{end}_!)^\sharp, \text{ping} : 1$.

Translating HGV to GV. As we saw in §2, unlike in HGV, equivalence in GV is not type-preserving. It follows that HGV types strictly more processes than GV. Let us revisit Lindley and Morris' example from §1 (adapted to use double-binders), where $\Gamma_1, \Gamma_2, \Gamma_3 \vdash_{\text{GV}} (\nu xx')(\nu yy')(\mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E})) : R_1 \sqcap R_2 \sqcap R_3$ with $\Gamma_1, x : S \vdash_{\text{GV}} \mathcal{C} : R_1$, $\Gamma_2, y : S' \vdash_{\text{GV}} \mathcal{D} : R_2$, and $\Gamma_3, x' : \bar{S}, y' : \bar{S}' \vdash_{\text{GV}} \mathcal{E} : R_3$.

The structurally-equivalent term $(\nu xx')(\nu yy')((\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E})$ is not typeable in GV, since we cannot split both channels over a single parallel composition:

$$\frac{\frac{\Gamma_1, \Gamma_2, x : S \not\vdash_{\text{GV}} \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2 \quad \Gamma_3, x' : \bar{S}, \langle y, y' \rangle : S^\sharp \not\vdash_{\text{GV}} \mathcal{E} : R_3}{\Gamma_1, \Gamma_2, \Gamma_3, \langle x, x' \rangle : S^\sharp, \langle y, y' \rangle : S^\sharp \not\vdash_{\text{GV}} (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E} : R_1 \sqcap R_2 \sqcap R_3}}{\frac{\Gamma_1, \Gamma_2, \Gamma_3, \langle x, x' \rangle : S^\sharp \not\vdash_{\text{GV}} (\nu yy')((\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}) : R_1 \sqcap R_2 \sqcap R_3}{\Gamma_1, \Gamma_2, \Gamma_3 \not\vdash_{\text{GV}} (\nu xx')(\nu yy')((\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}) : R_1 \sqcap R_2 \sqcap R_3}}$$

However, we *can* type this process in HGV:

$$\frac{\frac{\frac{\Gamma_1, x : S \vdash \mathcal{C} : R_1 \quad \Gamma_2, y : S' \vdash \mathcal{D} : R_2}{\Gamma_1, x : S \parallel \Gamma_2, y : S' \vdash \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2} \quad \Gamma_3, x' : \bar{S}, y' : \bar{S}' \vdash \mathcal{E} : R_3}{\Gamma_1, x : S \parallel \Gamma_2, y : S' \parallel \Gamma_3, x' : \bar{S}, y' : \bar{S}' \vdash (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E} : R_1 \sqcap R_2 \sqcap R_3}}{\Gamma_1, x : S \parallel \Gamma_2, \Gamma_3, x' : \bar{S} \vdash (\nu yy')((\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}) : R_1 \sqcap R_2 \sqcap R_3}}{\Gamma_1, \Gamma_2, \Gamma_3 \vdash (\nu xx')(\nu yy')((\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}) : R_1 \sqcap R_2 \sqcap R_3}$$

Note in particular the shaded hyper-environment, which includes hyper-environment separators to separate endpoints x and x' , as well as y and y' . It follows that, unlike in GV, *both* channels can be split over the same parallel composition. Similarly, the hyper-environment separator allows $\mathcal{C}$ and $\mathcal{D}$ to be composed *without* sharing any channels.

Although HGV types more processes, every well-typed HGV configuration typeable under a singleton hyper-environment Γ is *equivalent* to a well-typed GV configuration, which we show using tree canonical forms.

Proposition 4.5. *Suppose $\Gamma \vdash \mathcal{C} : R$ where $\mathcal{C}$ is in tree canonical form. Then, $\Gamma \vdash_{\text{GV}} \mathcal{C} : R$.*

Proof. By induction on the derivation of $\Gamma \vdash \mathcal{C} : R$, making use of an inductive definition of tree canonical forms. See Appendix B for details. $\square$

Remark 4.6. It is not the case that every HGV configuration typeable under an *arbitrary* hyper-environment $\mathcal{H}$ is equivalent to a well-typed GV configuration. This is because open HGV configurations can form *forest* process structures, whereas (even open) GV configurations must form a *tree* process structure.

Since we can write all well-typed HGV configurations in canonical form, and HGV tree canonical forms are typeable in GV, it follows that every well-typed HGV configuration typeable under a single type environment is equivalent to a well-typed GV configuration.

Typing rules for processes

$P \vdash \mathcal{G}$				
TP-LINK	TP-NEW	TP-PAR	TP-HALT	TP-CLOSE
$\frac{}{x \leftrightarrow^A y \vdash x : A, y : A^\perp}$	$\frac{P \vdash \mathcal{G} \parallel \Gamma, x : A \parallel \Delta, y : A^\perp}{(\nu xy)P \vdash \mathcal{G} \parallel \Gamma, \Delta}$	$\frac{P \vdash \mathcal{G} \quad Q \vdash \mathcal{H}}{P \parallel Q \vdash \mathcal{G} \parallel \mathcal{H}}$	$\frac{}{\mathbf{0} \vdash \emptyset}$	$\frac{P \vdash \emptyset}{x[] \cdot P \vdash x : \mathbf{1}}$
TP-WAIT	TP-SEND	TP-RECV	TP-OFFER-ABSURD	
$\frac{P \vdash \Gamma}{x().P \vdash \Gamma, x : \perp}$	$\frac{P \vdash \Gamma, y : A \parallel \Delta, x : B}{x[y].P \vdash \Gamma, \Delta, x : A \otimes B}$	$\frac{P \vdash \Gamma, y : A, x : B}{x(y).P \vdash \Gamma, x : A \wp B}$		
TP-SELECT-INL	TP-SELECT-INR	TP-OFFER		
$\frac{P \vdash \Gamma, x : A}{x \triangleleft \text{inl}.P \vdash \Gamma, x : A \oplus B}$	$\frac{P \vdash \Gamma, x : B}{x \triangleleft \text{inr}.P \vdash \Gamma, x : A \oplus B}$	$\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B}{x \triangleright \{\text{inl} : P; \text{inr} : Q\} \vdash \Gamma, x : A \& B}$		

Duality

$A^\perp$				
$(A \otimes B)^\perp = A^\perp \wp B^\perp$	$(\mathbf{1})^\perp = \perp$	$(A \oplus B)^\perp = A^\perp \& B^\perp$	$(\mathbf{0})^\perp = \top$	
$(A \wp B)^\perp = A^\perp \otimes B^\perp$	$(\perp)^\perp = \mathbf{1}$	$(A \& B)^\perp = A^\perp \oplus B^\perp$	$(\top)^\perp = \mathbf{0}$	

FIGURE 6. HCP, duality and typing rules for processes.

Corollary 4.7. *If $\Gamma \vdash \mathcal{C} : R$, then there exists some $\mathcal{D}$ such that $\mathcal{C} \equiv \mathcal{D}$ and $\Gamma \vdash_{\text{GV}} \mathcal{D} : R$.*

5. RELATION BETWEEN HGV AND HCP

In this section, we explore two translations, from HGV to HCP and from HCP to HGV, together with their operational correspondence results.

Hypersequent CP. HCP [MP18, KMP19b] is a session-typed process calculus with a correspondence to CLL, which exploits hypersequents to fix extensibility and modularity issues with CP.

Types (A, B) consist of the connectives of linear logic: the multiplicative operators $(\otimes, \wp)$ and units $(\mathbf{1}, \perp)$ and the additive operators $(\oplus, \&)$ and units $(\mathbf{0}, \top)$.

$$A, B ::= \mathbf{1} \mid \perp \mid \mathbf{0} \mid \top \mid A \otimes B \mid A \wp B \mid A \oplus B \mid A \& B$$

Type environments (Γ, Δ) associate names with types. Hyper-environments $(\mathcal{G}, \mathcal{H})$ are collections of type environments. The empty type environment and hyper-environment are written $\cdot$ and $\emptyset$, respectively. Names in type and hyper-environments must be unique and environments may be combined, written Γ, Δ and $\mathcal{G} \parallel \mathcal{H}$, only if they are disjoint.

Processes (P, Q) are a variant of the π -calculus with forwarding [San96, Bor98], bound output [San96], and double binders [Vas12]. The syntax of processes is given by the typing rules (Figure 6), which are standard for HCP [MP18, KMP19b]: $x \leftrightarrow^A y$ forwards messages between x and y ; $(\nu xy)P$ creates a channel with endpoints x and y , and continues as P ; $P \parallel Q$ composes P and Q in parallel; $\mathbf{0}$ is the terminated process; $x[y].P$ creates a new channel, outputs one endpoint over x , binds the other to y , and continues as P ; $x(y).P$ receives a channel endpoint, binds it to y , and continues as P ; $x[] \cdot P$ and $x().P$ close x and continue as P ; $x \triangleleft \text{inl}.P$ and $x \triangleleft \text{inr}.P$ make a binary choice; $x \triangleright \{\text{inl} : P; \text{inr} : Q\}$ offers a binary choice; and $x \triangleright \{\}$ offers a nullary choice. As HCP is synchronous, the only difference between $x[y].P$ and $x(y).P$ is their typing (and similarly for $x[] \cdot P$ and $x().P$). We write *unbound* send as $x(y).P$ (short for $x[z].(y \leftrightarrow z \parallel P)$), and synchronisation as $\bar{x}.P$ (short for

Action rules

$$\text{ACT-PREF} \quad \pi.P \xrightarrow{\pi} P$$

$$\text{ACT-LINK}_1 \quad x \leftrightarrow y \xrightarrow{x \leftrightarrow y} \mathbf{0}$$

$$\text{ACT-LINK}_2 \quad x \leftrightarrow y \xrightarrow{y \leftrightarrow x} \mathbf{0}$$

$$\text{ACT-OFF-INL} \quad x \triangleright \{\text{inl} : P; \text{inr} : Q\} \xrightarrow{x \triangleright \text{inl}} P$$

$$\text{ACT-OFF-INR} \quad x \triangleright \{\text{inl} : P; \text{inr} : Q\} \xrightarrow{x \triangleright \text{inr}} Q$$

Communication Rules

$$\text{ALP-LINK} \quad \frac{P \xrightarrow{x \leftrightarrow z} P'}{(\nu xy)P \xrightarrow{\alpha} P'\{z/y\}}$$

$$\text{BET-SEND} \quad \frac{P \xrightarrow{x[x']\|y(y')} P'}{(\nu xy)P \xrightarrow{\beta} (\nu xy)(\nu x'y')P'}$$

$$\text{BET-CLOSE} \quad \frac{P \xrightarrow{x[]\|y()} P'}{(\nu xy)P \xrightarrow{\beta} P'}$$

$$\text{BET-INL} \quad \frac{P \xrightarrow{x \triangleleft \text{inl} \| y \triangleright \text{inl}} P'}{(\nu xy)P \xrightarrow{\beta} (\nu xy)P'}$$

$$\text{BET-INR} \quad \frac{P \xrightarrow{x \triangleleft \text{inr} \| y \triangleright \text{inr}} P'}{(\nu xy)P \xrightarrow{\beta} (\nu xy)P'}$$

Structural Rules

$$\text{STR-RES} \quad \frac{P \xrightarrow{\ell} P' \quad x, y \notin \text{cn}(\ell)}{(\nu xy)P \xrightarrow{\ell} (\nu xy)P'}$$

$$\text{STR-PAR}_1 \quad \frac{P \xrightarrow{\ell} P' \quad \text{bn}(\ell) \cap \text{fn}(Q) = \emptyset}{P \parallel Q \xrightarrow{\ell} P' \parallel Q}$$

$$\text{STR-PAR}_2 \quad \frac{Q \xrightarrow{\ell} Q' \quad \text{bn}(\ell) \cap \text{fn}(P) = \emptyset}{P \parallel Q \xrightarrow{\ell} P \parallel Q'}$$

$$\text{STR-SYN} \quad \frac{P \xrightarrow{\ell} P' \quad Q \xrightarrow{\ell'} Q' \quad \text{bn}(\ell) \cap \text{bn}(\ell') = \emptyset}{P \parallel Q \xrightarrow{\ell \parallel \ell'} P' \parallel Q'}$$

FIGURE 7. HCP, label transition semantics.

$x[z].(z[].\mathbf{0} \parallel P)$ and $x.P$ (short for $x(z).z().P$). Duality is standard and is involutive, *i.e.*, $(A^\perp)^\perp = A$.

We define a standard structural congruence ($\equiv$) similar to that of HGV, *i.e.*, parallel composition is commutative and associative, we can commute name restrictions, swap the order of endpoints, swap links, and have scope extrusion (similar to Figure 4). Note that since we base our formal developments on an LTS semantics, structural congruence is not required for reduction.

$$\begin{aligned} x \leftrightarrow^A y &\equiv y \leftrightarrow^{A^\perp} x & P \parallel \mathbf{0} &\equiv P & P \parallel Q &\equiv Q \parallel P & P \parallel (Q \parallel R) &\equiv (P \parallel Q) \parallel R \\ (\nu xx')(\nu yy')P &\equiv (\nu yy')(\nu xx')P & (\nu xy)P &\equiv (\nu yx)P \\ (\nu xy)(P \parallel Q) &\equiv P \parallel (\nu xy)Q & \text{if } x, y &\notin \text{fv}(P) \end{aligned}$$

We define the labelled transition system for HCP as a small refinement of the LTS for the additive-multiplicative fragment of the π LL calculus introduced by Montesi and Peressotti [MP21], in turn inspired by their previous system *CT* [MP18]. The LTS is identical, save for the fact that we distinguish two types of internal actions. Action labels l represent the actions that a process can fire. Prefixes π are a convenient subset of action labels which can be written as prefixes to processes, *i.e.*, $\pi.P$. Transition labels ℓ include

action labels and the parallel composition of two action labels, along with *internal* actions α , β , and τ . The LTS gives rise to two types of internal action: α represents only the evaluation of links as *renaming*, and β represents only *communication*. Labels τ arise only due to saturated transition (Definition 5.4) and are not produced by the rules in the LTS.

Prefixes	$\pi ::= x[y] \mid x[] \mid x(y) \mid x() \mid x \triangleleft \text{inl} \mid x \triangleleft \text{inr}$
Action Labels	$l ::= \pi \mid x \leftrightarrow y \mid x \triangleright \text{inl} \mid x \triangleright \text{inr}$
Transition Labels	$\ell ::= l \mid l \parallel l' \mid \alpha \mid \beta$

We let ℓ_x range over labels on x : $x \leftrightarrow y$, $x[y]$, $x[]$, *etc.* Labelled transition $\xrightarrow{\ell}$ is defined in Figure 7. We write $\xrightarrow{\ell} \xrightarrow{\ell'}$ for the composition of $\xrightarrow{\ell}$ and $\xrightarrow{\ell'}$, $\xrightarrow{\ell^+}$ for the transitive closure of $\xrightarrow{\ell}$, and $\xrightarrow{\ell^*}$ for the reflexive-transitive closure of $\xrightarrow{\ell}$. We write $\text{bn}(\ell)$ and $\text{fn}(\ell)$ for the bound and free names contained in ℓ , respectively. We write $\text{cn}(\ell)$ for all names in ℓ , *i.e.*, $\text{cn}(\ell) = \text{fn}(\ell) \cup \text{bn}(\ell)$.

Metatheory. Transitions preserve typeability. Since internal actions occur only under binders, they are typable under the same hyper-environment.

Theorem 5.1 (Type Preservation). *Suppose $P \vdash \mathcal{G}$ and $P \xrightarrow{\ell} Q$.*

- *If ℓ is internal, then $Q \vdash \mathcal{G}$.*
- *If ℓ is not internal, then there exists some $\mathcal{H}$ such that $Q \vdash \mathcal{H}$.*

Proof. Following the approach of [KMP19a, MP18, MP21], type preservation is established by defining proof transformations on typing derivations of each reducing process. The only difference with respect to [MP18, MP21] arises due to our separate treatment of α and β actions, which does not materially impact the proof. $\square$

Similarly, our LTS for HCP satisfies progress. Following [KMP19a, MP21], the key intermediate step is to note that for every type environment in a hyper-environment, there is some free name which can be acted upon. Again, the stratification of internal actions does not materially impact the proof.

Theorem 5.2 (Progress). *If $P \vdash \mathcal{H}$ and $P \neq \mathbf{0}$, then there exist some ℓ, Q such that $P \xrightarrow{\ell} Q$.*

Behavioural Theory. The behavioural theory for HCP follows Kokke *et al.* [KMP19a], except that we distinguish two subrelations of weak bisimilarity, following the subtypes of internal actions.

Definition 5.3 (Strong bisimulation and strong bisimilarity). A symmetric relation $\mathcal{R}$ on processes is a *strong bisimulation* if $P \mathcal{R} Q$ implies that if $P \xrightarrow{\ell} P'$, then $Q \xrightarrow{\ell} Q'$ for some Q' such that $P' \mathcal{R} Q'$. *Strong bisimilarity* is the largest relation $\sim$ that is a strong bisimulation.

Definition 5.4 (Saturated transition). The $\mathcal{L}$ -saturated transition relation, for $\mathcal{L} \subseteq \{\alpha, \beta\}$, is the smallest relation $\Rightarrow_{\mathcal{L}}$ closed under the following rules, with saturated transition labels ℓ ranging over transition labels and the distinguished label τ :

$$\frac{}{P \Rightarrow_{\mathcal{L}} P} \quad \frac{P \xRightarrow{\tau}_{\mathcal{L}} P' \xrightarrow{\ell} Q' \xRightarrow{\tau}_{\mathcal{L}} Q \quad \ell \in \mathcal{L}}{P \xRightarrow{\tau}_{\mathcal{L}} Q} \quad \frac{P \xRightarrow{\tau}_{\mathcal{L}} P' \xrightarrow{\ell} Q' \xRightarrow{\tau}_{\mathcal{L}} Q \quad \ell \notin \mathcal{L}}{P \xRightarrow{\ell}_{\mathcal{L}} Q}$$

We write $\Rightarrow_{\ell}$ as shorthand for $\Rightarrow_{\{\ell\}}$, and we write $\Rightarrow$ as shorthand for $\Rightarrow_{\{\alpha, \beta\}}$.

Definition 5.5 (Weak bisimulation and weak bisimilarity). A symmetric relation $\mathcal{R}$ on processes is an $\mathcal{L}$ -bisimulation, for $\mathcal{L} \subseteq \{\alpha, \beta\}$, if $P \mathcal{R} Q$ implies that if $P \xRightarrow{\ell'}_{\mathcal{L}} P'$, then $Q \xRightarrow{\ell'}_{\mathcal{L}} Q'$ for some Q' such that $P' \mathcal{R} Q'$. The $\mathcal{L}$ -bisimilarity relation is the largest relation $\approx_{\mathcal{L}}$ that is an $\mathcal{L}$ -bisimulation. We write $\approx$ as shorthand for $\approx_{\{\alpha, \beta\}}$.

Lemma 5.6. *Structural congruence, strong bisimilarity and the various forms of weak bisimilarity are related as follows:*

$$\equiv \subset \sim \quad \sim \subset \approx \quad \sim \subset \approx_{\alpha} \quad \sim \subset \approx_{\beta}$$

Differences with previous version. The LTS in Figure 7 is similar to that in the previous version of this work [FKD⁺21], with the exception that we have removed the rules TAU-ALP and TAU-BET:

$$\begin{array}{c} \text{TAU-ALP} \\ \hline \frac{P \xrightarrow{\alpha} P'}{P \xrightarrow{\tau} P'} \end{array} \quad \begin{array}{c} \text{TAU-BET} \\ \hline \frac{P \xrightarrow{\beta} P'}{P \xrightarrow{\tau} P'} \end{array}$$

To see why these rules are problematic, consider processes $P = (\nu xy)(z \leftrightarrow x \parallel y[].\mathbf{0})$ and $Q = z[].\mathbf{0}$. Following Definition 5.5, P and Q are α -bisimilar, as P only has the α -transition $P \xrightarrow{\alpha} Q$ and Q has no transitions. In the previous version, TAU-ALP gave P the derived τ -transition $P \xrightarrow{\tau} Q$, which meant that $P \not\approx_{\alpha} Q$, as $Q \not\xrightarrow{\tau} Q$. Therefore TAU-ALP collapses $\approx_{\alpha}$ to $\sim$ and TAU-BET collapses $\approx_{\beta}$ to $\sim$.

The solution we adopted was to remove TAU-ALP and TAU-BET from the label transition relation $\longrightarrow$, and instead lift α - and β -transitions to τ -transitions in the definition of saturated transition¹.

Translating HGV to HCP. We factor the translation from HGV to HCP into two translations: (1) a translation into HGV*, a fine-grain call-by-value [LPT03] variant of HGV, which makes control flow explicit; and (2) a translation from HGV* to HCP. In so doing, we can concentrate on the essence of the translations as opposed to concerning ourselves with administrative reductions.

HGV*. We define HGV* as a refinement of HGV in which any non-trivial term must be named by a let-binding before being used. While **let** is syntactic sugar in HGV, it is part of the core language in HGV*. Correspondingly, the reduction rule for **let** follows from the encoding in HGV, i.e., $\mathbf{let } x = V \mathbf{ in } M \longrightarrow_{\mathbf{M}} M\{V/x\}$.

Terms	$L, M, N ::= V \mid \mathbf{let } x = M \mathbf{ in } N \mid V W$
	$\mid \mathbf{let } () = V \mathbf{ in } M \mid \mathbf{let } (x, y) = V \mathbf{ in } M$
	$\mid \mathbf{absurd } V \mid \mathbf{case } V \{ \mathbf{inl } x \mapsto M; \mathbf{inr } y \mapsto N \}$
Values	$V, W ::= x \mid K \mid \lambda x. M \mid () \mid (V, W) \mid \mathbf{inl } V \mid \mathbf{inr } V$
Evaluation contexts	$E ::= \square \mid \mathbf{let } x = E \mathbf{ in } M$
Thread contexts	$\mathcal{F} ::= \phi E$

¹We thank Marco Peressotti for notifying us of the error and suggesting the fix.

Remark 5.7. Fine-grain call-by-value λ -calculi typically include an explicit **return** V construct to embed values into the term language. As there is no difference between the shapes of the value and term typing judgements, we allow ourselves to embed values directly for simplicity.

We can *naïvely* translate HGV to HGV* ($\langle \cdot \rangle$) by let-binding each subterm in a value position, *e.g.*, $\langle \text{inl } M \rangle = \text{let } z = \langle M \rangle \text{ in inl } z$.

Definition 5.8 (Naïve translation of HGV to HGV*).

$\langle x \rangle$	$= x$
$\langle \lambda x. M \rangle$	$= \lambda x. \langle M \rangle$
$\langle L \ M \rangle$	$= \text{let } x = \langle L \rangle \text{ in let } y = \langle M \rangle \text{ in } x \ y$
$\langle () \rangle$	$= ()$
$\langle \text{let } () = L \text{ in } M \rangle$	$= \text{let } z = \langle L \rangle \text{ in let } () = z \text{ in } \langle M \rangle$
$\langle (M, N) \rangle$	$= \text{let } x = \langle M \rangle \text{ in let } y = \langle N \rangle \text{ in } (x, y)$
$\langle \text{let } (x, y) = L \text{ in } M \rangle$	$= \text{let } z = \langle L \rangle \text{ in let } (x, y) = z \text{ in } \langle M \rangle$
$\langle \text{inl } M \rangle$	$= \text{let } z = \langle M \rangle \text{ in inl } z$
$\langle \text{inr } M \rangle$	$= \text{let } z = \langle M \rangle \text{ in inr } z$
$\langle \text{case } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} \rangle$	$= \text{let } z = \langle L \rangle \text{ in case } z \{ \text{inl } x \mapsto \langle M \rangle; \text{inr } y \mapsto \langle N \rangle \}$
$\langle \text{absurd } L \rangle$	$= \text{let } z = \langle L \rangle \text{ in absurd } z$

Standard techniques can be used to avoid administrative redexes [Plø75, DMN07]. We give a full definition of HGV* in Appendix C.

HGV* to HCP. The translation from HGV* to HCP is given in Figure 8. All control flow is encapsulated in values and let-bindings. We define a pair of translations on types, $\llbracket \cdot \rrbracket$ and $\llbracket \cdot \rrbracket^\perp$, such that $\llbracket T \rrbracket = \llbracket T \rrbracket^\perp$. We extend these translations pointwise to type environments and hyper-environments. We define translations on configurations ($\llbracket \cdot \rrbracket_r^c$), terms ($\llbracket \cdot \rrbracket_r^m$) and values ($\llbracket \cdot \rrbracket_r^v$), where r is a fresh name denoting a distinguished output channel.

We translate an HGV sequent $\mathcal{G} \parallel \Gamma \vdash C : T$ as $\llbracket C \rrbracket_r^c \vdash \llbracket \mathcal{G} \rrbracket \parallel \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket^\perp$, where Γ is the type environment corresponding to the main thread. The translation of computations includes synchronisation action in order to faithfully simulate a call-by-value reduction strategy. The (term) translation of a value $\llbracket V \rrbracket_r^m$ immediately pings the output channel r to announce that it is a value. The translation of a let-binding $\llbracket \text{let } w = M \text{ in } N \rrbracket_r^m$ first evaluates M to a value, which then pings the internal channel x/x' and unblocks the continuation $x. \llbracket N \rrbracket_r^m$. The translations of main and child threads each make use of an internal result channel. The translation of a child thread consumes the yielded unit endpoint once the child thread has terminated. The translation of the main thread forwards the result value along the external output channel once the main thread has terminated.

There are two changes with respect to the translation of our earlier paper [FKD⁺21]. First, in the earlier work the translation of the main thread output directly to the external output channel instead of forwarding via an intermediary as in the current translation. This change is purely aesthetic. Second, in the earlier work the translation of **fork** was not sufficiently concurrent. Correspondingly there was an error in the case of the operational correspondence proof which is fixed in the current paper.

Lemma 5.9 (Type Preservation).

- (1) If $\Gamma \vdash V : T$, then $\llbracket V \rrbracket_r^v \vdash \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket^\perp$.
- (2) If $\Gamma \vdash M : T$, then $\llbracket M \rrbracket_r^m \vdash \llbracket \Gamma \rrbracket, r : \mathbf{1} \otimes \llbracket T \rrbracket^\perp$.

Translation on types

 $\llbracket T \rrbracket$ and $\llbracket T \rrbracket^\vee$

$$\begin{aligned}
\llbracket !T.S \rrbracket &= \llbracket T \rrbracket^\perp \otimes \llbracket S \rrbracket & \llbracket \text{end}_! \rrbracket &= \mathbf{1} & \llbracket T \rrbracket &= \llbracket T \rrbracket^\perp, \\
\llbracket ?T.S \rrbracket &= \llbracket T \rrbracket^\perp \wp \llbracket S \rrbracket & \llbracket \text{end}_? \rrbracket &= \perp & & \text{if } T \text{ is not a session type} \\
\llbracket T \times U \rrbracket &= \llbracket T \rrbracket \otimes \llbracket U \rrbracket & \llbracket \mathbf{1} \rrbracket &= \mathbf{1} & \llbracket T \multimap U \rrbracket &= \llbracket T \rrbracket^\perp \wp (\mathbf{1} \otimes \llbracket U \rrbracket) \\
\llbracket T + U \rrbracket &= \llbracket T \rrbracket \oplus \llbracket U \rrbracket & \llbracket \mathbf{0} \rrbracket &= \mathbf{0} & \llbracket S \rrbracket &= \llbracket S \rrbracket^\perp
\end{aligned}$$

Translation on configurations, terms, and values

 $\llbracket \mathcal{C} \rrbracket_r^c$, $\llbracket M \rrbracket_r^m$, and $\llbracket V \rrbracket_r^v$

$$\begin{aligned}
\llbracket \circ M \rrbracket_r^c &= (\nu y y')(\llbracket M \rrbracket_{y'}^m \parallel y'.y'().\mathbf{0}) & \llbracket (\nu x x')\mathcal{C} \rrbracket_r^c &= (\nu x x')\llbracket \mathcal{C} \rrbracket_r^c & \llbracket x \xrightarrow{z} y \rrbracket_r^c &= \bar{z}.z().x \leftrightarrow y \\
\llbracket \bullet M \rrbracket_r^c &= (\nu y y')(\llbracket M \rrbracket_{y'}^m \parallel y'.y' \leftrightarrow r) & \llbracket \mathcal{C} \parallel \mathcal{D} \rrbracket_r^c &= \llbracket \mathcal{C} \rrbracket_r^c \parallel \llbracket \mathcal{D} \rrbracket_r^c \\
\llbracket x \rrbracket_r^v &= r \leftrightarrow x & \llbracket () \rrbracket_r^v &= r().\mathbf{0} & \llbracket \text{inl } V \rrbracket_r^v &= r \triangleleft \text{inl}.\llbracket V \rrbracket_r^v \\
\llbracket \lambda x.M \rrbracket_r^v &= r(x).\llbracket M \rrbracket_r^m & \llbracket (V, W) \rrbracket_r^v &= r[x].(\llbracket V \rrbracket_x^v \parallel \llbracket W \rrbracket_r^v) & \llbracket \text{inr } V \rrbracket_r^v &= r \triangleleft \text{inr}.\llbracket V \rrbracket_r^v \\
\llbracket V \ W \rrbracket_r^m &= (\nu x x')(\nu y y')(y\langle x \rangle.r \leftrightarrow y \parallel \llbracket V \rrbracket_{y'}^v \parallel \llbracket W \rrbracket_{x'}^v) \\
\llbracket \text{let } () = V \text{ in } M \rrbracket_r^m &= (\nu x x')(x().\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{x'}^v) \\
\llbracket \text{let } (x, y) = V \text{ in } M \rrbracket_r^m &= (\nu y y')(y(x).\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{y'}^v) \\
\llbracket \text{case } V \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} \rrbracket_r^m &= (\nu x x')(x \triangleright \{ \text{inl} : \llbracket M \rrbracket_r^m; \text{inr} : \llbracket N \{x/y\} \rrbracket_r^m \} \parallel \llbracket V \rrbracket_{x'}^v) \\
\llbracket \text{absurd } V \rrbracket_r^m &= (\nu x x')(x \triangleright \{ \} \parallel \llbracket V \rrbracket_{x'}^v) \\
\llbracket \text{let } x = M \text{ in } N \rrbracket_r^m &= (\nu x x')(x.\llbracket N \rrbracket_r^m \parallel \llbracket M \rrbracket_{x'}^m) \\
\llbracket V \rrbracket_r^m &= \bar{r}.\llbracket V \rrbracket_r^v \\
\llbracket \text{link} \rrbracket_r^v &= r(y).y(x).\bar{r}.r().x \leftrightarrow y \\
\llbracket \text{fork} \rrbracket_r^v &= (\nu y y')(r(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \parallel y'(x).x\langle y' \rangle.x.x().\mathbf{0}) \\
\llbracket \text{send} \rrbracket_r^v &= r(y).y(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \\
\llbracket \text{recv} \rrbracket_r^v &= r(x).x(y).\bar{r}.r\langle y \rangle.r \leftrightarrow x \\
\llbracket \text{wait} \rrbracket_r^v &= r(x).x().\bar{r}.r().\mathbf{0}
\end{aligned}$$

FIGURE 8. Translation from HGV* to HCP.

(3) If $\mathcal{G} \parallel \Gamma \vdash \mathcal{C} : T$, where Γ is the type environment for the main thread in $\mathcal{C}$, then $\llbracket \mathcal{C} \rrbracket_r^c \vdash \llbracket \mathcal{G} \rrbracket \parallel \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket^\perp$.

Lemma 5.10 (Substitution). *If M is a well-typed term with $w \in \text{fv}(M)$, and V is a well-typed value, then $(\nu w w')(\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{w'}^v) \approx_\alpha \llbracket M\{V/w\} \rrbracket_r^m$.*

Theorem 5.11 (Operational Correspondence). *Suppose $\mathcal{C}$ is a well-typed configuration.*

- (1) (Preservation of reductions) *If $\mathcal{C} \longrightarrow \mathcal{C}'$, then there exists a P such that $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\beta^+}_\alpha P$ and $P \approx_\alpha \llbracket \mathcal{C}' \rrbracket_r^c$; and*
- (2) (Reflection of transitions)
 - *if $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\alpha} P$, then $P \approx_\alpha \llbracket \mathcal{C} \rrbracket_r^c$; and*
 - *if $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\beta} P$, then there exists a $\mathcal{C}'$ and a P' such that $\mathcal{C} \longrightarrow \mathcal{C}'$ and $P \xrightarrow{\beta^*}_\alpha P'$ and $P' \approx_\alpha \llbracket \mathcal{C}' \rrbracket_r^c$. Furthermore, $\mathcal{C}'$ is unique up to structural congruence.*

The proof is in Appendix C. One might strive for a tighter operational correspondence here, but our current translation generates multiple administrative β -transitions. The only term reduction that translates to multiple β -transitions is the one for let-bindings. This is because we choose to encode synchronisation using two β -transitions. We could adjust the accounting here by treating synchronisation as a single β -transition or its own special kind of

administrative transition. Many more administrative reductions arise from the configuration translation. These are due to a combination of synchronisations and also the fact that we use constants along with pairs and application for our communication primitives instead of building-in fully-applied communication primitives.

Translating HCP to HGV. We cannot translate HCP processes to HGV terms directly: HGV's term language only supports **fork** (see Section 7 for further discussion), so there is no way to translate an individual name restriction or parallel composition. However, we can still translate HCP into HGV via the composition of known translations.

HCP into CP: We must first reunite each parallel composition with its corresponding name restriction, *i.e.*, translate to CP using the *disentanglement* translation shown by Kokke *et al.* [KMP19b, Lemma 4.7]. The result is a collection of independent CP processes.

CP into GV: Next, we can translate each CP process into a GV configuration using (a variant of) Lindley and Morris' translation [LM15, Figure 8].

GV into HGV: Finally, we can use our embedding of GV into HGV (Theorem 4.3) to obtain a collection of well-typed HGV configurations, which can be composed using TC-PAR to result in a single well-typed HGV configuration.

The translation from HCP into CP and the embedding of GV into HGV preserve and reflect reduction. However, as previously mentioned, Lindley and Morris's original translation from CP to GV preserves but does *not* reflect reduction due to an asynchronous encoding of choice. By adapting their translation to use a synchronous encoding of choice (Section 3), we obtain a translation from CP to GV that both preserves and reflects reduction. Thus, composing all three translations together we obtain a translation from HCP to HGV that preserves and reflects reduction.

6. EXTENSIONS

In this section, we outline three extensions to HGV that exploit generalising the tree structure of processes to a forest structure. These extensions are of particular interest since HGV already supports a core aspect of forest structure, enabling its full utilisation merely through the addition of a structural rule. In contrast, to extend GV with forest structure one must distinguish two distinct introduction rules for parallel composition [LM15, Fow19]. Other extensions to GV such as shared channels [LM15], polymorphism [LM17], and recursive session types [LM16] adapt to HGV almost unchanged.

From trees to forests. The TC-PAR rule allows two processes to be composed in parallel if they are typeable under separate hyper-environments. In a closed program, hyper-environment separators are introduced by TC-RES, meaning that each process must be connected by a channel.

The following TC-MIX rule allows two type environments Γ_1, Γ_2 to be split by a hyper-environment separator *without* a channel connecting them, and is inspired by Girard's Mix rule [Gir87]; in the concurrent setting, Mix can be interpreted as concurrency *without* communication [LM15, ALM16]. TC-Mix admits a much simpler treatment of **link** and provides a crucial ingredient for handling exceptional behaviour.

$$\text{TC-Mix} \quad \frac{\mathcal{G} \parallel \Gamma_1 \parallel \Gamma_2 \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma_1, \Gamma_2 \vdash \mathcal{C} : R}$$

Atkey *et al.* [ALM16] show that conflating the $\mathbf{1}$ and $\perp$ types in CP (which correspond respectively to the **end**_! and **end**_? types in GV) is logically equivalent to adding the MIX rule and a 0-MIX rule (used to type an empty process). It follows that in the presence of TC-Mix, we use self-dual **end** type; in the GV setting, by using a self-dual **end** type, we decouple closing a channel from process termination. We therefore refine the TC-CHILD rule and the type schema for **fork** to ensure that each child thread returns the unit value, and replace the **wait** constant with a **close** constant which eliminates an endpoint of type **end**.

$$\text{fork} : (S \multimap \mathbf{1}) \multimap \bar{S} \quad \text{close} : \text{end} \multimap \mathbf{1} \quad \text{TC-CHILD} \quad \frac{\Gamma \vdash M : \mathbf{1}}{\Gamma \vdash \circ M : \circ}$$

$$\text{E-CLOSE} \quad (\nu xy)(E[\text{close } x] \parallel E'[\text{close } y]) \longrightarrow E[()] \parallel E'[()]$$

Given TC-Mix, we might expect a term-level construct **spawn** : $(\mathbf{1} \multimap \mathbf{1}) \multimap \mathbf{1}$ which spawns a parallel thread without a connecting channel. We can encode such a construct using **fork** and **close** (assuming fresh x and y):

$$\text{spawn } M \triangleq \text{let } x = \text{fork}(\lambda y. \text{close } y; M) \text{ in } \text{close } x$$

Assuming the encoded **spawn** is running in a main thread, after two reduction steps, we are left with the configuration:

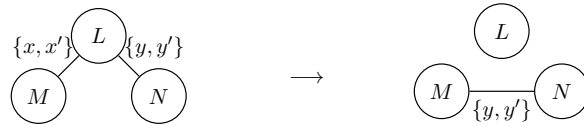
$$\frac{\frac{\frac{\cdot \vdash M : \mathbf{1}}{\cdot \vdash \circ M : \circ} \text{TC-CHILD} \quad \frac{\cdot \vdash M : \mathbf{1}}{\cdot \vdash \bullet() : \mathbf{1}} \text{TC-MAIN}}{\cdot \parallel \cdot \vdash \circ M \parallel \bullet() : \mathbf{1}} \text{TC-PAR}}{\cdot \vdash \circ M \parallel \bullet() : \mathbf{1}} \text{TC-Mix}$$

Note the essential use of TC-Mix to insert a hyper-environment separator.

The addition of TC-Mix does not affect preservation or progress. The result follows from routine adaptations of the proof of Theorem 3.3 and Theorem 3.20.

By relaxing the tree process structure restriction using TC-Mix, we can obtain a more efficient treatment of **link**, and can support the treatment of exceptions advocated by Fowler *et al.* [FLMD19].

A simpler link. The **link** (x, y) construct forwards messages from x to y and vice-versa. Consider threads $L = F[\text{link } (x, y)]$, M , N , where L connects to M by x and to N by y .



The result of link reduction has forest structure. Well-typed closed programs in both GV and HGV must *always* maintain tree structure. Different versions of GV do so in various unsatisfactory ways: one is pre-emptive blocking [LM15], which breaks confluence; another is two-stage linking (Figure 4), which defers forwarding via a special link thread [LM16].

Lindley and Morris [LM15] implement **link** using the following rule (modified here to use a double-binder formulation):

$$(\nu xx')(F[\mathbf{link}(x, y)] \parallel \mathcal{F}'[M]) \longrightarrow (\nu xx')(F[x] \parallel \mathcal{F}'[\mathbf{wait} x'; M\{y/x'\}]) \quad \text{where } x' \in \text{fv}(M)$$

The first thread will eventually reduce to $\circ x$, at which point the second thread will synchronise to eliminate x and x' and then evaluate the continuation M with endpoint y substituted for x' . Unfortunately, this formulation of **link** preemptively inhibits reduction in the second thread, since the evaluation rule inserts a blocking **wait**. The resulting system does not satisfy the diamond property.

HGV uses the incarnation of **link** advocated by Lindley and Morris [LM16], where linking is split into two stages: the first generates a fresh pair of endpoints z, z' and a link thread of the form $x \xrightarrow{z'} y$, and returns z to the calling thread. Once the calling thread has evaluated to a value (which must by typing be z), then the link substitution can take place. This formulation recovers confluence, but we still lose a degree of concurrency: communication on y is blocked until the linking thread has fully evaluated. In an ideal implementation, the behaviour of the linking thread would be irrelevant to the remainder of the configuration. The operation requires additional runtime syntax and thus complicates the metatheory.

The above issues are symptomatic of the fact that the process structure after a link takes place is a forest rather than a tree. However, with TC-Mix, we can refine the type schema for **link** to $(S \times \bar{S}) \multimap \mathbf{1}$ and we can use the following rule:

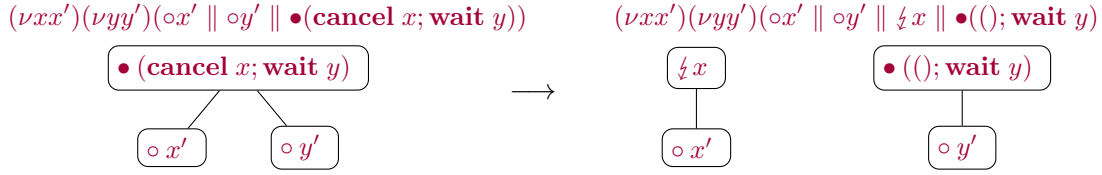
$$(\nu xx')(F[\mathbf{link}(x, y)] \parallel \phi N) \longrightarrow F[()] \parallel \phi N\{y/x'\}$$

This formulation enables immediate substitution, maximimising concurrency. A variant of HGV replacing E-REIFY-LINK and E-COMM-LINK with E-LINK-MIX retains HGV's metatheory.

Exceptions. In order to support exceptions in the presence of linear endpoints [FLMD19, MV18] we must have a way of *cancelling* an endpoint. Mostrous and Vasconcelos [MV18] describe a process calculus allowing the *explicit cancellation* of a channel endpoint, accounting for exceptional scenarios such as a client disconnecting, or a thread encountering an unrecoverable error. Attempting to communicate with a cancelled endpoint raises an exception. Fowler *et al.* [FLMD19] extend these ideas to the functional setting, introducing Exceptional GV (EGV). EGV supports exceptional behaviour by adding three term-level constructs:

- a new constant, **cancel** : $S \multimap \mathbf{1}$, which allows us to discard an arbitrary session endpoint with type S
- a construct **raise**, which raises an exception
- an exception handling construct **try** L **as** x **in** M **otherwise** N in the style of Benton & Kennedy [BK01], which attempts possibly-failing computation L , binding the result to x in success continuation M if successful and evaluating N if an exception is raised

Cancellation generates a *zapper thread* ($\cancel{x}$) which severs a tree topology into a forest as in the following example.



The configuration on the left has a tree process structure. However, after reduction, we obtain the configuration on the right which is clearly a forest and thus needs TC-Mix to be typeable. We have described a *synchronous* version of EGV, but extending our treatment to asynchrony as in the work of [FLMD19] is a routine adaptation.

7. CAN WE SEPARATE FORK?

Hyper-environments allow us to cleanly separate name restriction and parallel composition in process configurations. A natural follow-on question is whether we could use the same technique at the level of *terms* in order to split **fork** into separate constructs for creating a channel and spawning a process. As tantalising a prospect this is, we argue that the disadvantages outweigh the benefits.

Suppose we were to extend term typing to allow hyper-environments, $\mathcal{G} \vdash M : T$, and were to introduce terms **let** $\langle x, x' \rangle = \text{new in } M$ to create a channel and **let** $\langle \rangle = \text{spawn } M \text{ in } N$ to spawn a thread, with the following typing rules:

$$\begin{array}{c} \text{TM-LETNEW} \\ \frac{\mathcal{G} \parallel \Gamma_1, x : S \parallel \Gamma_2, x' : \bar{S} \vdash M : T}{\mathcal{G} \parallel \Gamma_1, \Gamma_2 \vdash \text{let } \langle x, x' \rangle = \text{new in } M : T} \end{array} \quad \begin{array}{c} \text{TM-LETSPAWN} \\ \frac{\mathcal{G} \vdash M : \text{end}_! \quad \mathcal{H} \vdash N : T}{\mathcal{G} \parallel \mathcal{H} \vdash \text{let } \langle \rangle = \text{spawn } M \text{ in } N : T} \end{array}$$

These rather ad-hoc rules mirror hypersequent cut and hypersequent composition: TM-LETNEW creates a new channel with endpoints x and x' , and requires them to be used in separate threads in the continuation M ; and TM-LETSPAWN takes a term M , spawns it as a child thread, and continues as N . Using these rules, we can encode **fork** M as **let** $\langle x, x' \rangle = \text{new in let } \langle \rangle = \text{spawn } (M \ x) \text{ in } x'$.

Where else can we allow hyper-environments? In HCP, we have two options: (1) if we restrict *all logical rules* to singleton hypersequents and allow hyper-environments only in the rules for name restriction and parallel composition, we can use standard sequential semantics [MP18, KMP19b]; but (2) if we allow hyper-environments in *any logical rules*, we must use a semantics which allows the corresponding actions to be delayed [KMP19a]. This is unlikely to be a property of logical rules, but rather due to the fact that the logical rules correspond exactly to the communication actions—which block reduction—and the structural rules to name restriction and parallel composition—which do not. Therefore, we expect the positions where hypersequents can safely occur to follow from the structure of evaluation contexts and whether any blocking term perform a communication action.

Regardless of our choice, we would be left with restrictions on the syntax of terms that seem sensible in a process calculus, but are surprising in a λ -calculus. In the strictest variant, where we disallow hyper-environments in all but the above two rules, uses of TM-LETNEW and TM-LETSPAWN may be interleaved, but no other construct may appear between a

TM-LETNEW and its corresponding TM-LETSPAWN. Consider the following terms, where M uses x and y , and N uses x' . Term (7.1) may be well-typed, but (7.2) is always ill-typed:

$$\text{let } y = 1 \text{ in let } \langle x, x' \rangle = \text{new in let } \langle \rangle = \text{spawn } M \text{ in } N \quad (7.1)$$

$$\text{let } \langle x, x' \rangle = \text{new in let } y = 1 \text{ in let } \langle \rangle = \text{spawn } M \text{ in } N \quad (7.2)$$

Note that $\text{let } \langle x, x' \rangle = \text{new in } M$ is a single, monolithic term constructor—exactly what hypersequents were meant to prevent! However, if we attempt to decompose these constructors, we find that these are not the regular product and unit types.

8. RELATED WORK

Session Types and Functional Languages. Session types were originally introduced in the context of process calculi [Hon93, THK94, HVK98], however they have been vastly integrated also in functional calculi, a line of work initiated by Gay and collaborators [VRG04, VGR06, GV10]. This family of calculi builds session types directly into a lambda calculus. Toninho *et al.* [TCP13] take an alternative approach, stratifying their system into a session-typed process calculus and a separate functional calculus. There are many pragmatic embeddings of session type systems in existing functional programming languages [NT04, PT08, SE08, IYA10, OY16, KD21a]. A detailed survey is given by Orchard and Yoshida [OY17].

Propositions as Sessions. When Girard introduced linear logic [Gir87] he suggested a connection with concurrency. Abramsky [Abr94] and Bellin and Scott [BS94] give embeddings of linear logic proofs in π -calculus, where cut reduction is simulated by π -calculus reduction. Both embeddings interpret tensor as parallel composition. The correspondence with π -calculus is not tight in that these systems allow independent prefixes to be reordered. Caires and Pfenning [CP10] give a propositions as types correspondence between dual intuitionistic linear logic and a session-typed π -calculus called π DILL. They interpret tensor as output. The correspondence with π -calculus is tight in that independent prefixes may not be reordered. With CP [Wad14], Wadler adapts π DILL to classical linear logic. Aschieri and Genco [AG20] give an interpretation of classical multiplicative linear logic as concurrent functional programs. They interpret $\wp$ as parallel composition, and the connection to session types is less direct.

Priority-based Calculi. Systems such as π DILL, CP, and GV (and indeed HCP and HGV) ensure deadlock freedom by exploiting the type system to statically impose a tree structure on the communication topology — there can be at most one communication channel between any two processes. Another line of work explores a more liberal approach to deadlock freedom enabling some cyclic communication topologies, where deadlock freedom is guaranteed via *priorities*, which impose an order on actions. Priorities were introduced by Kobayashi and Padovani [Kob06, Pad14] and adopted by Dardha and Gay [DG18] in Priority CP (PCP), and Kokke and Dardha in Priority GV (PGV) [KD21b]. Dezani *et al.* [DCdY07] and Vieira and Vasconcelos [VV13] use a partial order on channels to guarantee deadlock freedom, following Kobayashi's work [Kob06]. Later on Dezani *et al.* [DCMYD06] guarantee progress by allowing only one active session at a time. Carbone *et al.* [CDM14] use catalysers to show that progress is a compositional form of lock freedom for standard typed π calculus. The authors describe how this technique can be used for session typed

π -calculus by using the encoding of session types to linear types [DGS17, Dar14, Dar16]. Dardha and Perez [DP22] compare the different calculi and techniques for deadlock freedom using CP and CLL as a yardstick and showing that the class of processes in CP is strictly included in the class of processes typed by Kobayashi [Kob06].

Graph-theoretic Approaches. Carbone and Debois [CD10] define a graph-theoretic approach for a session typed π -calculus. They define an explicit dependency graph defined inductively on the structure of a process, in contrast to our approach of inducing a graph on type environments given a co-name set. They ensure progress for processes with acyclic graphs using a *catalyser*, which provides a missing counterpart to a process. Jacobs *et al.* [JBK22a] also define a graph-theoretic approach to deadlock freedom, but differently from Carbone and Debois, their work is based on separation logic. A line of work on many-writer, single-reader process calculi [Pad18, dP18] uses explicit dependency graphs to both ensure resource separation and guarantee deadlock freedom, however it is not immediate how to apply this approach to functional calculi.

9. CONCLUSION AND FUTURE WORK

HGV exploits hypersequents to resolve fundamental modularity issues with GV. As a consequence, we have obtained a tight operational correspondence between HGV and HCP. HGV is a modular and extensible core calculus for functional programming with *binary* session types. In future we intend to apply hypersequents to *multiparty* versions of CP [CLM⁺16] and GV [JBK22b] to exhibit a similarly strong operational correspondence.

ACKNOWLEDGEMENTS

We are deeply grateful to Marco Peressotti for discovering an error in our presentation of the LTS for HCP in the original conference paper, and for suggesting a fix via the definition of saturated transitions. We thank the CONCUR and LMCS reviewers for their helpful comments and suggestions. This work was supported by EPSRC grants EP/K034413/1 (ABCD), EP/T014628/1 (STARDUST), EP/L01503X/1 (CDT PPar), ERC Consolidator Grant 682315 (Skye), UKRI Future Leaders Fellowship MR/T043830/1 (EHOP), a UK Government ISCF Metrology Fellowship grant, EU HORIZON 2020 MSCA RISE project 778233 (BehAPI), and NSF grant CCF-2044815.

REFERENCES

- [Abr94] Samson Abramsky. Proofs as processes. *Theoretical Computer Science*, 135(1):5–9, 1994. doi:10.1016/0304-3975(94)00103-0.
- [AG20] Federico Aschieri and Francesco A. Genco. Par means parallel: multiplicative linear logic proofs as concurrent functional programs. *Proc. ACM Program. Lang.*, 4(POPL):18:1–18:28, 2020. doi:10.1145/3371086.
- [ALM16] Robert Atkey, Sam Lindley, and J. Garrett Morris. Conflation confers concurrency. In *A List of Successes That Can Change the World*, volume 9600 of *Lecture Notes in Computer Science*, pages 32–55. Springer, 2016. doi:10.1007/978-3-319-30936-1_2.
- [Avr91] Arnon Avron. Hypersequents, logical consequence and intermediate logics for concurrency. *Ann. Math. Artif. Intell.*, 4:225–248, 1991. doi:10.1007/BF01531058.
- [BK01] Nick Benton and Andrew Kennedy. Exceptional syntax. *J. Funct. Program.*, 11(4):395–410, 2001. doi:10.1017/S0956796801004099.

- [Bor98] Michele Boreale. On the expressiveness of internal mobility in name-passing calculi. *Theoretical Computer Science*, 195(2):205–226, 3 1998. doi:10.1016/s0304-3975(97)00220-x.
- [BS94] Gianluigi Bellin and Philip J. Scott. On the pi-calculus and linear logic. *Theoretical Computer Science*, 135(1):11–65, 1994. doi:10.1016/0304-3975(94)00104-9.
- [CD10] Marco Carbone and Søren Debois. A graphical approach to progress for structured communication in web services. In *Proc. of ICE*, volume 38 of *EPTCS*, pages 13–27, 2010. doi:10.4204/EPTCS.38.4.
- [CDM14] Marco Carbone, Ornella Dardha, and Fabrizio Montesi. Progress as compositional lock-freedom. In *Proc. of COORDINATION*, volume 8459 of *Lecture Notes in Computer Science*, pages 49–64. Springer, 2014. doi:10.1007/978-3-662-43376-8_4.
- [CLM⁺16] Marco Carbone, Sam Lindley, Fabrizio Montesi, Carsten Schürmann, and Philip Wadler. Coherence generalises duality: A logical explanation of multiparty session types. In *Proc. of CONCUR*, volume 59 of *LIPIcs*, pages 33:1–33:15. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2016. doi:10.4230/LIPIcs.CONCUR.2016.33.
- [CMS18] Marco Carbone, Fabrizio Montesi, and Carsten Schürmann. Choreographies, logically. *Distributed Comput.*, 31(1):51–67, 2018. doi:10.1007/s00446-017-0295-1.
- [CP10] Luís Caires and Frank Pfenning. Session types as intuitionistic linear propositions. In *Proc. of CONCUR*, volume 6269 of *LNCS*, pages 222–236. Springer, 2010. doi:10.1007/978-3-642-15375-4_16.
- [Dar14] Ornella Dardha. Recursive session types revisited. In *Proc. of BEAT*, volume 162 of *Electron. Proc. in Theor. Comput. Sci.*, pages 27–34, 2014. doi:10.4204/EPTCS.162.4.
- [Dar16] Ornella Dardha. *Type Systems for Distributed Programs: Components and Sessions*, volume 7 of *Atlantis Studies in Computing*. Springer / Atlantis Press, 2016. doi:10.2991/978-94-6239-204-5.
- [DCdY07] Mariangiola Dezani-Ciancaglini, Ugo de’Liguoro, and Nobuko Yoshida. On progress for structured communications. In *Proc. of TGC*, volume 4912 of *Lect. Notes Comput. Sci.*, pages 257–275. Springer, 2007. doi:10.1007/978-3-540-78663-4_18.
- [DCMYD06] Mariangiola Dezani-Ciancaglini, Dimitris Mostrous, Nobuko Yoshida, and Sophia Drossopoulou. Session types for object-oriented languages. In *Proc. of ECOOP*, volume 4067 of *Lect. Notes Comput. Sci.*, pages 328–352. Springer, 2006. doi:10.1007/11785477_20.
- [DG18] Ornella Dardha and Simon J. Gay. A new linear logic for deadlock-free session-typed processes. In *Proc. of FoSSaCS*, volume 10803 of *LNCS*, pages 91–109. Springer, 2018. doi:10.1007/978-3-319-89366-2_5.
- [DGS17] Ornella Dardha, Elena Giachino, and Davide Sangiorgi. Session types revisited. *Inf. Comput.*, 256:253–286, 2017. doi:10.1016/j.ic.2017.06.002.
- [DMN07] Olivier Danvy, Kevin Millikin, and Lasse R. Nielsen. On one-pass CPS transformations. *J. Funct. Program.*, 17(6):793–812, 2007. doi:10.1017/S0956796807006387.
- [dP18] Ugo de’Liguoro and Luca Padovani. Mailbox types for unordered interactions. In *Proc. of ECOOP*, volume 109 of *LIPIcs*, pages 15:1–15:28. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2018. doi:10.4230/LIPIcs.ECOOP.2018.15.
- [DP22] Ornella Dardha and Jorge A. Pérez. Comparing type systems for deadlock freedom. *J. Log. Algebraic Methods Program.*, 124:100717, 2022. doi:10.1016/j.jlamp.2021.100717.
- [FKD⁺21] Simon Fowler, Wen Kokke, Ornella Dardha, Sam Lindley, and J. Garrett Morris. Separating sessions smoothly. In *Proc. of CONCUR*, volume 203 of *LIPIcs*, pages 36:1–36:18. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2021. doi:10.4230/LIPIcs.CONCUR.2021.36.
- [FLMD19] Simon Fowler, Sam Lindley, J. Garrett Morris, and Sára Decova. Exceptional asynchronous session types: session types without tiers. *Proc. ACM Program. Lang.*, 3(POPL):28:1–28:29, 2019. doi:10.1145/3290341.
- [Fow19] Simon Fowler. *Typed Concurrent Functional Programming with Channels, Actors, and Sessions*. PhD thesis, University of Edinburgh, 2019.
- [Gir87] Jean-Yves Girard. Linear logic. *Theor. Comput. Sci.*, 50:1–102, 1987. doi:10.1016/0304-3975(87)90045-4.
- [GV10] Simon J. Gay and Vasco T. Vasconcelos. Linear type theory for asynchronous session types. *Journal of Functional Programming*, 20(1):19–50, 2010. doi:10.1017/S0956796809990268.

- [Hon93] Kohei Honda. Types for dyadic interaction. In *CONCUR*, volume 715 of *Lecture Notes in Computer Science*, pages 509–523. Springer, 1993. doi:10.1007/3-540-57208-2_35.
- [HVK98] Kohei Honda, Vasco Thudichum Vasconcelos, and Makoto Kubo. Language primitives and type discipline for structured communication-based programming. In *Proc. of ESOP*, volume 1381 of *LNCS*, pages 122–138. Springer, 1998. doi:10.1007/BFb0053567.
- [ITT⁺19] Atsushi Igarashi, Peter Thiemann, Yuya Tsuda, Vasco T. Vasconcelos, and Philip Wadler. Gradual session types. *J. Funct. Program.*, 29:e17, 2019. doi:10.1017/S0956796819000169.
- [IYA10] Keigo Imai, Shoji Yuen, and Kiyoshi Agusa. Session type inference in Haskell. In *Proc. of PLACES*, volume 69 of *EPTCS*, pages 74–91, 2010. doi:10.4204/EPTCS.69.6.
- [JBK22a] Jules Jacobs, Stephanie Balzer, and Robbert Krebbers. Connectivity graphs: a method for proving deadlock freedom based on separation logic. *Proc. ACM Program. Lang.*, 6(POPL):1–33, 2022.
- [JBK22b] Jules Jacobs, Stephanie Balzer, and Robbert Krebbers. Multiparty GV: functional multiparty session types with certified deadlock freedom. *Proc. ACM Program. Lang.*, 6(ICFP):466–495, 2022.
- [KD21a] Wen Kokke and Ornela Dardha. Deadlock-free session types in linear haskell. In *Proc. of Haskell*, pages 1–13. ACM, 2021. doi:10.1145/3471874.3472979.
- [KD21b] Wen Kokke and Ornela Dardha. Prioritise the best variation. In *Proc. of FORTE*, volume 12719 of *Lecture Notes in Computer Science*, pages 100–119. Springer, 2021. doi:10.1007/978-3-030-78089-0_6.
- [KMP19a] Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Better late than never: A fully-abstract semantics for classical processes. *PACMPL*, 3(POPL), 2019. doi:10.1145/3290337.
- [KMP19b] Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Taking linear logic apart. In Thomas Ehrhard, Maribel Fernández, Valeria de Paiva, and Lorenzo Tortora de Falco, editors, *Linearity-TLLA*, volume 292 of *EPTCS*, pages 90–103. Open Publishing Association, 2019. doi:10.4204/EPTCS.292.5.
- [Kob03] Naoki Kobayashi. Type systems for concurrent programs. In Bernhard K. Aichernig and Tom Maibaum, editors, *Formal Methods at the Crossroads. From Panacea to Foundational Support: 10th Anniversary Colloquium of UNU/IIST, the International Institute for Software Technology of The United Nations University, Lisbon, Portugal, March 18-20, 2002. Revised Papers*, pages 439–453. Springer Berlin Heidelberg, 2003. doi:10.1007/978-3-540-40007-3_26.
- [Kob06] Naoki Kobayashi. A new type system for deadlock-free processes. In *Proc. of CONCUR*, volume 4137 of *LNCS*, pages 233–247. Springer, 2006. doi:10.1007/11817949_16.
- [LM99] Jean-Jacques Lévy and Luc Maranget. Explicit substitutions and programming languages. In *Foundations of Software Technology and Theoretical Computer Science, 1999*, volume 1738 of *LNCS*. Springer, 1999. doi:10.1007/3-540-46691-6_14.
- [LM15] Sam Lindley and J. Garrett Morris. A semantics for propositions as sessions. In Jan Vitek, editor, *Programming Languages and Systems*, pages 560–584. Springer Berlin Heidelberg, 2015. doi:10.1007/978-3-662-46669-8_23.
- [LM16] Sam Lindley and J. Garrett Morris. Talking bananas: Structural recursion for session types. *SIGPLAN Not.*, 51(9):434–447, 2016. doi:10.1145/3022670.2951921.
- [LM17] Sam Lindley and J. Garrett Morris. Lightweight functional session types. In Simon Gay and Antonio Ravara, editors, *Behavioural Types: from Theory to Tools*, chapter 12, pages 265–286. River publishers, 2017.
- [LPT03] Paul Blain Levy, John Power, and Hayo Thielecke. Modelling environments in call-by-value programming languages. *Inf. Comput.*, 185(2):182–210, 2003. doi:10.1016/S0890-5401(03)00088-9.
- [Mon13] Fabrizio Montesi. *Choreographic Programming*. PhD thesis, IT University of Copenhagen, 2013.
- [MP18] Fabrizio Montesi and Marco Peressotti. Classical transitions. *CoRR*, abs/1803.01049, 2018. URL: <https://arxiv.org/abs/1803.01049>, arXiv:1803.01049.
- [MP21] Fabrizio Montesi and Marco Peressotti. Linear logic, the π -calculus, and their metatheory: A recipe for proofs as processes. *CoRR*, abs/2106.11818, 2021. URL: <https://arxiv.org/abs/2106.11818>, arXiv:2106.11818.
- [MV18] Dimitris Mostrous and Vasco T. Vasconcelos. Affine sessions. *Log. Methods Comput. Sci.*, 14(4), 2018. doi:10.1007/978-3-662-43376-8_8.

- [NT04] Matthias Neubauer and Peter Thiemann. An implementation of session types. In *Proc. of PADL*, volume 3057 of *Lecture Notes in Computer Science*, pages 56–70. Springer, 2004. doi:10.1007/978-3-540-24836-1_5.
- [OY16] Dominic A. Orchard and Nobuko Yoshida. Effects as sessions, sessions as effects. In *Proc. of POPL*, pages 568–581. ACM, 2016. doi:10.1145/2837614.2837634.
- [OY17] Dominic Orchard and Nobuko Yoshida. Session types with linearity in Haskell. *Behavioural Types: from Theory to Tools*, page 219, 2017.
- [Pad14] Luca Padovani. Deadlock and Lock Freedom in the Linear π -Calculus. In *Proc. of CSL-LICS*, pages 72:1–72:10. ACM, 2014. doi:10.1145/2603088.2603116.
- [Pad18] Luca Padovani. Deadlock-free typestate-oriented programming. *Art Sci. Eng. Program.*, 2(3):15, 2018. doi:10.22152/programming-journal.org/2018/2/15.
- [Plo75] Gordon D. Plotkin. Call-by-name, call-by-value and the lambda-calculus. *Theor. Comput. Sci.*, 1(2):125–159, 1975. doi:10.1016/0304-3975(75)90017-1.
- [PT08] Riccardo Pucella and Jesse A. Tov. Haskell session types with (almost) no class. In *Proc. of Haskell*. ACM, 2008. doi:10.1145/1411286.1411290.
- [Rey00] John C. Reynolds. The meaning of types—from intrinsic to extrinsic semantics. Technical Report RS-00-32, BRICS, 2000.
- [San96] Davide Sangiorgi. π -calculus, internal mobility, and agent-passing calculi. *Theoretical Computer Science*, 167(1-2):235–274, 1996. doi:10.1016/0304-3975(96)00075-8.
- [SE08] Matthew Sackman and Susan Eisenbach. Session types in Haskell: Updating message passing for the 21st century. Unpublished manuscript, 2008.
- [TCP13] Bernardo Toninho, Luís Caires, and Frank Pfenning. Higher-order processes, functions, and sessions: A monadic integration. In *ESOP*, volume 7792 of *Lecture Notes in Computer Science*, pages 350–369. Springer, 2013. doi:10.1007/978-3-642-37036-6_20.
- [THK94] Kaku Takeuchi, Kohei Honda, and Makoto Kubo. An interaction-based language and its typing system. In *Proc. of PARLE*, volume 817 of *LNCS*, pages 398–413. Springer, 1994. doi:10.1007/3-540-58184-7_118.
- [TV20] Peter Thiemann and Vasco T. Vasconcelos. Label-dependent session types. *Proceedings of the ACM on Programming Languages*, 4(POPL):1–29, 2020. doi:10.1145/3371135.
- [Vas12] Vasco T. Vasconcelos. Fundamentals of session types. *Inf. Comput.*, 217:52–70, 2012. doi:10.1016/j.ic.2012.05.002.
- [VGR06] Vasco Thudichum Vasconcelos, Simon J. Gay, and Antonio Ravara. Type checking a multi-threaded functional language with session types. *Theor. Comput. Sci.*, 368(1-2):64–87, 2006. doi:10.1016/j.tcs.2006.06.028.
- [VRG04] Vasco Vasconcelos, Antonio Ravara, and Simon J. Gay. Session types for functional multithreading. In *CONCUR*, volume 3170 of *LNCS*, pages 497–511. Springer, 2004. doi:10.1007/978-3-540-28644-8_32.
- [VV13] Hugo Torres Vieira and Vasco Thudichum Vasconcelos. Typing progress in communication-centred systems. In *Proc. of COORDINATION*, volume 7890 of *Lect. Notes Comput. Sci.*, pages 236–250. Springer, 2013. doi:10.1007/978-3-662-43376-8_10.
- [Wad14] Philip Wadler. Propositions as sessions. *Journal of Functional Programming*, 24(2-3):384–418, 2014. doi:10.1145/2398856.2364568.

APPENDICES

Appendix A. Omitted Proofs for Section 3: Hypersequent GV	31
A.1. Derived typing rules for syntactic sugar	31
A.2. Preservation Proof	31
A.3. Progress	37
Appendix B. Omitted Proofs for section 4: Relation between HGV and GV	38
Appendix C. Omitted Proofs for section 5: Relation between HGV and CP	40
C.1. Full definition of HGV*	40
C.2. Translating HGV* to HCP	41

APPENDIX A. OMITTED PROOFS FOR SECTION 3: HYPERSEQUENT GV

In this Appendix we give full definitions and proofs for Section 3.

$\frac{\text{T-SEQ} \quad \Gamma \vdash M : \mathbf{1} \quad \Delta \vdash N : T}{\Gamma, \Delta \vdash M; N : T}$	$\frac{\text{T-LAMUNIT} \quad \Gamma \vdash M : T}{\Gamma \vdash \lambda().M : \mathbf{1} \multimap T}$	$\frac{\text{T-LAMPAIR} \quad \Gamma, x : T, y : T' \vdash M : U}{\Gamma \vdash \lambda(x, y).M : T \times T' \multimap U}$
$\frac{\text{T-LET} \quad \Gamma \vdash M : T \quad \Delta, x : T \vdash N : U}{\Gamma, \Delta \vdash \text{let } x = M \text{ in } N : U}$	$\frac{\text{T-SELECT-INL}}{\cdot \vdash \text{select inl} : S \oplus S' \multimap S}$	
$\frac{\text{T-SELECT-INR}}{\cdot \vdash \text{select inr} : S \oplus S' \multimap S'}$	$\frac{\text{T-OFFER} \quad \Gamma \vdash L : S \& S' \quad \Delta, x : S \vdash M : T \quad \Delta, y : S' \vdash N : T}{\Gamma, \Delta \vdash \text{offer } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} : T}$	
	$\frac{\text{T-OFFER-ABSURD} \quad \Gamma \vdash L : \&\{\}}{\Gamma, \Delta \vdash \text{offer } L \{ \} : T}$	

FIGURE 9. Derived rules for syntactic sugar

A.1. Derived typing rules for syntactic sugar. The main body makes use of syntactic sugar, and encodings of branching and selection. Figure 9 shows the derived typing rules.

A.2. Preservation Proof. Next, we detail the proof of preservation. We begin with the usual lemmas to manipulate evaluation contexts, and the usual substitution lemma.

Lemma A.1 (Subterm typeability). *Suppose $\mathbf{D}$ is a derivation of $\Gamma \vdash E[M] : T$. Then, there exist Γ_1, Γ_2 such that $\Gamma = \Gamma_1, \Gamma_2$, a type U , and some subderivation $\mathbf{D}'$ of $\mathbf{D}$ concluding $\Gamma_2 \vdash M : U$, where the position of $\mathbf{D}'$ in $\mathbf{D}$ coincides with the position of the hole in $\mathbf{D}$.*

Proof. By induction on the structure of E . □

Lemma A.2 (Replacement, Evaluation Contexts). *If:*

- $\mathbf{D}$ is a derivation of $\Gamma_1, \Gamma_2 \vdash E[M] : T$
 - $\mathbf{D}'$ is a subderivation of $\mathbf{D}$ concluding $\Gamma_2 \vdash M : U$
 - The position of $\mathbf{D}'$ in $\mathbf{D}$ corresponds to that of the hole in E
 - $\Gamma_3 \vdash N : U$
 - Γ_1, Γ_3 is defined
- then $\Gamma_1, \Gamma_3 \vdash E[N] : T$.

Proof. By induction on the structure of E . □

Lemma A.3 (Substitution). *If:*

- (1) $\Gamma_1, x : U \vdash M : T$
 - (2) $\Gamma_2 \vdash N : U$
 - (3) Γ_1, Γ_2 is defined
- then $\Gamma_1, \Gamma_2 \vdash M\{N/x\} : T$.

Proof. By induction on the derivation of $\Gamma_1, x : U \vdash M : T$. □

Preservation of typing under term reduction is standard.

Lemma A.4 (Preservation, $\longrightarrow_M$). *If $\Gamma \vdash M : T$ and $M \longrightarrow_M N$, then $\Gamma \vdash N : T$.*

Proof. A standard induction on the derivation of $\longrightarrow_M$. □

Runtime type merging is commutative and associative. We make use of these properties implicitly in the remainder of the proofs.

Lemma A.5.

- (1) $R_1 \sqcap R_2 \iff R_2 \sqcap R_1$
- (2) $R_1 \sqcap (R_2 \sqcap R_3) \iff (R_1 \sqcap R_2) \sqcap R_3$

Proof. Immediate from the definition of $\sqcap$. □

The first more major result is preservation of configuration typing under structural congruence.

Lemma A.6 (Preservation ($\equiv$)). *If $\mathcal{G} \vdash \mathcal{C} : R$ and $\mathcal{C} \equiv \mathcal{D}$, then $\mathcal{G} \vdash \mathcal{D} : R$.*

Proof. We consider the cases for the equivalence axioms; the congruence cases are straightforward applications of the IH.

Case (SC-PARASSOC).

$$\mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E}) \equiv (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}$$

$$\frac{\mathcal{G}_1 \vdash \mathcal{C} : R_1 \quad \frac{\mathcal{G}_2 \vdash \mathcal{D} : R_2 \quad \mathcal{G}_3 \vdash \mathcal{E} : R_3}{\mathcal{G}_2 \parallel \mathcal{G}_3 \vdash \mathcal{D} \parallel \mathcal{E} : R_2 \sqcap R_3}}{\mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \mathcal{G}_3 \vdash \mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E}) : R_1 \sqcap R_2 \sqcap R_3} \iff \frac{\frac{\mathcal{G}_1 \vdash \mathcal{C} : R_1 \quad \mathcal{G}_2 \vdash \mathcal{D} : R_2}{\mathcal{G}_1 \parallel \mathcal{G}_2 \vdash \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2} \quad \mathcal{G}_3 \vdash \mathcal{E} : R_3}{\mathcal{G}_1 \parallel \mathcal{G}_2 \parallel \mathcal{G}_3 \vdash (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E} : R_1 \sqcap R_2 \sqcap R_3}$$

Case (SC-PARCOMM).

$$\mathcal{C} \parallel \mathcal{D} \equiv \mathcal{D} \parallel \mathcal{C}$$

$$\frac{\mathcal{G} \vdash \mathcal{C} : R_1 \quad \mathcal{H} \vdash \mathcal{D} : R_2}{\mathcal{G} \parallel \mathcal{H} \vdash \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2} \iff \frac{\mathcal{H} \vdash \mathcal{D} : U \quad \mathcal{G} \vdash \mathcal{C} : T}{\mathcal{G} \parallel \mathcal{H} \vdash \mathcal{D} \parallel \mathcal{C} : R_1 \sqcap R_2}$$

Case (SC-NEWCOMM).

$$(\nu xx')(\nu yy')\mathcal{C} \equiv (\nu yy')(\nu xx')\mathcal{C}$$

Two illustrative subcases:

Subcase (1).

$$\frac{\frac{\mathcal{G} \parallel \Gamma_1, x : S \parallel \Gamma_2, x' : \bar{S} \parallel \Gamma_3, y : S' \parallel \Gamma_4, y' : \bar{S}' \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma_1, x : S \parallel \Gamma_2, x' : \bar{S} \parallel \Gamma_3, \Gamma_4 \vdash (\nu yy')\mathcal{C} : R}}{\mathcal{G} \parallel \Gamma_1, \Gamma_2 \parallel \Gamma_3, \Gamma_4 \vdash (\nu xx')(\nu yy')\mathcal{C} : R} \iff \frac{\mathcal{G} \parallel \Gamma_1, y : S' \parallel \Gamma_2, y' : \bar{S}' \parallel \Gamma_3, x : S \parallel \Gamma_4, x' : \bar{S} \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma_1, y : S' \parallel \Gamma_2, y' : \bar{S}' \parallel \Gamma_3, \Gamma_4 \vdash (\nu xx')\mathcal{C} : R}}{\mathcal{G} \parallel \Gamma_1, \Gamma_2 \parallel \Gamma_3, \Gamma_4 \vdash (\nu yy')(\nu xx')\mathcal{C} : R}$$

Subcase (2).

$$\frac{\frac{\mathcal{G} \parallel \Gamma_1, x : S, y : S' \parallel \Gamma_2, y' : \bar{S}' \parallel \Gamma_3, x' : \bar{S} \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma_1, \Gamma_2, x : S \parallel \Gamma_3, x' : \bar{S} \vdash (\nu yy')\mathcal{C} : R}}{\mathcal{G} \parallel \Gamma_1, \Gamma_2, \Gamma_3 \vdash (\nu xx')(\nu yy')\mathcal{C} : R} \iff \frac{\mathcal{G} \parallel \Gamma_1, x : S, y : S' \parallel \Gamma_2, y' : \bar{S}' \parallel \Gamma_3, x' : \bar{S} \vdash \mathcal{C} : R}{\mathcal{G} \parallel \Gamma_1, \Gamma_3, y : S' \parallel \Gamma_2, y' : \bar{S}' \vdash (\nu xx')\mathcal{C} : R}}{\mathcal{G} \parallel \Gamma_1, \Gamma_2, \Gamma_3 \vdash (\nu yy')(\nu xx')\mathcal{C} : R}$$

Case (SC-NEWSWAP).

$$(\nu xy)\mathcal{C} \equiv (\nu yx)\mathcal{C}$$

Follows immediately since hyper-environments are treated as unordered.

Case (SC-SCOPEEXT).

$$\mathcal{C} \parallel (\nu xy)\mathcal{D} \equiv (\nu xy)(\mathcal{C} \parallel \mathcal{D})$$

(where $x, y \notin \text{fv}(\mathcal{C})$)

$$\frac{\frac{\mathcal{G} \vdash \mathcal{C} : R_1 \quad \mathcal{H} \parallel \Gamma_1, x : S \parallel \Gamma_2, y : \bar{S} \vdash \mathcal{D} : R_2}{\mathcal{G} \vdash \mathcal{C} : R_1 \quad \mathcal{H} \parallel \Gamma_1, \Gamma_2 \vdash (\nu xy)\mathcal{D} : R_2}}{\mathcal{G} \parallel \mathcal{H} \parallel \Gamma_1, \Gamma_2 \vdash \mathcal{C} \parallel (\nu xy)\mathcal{D} : R_1 \sqcap R_2} \iff \frac{\mathcal{G} \vdash \mathcal{C} : R_1 \quad \mathcal{H} \parallel \Gamma_1, x : S \parallel \Gamma_2, y : \bar{S} \vdash \mathcal{D} : R_2}{\mathcal{G} \parallel \mathcal{H} \parallel \Gamma_1, x : S \parallel \Gamma_2, y : \bar{S} \vdash \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2}}{\mathcal{G} \parallel \mathcal{H} \parallel \Gamma_1, \Gamma_2 \vdash (\nu xy)(\mathcal{C} \parallel \mathcal{D}) : R_1 \sqcap R_2}$$

Case (SC-LINKCOMM).

$$x \xleftrightarrow{\tau} y \equiv y \xleftrightarrow{\tau} x$$

Assumption:

$$\frac{}{x : S, y : \bar{S} \vdash x \xleftrightarrow{\tau} y : \circ}$$

By dualising both variables, we have that $x : \bar{S}, y : \bar{\bar{S}}$. Since duality is an involution, we can show $x : S, y : \bar{S} \iff x : \bar{S}, y : S$.

Thus:

$$\overline{y : S, x : \bar{S} \vdash y \overset{z}{\leftrightarrow} x : \circ}$$

The reasoning for the symmetric case is identical. $\square$

The next result shows that configuration typeability is preserved under configuration reduction. Note that this lemma makes crucial use of Lemma A.6 due to E-EQUIV.

Lemma A.7 (Preservation ($\longrightarrow$)). *If $\mathcal{G} \vdash \mathcal{C} : R$ and $\mathcal{C} \longrightarrow \mathcal{D}$, then $\mathcal{G} \vdash \mathcal{D} : R$.*

Proof. By induction on the derivation of $\mathcal{C} \longrightarrow \mathcal{D}$. Where there is a choice for ϕ , we prove the case for $\phi = \bullet$ and expand $\mathcal{T}[M]$ to $\bullet(E[M])$ for some evaluation context E ; the other cases are similar.

Case (E-Reify-Fork).

$$\bullet E[\mathbf{fork} V] \longrightarrow (\nu xy)(\bullet E[x] \parallel \circ V y)$$

Assumption:

$$\frac{\Gamma \vdash E[\mathbf{fork} V] : T}{\Gamma \vdash \bullet E[\mathbf{fork} V] : T}$$

By Lemma A.1, there exist Γ_1, Γ_2, S such that $\Gamma = \Gamma_1, \Gamma_2$ and $\Gamma_1, \Gamma_2 \vdash E[\mathbf{fork} V] : T$ and:

$$\frac{\Gamma_2 \vdash V : S \multimap \mathbf{end}_!}{\Gamma_2 \vdash \mathbf{fork} V : \bar{S}}$$

By Lemma A.2:

$$\frac{\Gamma_1, x : \bar{S} \vdash E[x] : T}{\Gamma_1, x : \bar{S} \vdash \bullet E[x] : T}$$

By TM-APP, $\Gamma_2, y : S \vdash V y : \mathbf{end}_!$ and so by TC-CHILD, $\Gamma_2, y : S \vdash V y : \circ$ Recomposing:

$$\frac{\frac{\Gamma_1, x : \bar{S} \vdash E[x] : T}{\Gamma_1, x : \bar{S} \vdash \bullet E[x] : T} \quad \frac{\Gamma_2, y : S \vdash V y : \mathbf{end}_!}{\Gamma_2, y : S \vdash \circ(V y) : \circ}}{\frac{\Gamma_1, x : \bar{S} \parallel \Gamma_2, y : S \vdash \bullet E[x] \parallel \circ(V y) : T}{\Gamma_1, \Gamma_2 \vdash (\nu xy)(\bullet E[x] \parallel \circ(V y)) : T}}$$

as required.

Case (E-Comm-Send).

$$(\nu xy)(\bullet E[\mathbf{send}(V, x)] \parallel \circ E'[\mathbf{recv} y]) \longrightarrow (\nu xy)(\bullet E[x] \parallel \circ E'[(V, y)])$$

Assumption:

$$\frac{\frac{\Gamma, x : S \vdash E[\mathbf{send}(V, x)] : U}{\Gamma, x : S \vdash \bullet E[\mathbf{send}(V, x)] : U} \quad \frac{\Gamma', y : \bar{S} \vdash E'[\mathbf{recv} y] : \mathbf{end}_!}{\Gamma', y : \bar{S} \vdash \circ E'[\mathbf{recv} y] : \circ}}{\frac{\Gamma, x : S \parallel \Gamma', y : \bar{S} \vdash \bullet E[\mathbf{send}(V, x)] \parallel \circ E'[\mathbf{recv} y] : U}{\Gamma, \Gamma' \vdash (\nu xy)(\bullet E[\mathbf{send}(V, x)] \parallel \circ E'[\mathbf{recv} y]) : U}}$$

By Lemma A.1, there exist Γ_1, Γ_2, S such that $\Gamma = \Gamma_1, \Gamma_2$, and $\Gamma_1, \Gamma_2, x : S \vdash E[\text{send}(V, x)] : U$ and:

$$\frac{\Gamma_2 \vdash V : T \quad x : !T.S' \vdash x : !T.S'}{\Gamma_2, x : !T.S' \vdash \text{send}(V, x) : S'}$$

With the knowledge that $S = !T.S'$, we can refine our original derivation:

$$\frac{\frac{\Gamma_1, \Gamma_2, x : !T.S' \vdash E[\text{send}(V, x)] : U}{\Gamma_1, \Gamma_2, x : !T.S' \vdash \bullet E[\text{send}(V, x)] : U} \quad \frac{\Gamma', y : ?T.\bar{S}' \vdash E'[\text{recv } y] : \text{end!}}{\Gamma', y : ?T.\bar{S}' \vdash \circ E'[\text{recv } y] : \circ}}{\frac{\Gamma_1, \Gamma_2, x : !T.S' \parallel \Gamma', y : ?T.\bar{S}' \vdash \bullet E[\text{send}(V, x)] \parallel \circ E'[\text{recv } y] : U}{\Gamma_1, \Gamma_2, \Gamma' \vdash (\nu xy)(\bullet E[\text{send}(V, x)] \parallel \circ E'[\text{recv } y]) : U}}$$

Again by Lemma A.1, we have that $\Gamma', y : ?T.\bar{S}' \vdash E'[\text{recv } y] : \text{end!}$ and:

$$\frac{y : ?T.\bar{S}' \vdash y : ?T.\bar{S}'}{y : ?T.\bar{S}' \vdash \text{recv } y : T \times \bar{S}'}$$

We can show:

$$\frac{\Gamma_2 \vdash V : T \quad y : \bar{S}' \vdash y : \bar{S}'}{\Gamma_2, y : \bar{S}' \vdash (V, y) : T \times \bar{S}'}$$

By Lemma A.2, we have that $\Gamma_2, \Gamma', y : \bar{S}' \vdash E'[(V, y)] : \bar{S}'$.

Recomposing:

$$\frac{\frac{\Gamma_1, x : S' \vdash E[x] : U}{\Gamma_1, x : S' \vdash \bullet E[x] : U} \quad \frac{\Gamma_2, \Gamma', y : \bar{S}' \vdash E'[(V, y)] : \text{end!}}{\Gamma_2, \Gamma', y : \bar{S}' \vdash \circ E'[(V, y)] : \circ}}{\frac{\Gamma_1, x : S' \parallel \Gamma_2, \Gamma', y : \bar{S}' \vdash \bullet E[x] \parallel \circ E'[(V, y)] : U}{\Gamma_1, \Gamma_2, \Gamma' \vdash (\nu xy)(\bullet E[x] \parallel \circ E'[(V, y)]) : U}}$$

as required.

Case (E-Comm-Close).

$$(\nu xy)(\mathcal{T}[\text{wait } x] \parallel \circ y) \longrightarrow \mathcal{T}[]$$

Taking $\mathcal{T} = \bullet E$, assumption:

$$\frac{\frac{\Gamma, x : \text{end}_? \vdash E[\text{wait } x] : T}{\Gamma, x : \text{end}_? \vdash \bullet E[\text{wait } x] : T} \quad \frac{y : \text{end!} \vdash y : \text{end!}}{y : \text{end!} \vdash \circ y : \circ}}{\frac{\Gamma, x : \text{end}_? \parallel y : \text{end!} \vdash \bullet E[\text{wait } x] \parallel \circ y : T}{\Gamma \vdash (\nu xy)(\bullet E[\text{wait } x] \parallel \circ y) : T}}$$

By Lemma A.1, we have that:

$$\frac{x : \text{end}_? \vdash x : \text{end}_?}{x : \text{end}_? \vdash \text{wait } x : \mathbf{1}}$$

By Lemma A.2, $\Gamma \vdash E[()] : T$.

Recomposing:

$$\frac{\Gamma \vdash E[()] : T}{\Gamma \vdash \bullet E[()] : T}$$

as required.

Case (E-Reify-Link).

$$F[\mathbf{link}(x, y)] \longrightarrow (\nu z z')(x \xleftrightarrow{z} y \parallel F[z'])$$

where z, z' fresh.

Taking $F = \bullet E$, we have that:

$$\frac{\Gamma \vdash E[\mathbf{link}(x, y)] : T}{\Gamma \vdash \bullet E[\mathbf{link}(x, y)] : T}$$

By Lemma A.1, we have that $\Gamma = \Gamma', x : S, y : \bar{S}$ such that:

$$\frac{\frac{x : S \vdash x : S \quad y : \bar{S} \vdash y : \bar{S}}{x : S, y : \bar{S} \vdash (x, y) : S \times \bar{S}}}{x : S, y : \bar{S} \vdash \mathbf{link}(x, y) : \circ}$$

By Lemma A.2, we have that $\Gamma', z : \mathbf{end}_! \vdash E[z] : T$.

Reconstructing:

$$\frac{\frac{z : \mathbf{end}_?, x : S, y : \bar{S} \vdash x \xleftrightarrow{z} y : \circ \quad \Gamma', z : \mathbf{end}_! \vdash \bullet E[z] : T}{z : \mathbf{end}_?, x : S, y : \bar{S} \parallel \Gamma', z : \mathbf{end}_! \vdash x \xleftrightarrow{z} y \parallel \bullet E[z] : T}}{\Gamma', x : S, y : \bar{S} \vdash (\nu z z')(x \xleftrightarrow{z} y \parallel \bullet E[z]) : T}$$

as required.

Case (E-Comm-Link).

$$(\nu z z')(\nu x x')(x \xleftrightarrow{z} y \parallel \circ z \parallel \bullet M) \longrightarrow \bullet(M\{y/x\})$$

Assumption:

$$\frac{\frac{x : S, y : \bar{S}, z : \mathbf{end}_? \vdash x \xleftrightarrow{z} y : \circ \quad \frac{\frac{z' : \mathbf{end}_! \vdash z : \mathbf{end}_! \quad \Gamma, x' : \bar{S} \vdash M : T}{z' : \mathbf{end}_! \vdash \circ z : \circ} \quad \Gamma, x' : \bar{S} \vdash \bullet M : T}{z' : \mathbf{end}_! \parallel \Gamma, x' : \bar{S} \vdash \circ z \parallel \bullet M : T}}{x : S, y : \bar{S}, z : \mathbf{end}_? \parallel z' : \mathbf{end}_! \parallel \Gamma, x' : \bar{S} \vdash x \xleftrightarrow{z} y \parallel \circ z' \parallel \bullet M : T}}{\Gamma, y : \bar{S}, z : \mathbf{end}_? \parallel z' : \mathbf{end}_! \vdash (\nu x x')(x \xleftrightarrow{z} y \parallel \circ z' \parallel \bullet M) : T}}{\Gamma, y : \bar{S} \vdash (\nu z z')(\nu x x')(x \xleftrightarrow{z} y \parallel \circ z' \parallel \bullet M) : T}$$

By Lemma A.3, $\Gamma, y' : \bar{S} \vdash M\{y/x'\} : T$, thus:

$$\frac{\Gamma, y' : \bar{S} \vdash M\{y/x'\} : T}{\Gamma, y' : \bar{S} \vdash \bullet M\{y/x'\} : T}$$

as required.

Case (E-Res).

$$(\nu xy)\mathcal{C} \longrightarrow (\nu xy)\mathcal{D} \quad \text{if } \mathcal{C} \longrightarrow \mathcal{D}$$

Immediate by the IH.

Case (E-Par).

$$\mathcal{C} \parallel \mathcal{D} \longrightarrow \mathcal{C}' \parallel \mathcal{D} \quad \text{if } \mathcal{C} \longrightarrow \mathcal{C}'$$

Immediate by the IH.

Case (E-Equiv).

$$\mathcal{C} \longrightarrow \mathcal{D} \quad \text{if } \mathcal{C} \equiv \mathcal{C}', \mathcal{C}' \longrightarrow \mathcal{D}', \text{ and } \mathcal{D}' \equiv \mathcal{D}$$

Assumption: $\mathcal{G} \vdash \mathcal{C} : R$.

By Lemma A.6, $\mathcal{G} \vdash \mathcal{C}' : R$.

By the IH, $\mathcal{G} \vdash \mathcal{D}' : R$.

By Lemma A.6, $\mathcal{G} \vdash \mathcal{D} : R$, as required.

Case (E-Lift-M).

$$\phi M \longrightarrow \phi N \quad \text{if } M \longrightarrow_M N$$

Immediate by Lemma A.4. □

Theorem 3.3 (Preservation).

- (1) If $\mathcal{G} \vdash \mathcal{C} : R$ and $\mathcal{C} \equiv \mathcal{D}$, then $\mathcal{G} \vdash \mathcal{D} : R$.
- (2) If $\mathcal{G} \vdash \mathcal{C} : R$ and $\mathcal{C} \longrightarrow \mathcal{D}$, then $\mathcal{G} \vdash \mathcal{D} : R$.

Proof. A direct corollary of Lemmas A.6 and A.7. □

A.3. Progress. Functional reduction satisfies progress: under an environment only containing runtime names, a term will either reduce, be a value, or be ready to perform a communication action.

Lemma A.8 (Progress, Terms). *If $\Psi \vdash M : T$, then either M is a value, or there exists some N such that $M \longrightarrow_M N$, or M can be written $E[N]$ for some $N \in \{\text{fork } V, \text{send } (V, W), \text{recv } V, \text{wait } V, \text{link } (V, W)\}$.*

Proof. A standard induction on the derivation of $\Psi \vdash M : T$. □

Note that tree canonical forms can be defined inductively:

$$\mathcal{F} ::= \phi M \mid (\nu xy)(\mathcal{A} \parallel \mathcal{F})$$

We assume the same requirement for configurations $\mathcal{F}$ as the non-inductive definition of tree canonical forms: i.e., that for a configuration $(\nu xy)(\mathcal{A} \parallel \mathcal{F})$, that $x \in \text{fv}(\mathcal{A})$.

Lemma 3.17 follows as a direct corollary of a slightly more verbose property, which follows from the inductive definition of TCFs.

Definition A.9 (Open progress). Suppose $\Psi \vdash \mathcal{F} : R$, where $\mathcal{F} \not\rightarrow$. We say that $\mathcal{F}$ *satisfies open progress* if:

- (1) $\mathcal{C} = (\nu xx')(\mathcal{A} \parallel \mathcal{F}')$, where:
 - (a) There exist Ψ_1, Ψ_2 such that $\Psi = \Psi_1, \Psi_2$
 - (b) $\Psi_1, x:S \vdash \mathcal{A} : \circ$ for some session type S , and $\text{blocked}(\mathcal{A}, y)$ for some $y \in \text{fv}(\Psi_1, x:S)$
 - (c) $\Psi_2, x':\bar{S} \vdash \mathcal{D} : R$, where $\mathcal{F}'$ satisfies open progress
- (2) $\mathcal{F} = \phi M$, and either M is a value, or $\text{blocked}(\phi M, x)$ for some $x \in \text{fv}(\Psi)$.

Lemma A.10 (Open progress). *If $\Psi \vdash \mathcal{F} : R$ and $\mathcal{F} \not\rightarrow$, then $\mathcal{F}$ satisfies open progress.*

Proof. By induction on the derivation of $\mathcal{G} \vdash \mathcal{F} : R$. By the definition of canonical forms, it must be the case that $\mathcal{C}$ is of the form $(\nu xy)(\mathcal{A} \parallel \mathcal{F}')$ or $\bullet M$.

We show the case where $\mathcal{C} = (\nu xy)(\circ M \parallel \mathcal{F}')$; the cases for $\mathcal{A} = x \xrightarrow{z} x'$ and $\mathcal{C} = \bullet M$ follow similar reasoning.

Assumption:

$$\frac{\frac{\Psi_1, x : S \vdash \mathcal{A} : \circ \quad \Psi_2, y : \bar{S} \vdash \mathcal{F}' : R}{\Psi_1, x : S \parallel \Psi_2, y : \bar{S} \vdash \mathcal{A} \parallel \mathcal{F}' : R}}{\Psi_1, \Psi_2 \vdash (\nu xy)(\circ M \parallel \mathcal{F}') : R}$$

In both cases, by the induction hypothesis, $\Psi_2, y : \bar{S} \vdash \mathcal{F}' : T$ satisfies open progress.

Subcase ($\mathcal{A} = \circ M$). By Lemma A.8, either M is a value, or M can be written $E[N]$ for some communication and concurrency construct $N \in \{\text{fork } V, \text{send } (V, W), \text{recv } V, \text{wait } V, \text{link } (V, W)\}$.

Otherwise, M is a communication or concurrency construct. If $N = \text{fork } V$, then reduction could occur by E-REIFY-FORK. If $N = \text{link } (V, W)$, then by the type schema for **link**, we have that **link** (V, W) must be of the form **link** (z, z') for $z, z' \in \text{fv}(\Psi, x : S)$ and could reduce by E-REIFY-LINK.

Otherwise, it must be the case that $\text{blocked}(\circ M, z)$ for some $z \in \text{fv}(\Psi_1, x : S)$.

Thus, $(\nu xy)(\circ M \parallel \mathcal{D})$ satisfies open progress, as required.

Subcase ($\mathcal{A} = z_2 \xrightarrow{z_1} z_3$). We have that $z_1, z_2, z_3 \in \text{fv}(\Psi_1, x : S)$, and the thread must be blocked by definition. $\square$

APPENDIX B. OMITTED PROOFS FOR SECTION 4: RELATION BETWEEN HGV AND GV

Theorem 4.3 (Typeability of GV configurations in HGV). *If $\Gamma \vdash_{\text{GV}} \mathcal{C} : R$, then there exists some $\mathcal{G}$ such that $\mathcal{G}$ is a splitting of Γ and $\mathcal{G} \vdash \mathcal{C} : R$.*

Proof. By induction on the derivation of $\Gamma \vdash \mathcal{C} : R$.

Case (TG-NEW). Assumption:

$$\frac{\Gamma, \langle y, y' \rangle : S^\# \vdash_{\text{GV}} \mathcal{C} : R}{\Gamma \vdash_{\text{GV}} (\nu yy') \mathcal{C} : R}$$

Suppose $\Gamma = \langle x_1, x'_1 \rangle : S_1^\sharp, \dots, \langle x_n, x'_n \rangle : S_n^\sharp$ (for clarity, without loss of generality, we assume the absence of non-session variables. As these are simply split between environments, they can be added orthogonally).

By the IH, we have that there exists some hyper-environment $\mathcal{G}$ such that $\mathcal{G} \vdash \mathcal{C} : R$, where $\mathcal{G}$ is a splitting of $\Gamma, \langle y, y' \rangle : S^\sharp$.

Since $\mathcal{G}$ is a splitting of $\mathcal{C}$, we know that $y : S \in \mathcal{G}$ and $y' : \bar{S} \in \mathcal{G}$, and that $\mathcal{G}$ has a tree structure with respect to names $\{\{x_1, x'_1\}, \dots, \{x_n, x'_n\}, \{y, y'\}\}$.

Since $\mathcal{G}$ has a tree structure, by definition we have that $\mathcal{G} = \mathcal{G}' \parallel \Gamma_1, y : S \parallel \Gamma_2, y' : \bar{S}$ for some $\mathcal{G}', \Gamma_1, \Gamma_2$, where $\mathcal{G}'$ has a tree structure.

By Lemma 3.12 (clause 1, left-to-right), $\mathcal{G}' \parallel \Gamma_1, \Gamma_2$ has a tree structure with respect to names $\{\{x_1, x'_1\}, \dots, \{x_n, x'_n\}\}$.

Thus, we can show:

$$\frac{\mathcal{G}' \parallel \Gamma_1, y : S \parallel \Gamma_2, y' : \bar{S} \vdash \mathcal{C} : R}{\mathcal{G}' \parallel \Gamma_1, \Gamma_2 \vdash (\nu y y') \mathcal{C} : R}$$

where $\mathcal{G}' \parallel \Gamma_1, \Gamma_2$ has a tree structure with respect to names $\{\{x_1, x'_1\}, \dots, \{x_n, x'_n\}\}$ and is therefore a splitting of Γ , as required.

Case (TG-CONNECT₁). Assumption:

$$\frac{\Gamma_1, y : S \vdash_{\text{GV}} \mathcal{C} : R_1 \quad \Gamma_2, y' : \bar{S} \vdash_{\text{GV}} \mathcal{D} : R_2}{\Gamma_1, \Gamma_2, \langle y, y' \rangle : S^\sharp \vdash_{\text{GV}} \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2}$$

Suppose $\Gamma_1 = \langle x_1, x'_1 \rangle : S_1^\sharp, \dots, \langle x_m, x'_m \rangle : S_m^\sharp$ and $\Gamma_2 = \langle x_{m+1}, x'_{m+1} \rangle : S_{m+1}^\sharp, \dots, \langle x_n, x'_n \rangle : S_n^\sharp$.

By the IH, there exist hyper-environments $\mathcal{G}, \mathcal{H}$ such that:

- (1) $\mathcal{G}$ is a splitting of $\Gamma_1, y : S$
- (2) $\mathcal{H}$ is a splitting of $\Gamma_2, y' : \bar{S}$
- (3) $\mathcal{G} \vdash_{\text{GV}} \mathcal{C} : R_1$
- (4) $\mathcal{H} \vdash_{\text{GV}} \mathcal{D} : R_2$

By the definition of splittings, $\mathcal{G}$ and $\mathcal{H}$ can be written $\mathcal{G} = \mathcal{G}' \parallel \Gamma'_1, y : S$ and $\mathcal{H} = \mathcal{H}' \parallel \Gamma'_2, y' : \bar{S}$ for some Γ'_1, Γ'_2 . Furthermore, $\mathcal{G}$ has a tree structure with respect to $\{\{x_1, x'_1\}, \dots, \{x_m, x'_m\}\}$ and $\mathcal{H}$ has a tree structure with respect to $\{\{x_{m+1}, x'_{m+1}\}, \dots, \{x_n, x'_n\}\}$.

By Lemma 3.12 (clause 2, left-to-right), $\mathcal{G}' \parallel \Gamma'_1, y : S \parallel \mathcal{H}' \parallel \Gamma'_2, y' : \bar{S}$ has a tree structure with respect to $\{\{x_1, x'_1\}, \dots, \{x_n, x'_n\}, \{y, y'\}\}$ and therefore $\mathcal{G} \parallel \mathcal{H}$ is a splitting of $\Gamma_1, \Gamma_2, \langle y, y' \rangle : S^\sharp$.

Recomposing in HGV:

$$\frac{\mathcal{G} \vdash \mathcal{C} : R_1 \quad \mathcal{H} \vdash \mathcal{D} : R_2}{\mathcal{G} \parallel \mathcal{H} \vdash \mathcal{C} \parallel \mathcal{D} : R_1 \sqcap R_2}$$

as required.

Case (TG-CONNECT₂). Similar to TG-CONNECT₁.

Case (TG-CHILD). Assumption:

$$\frac{\Gamma \vdash M : \mathbf{end}_!}{\Gamma \vdash_{\text{GV}} \circ M : \circ}$$

Since we mandated that variables of type $S^\sharp$ cannot appear in terms, there are no names of type $S^\sharp$ in Γ . Therefore, the singleton hyper-environment Γ is a valid splitting, and so we can conclude by TC-CHILD in HGV.

Case (TG-MAIN). Similar to TG-CHILD. $\square$

Proposition 4.5. *Suppose $\Gamma \vdash \mathcal{C} : R$ where $\mathcal{C}$ is in tree canonical form. Then, $\Gamma \vdash_{\text{GV}} \mathcal{C} : R$.*

Proof. By induction on the number of ν -bound names.

In the case that $n = 0$, the result follows immediately by TG-CHILD or TG-MAIN.

In the case that $n \geq 1$, we have that $\Gamma = \Gamma_1, \Gamma_2$ for some Γ_1, Γ_2 and:

$$\frac{\frac{\Gamma_1, x : S \vdash \circ L : \circ \quad \Gamma_2, y : \bar{S} \vdash \mathcal{D} : R}{\Gamma_1, x : S \parallel \Gamma_2, y : \bar{S} \vdash \circ L \parallel \mathcal{D} : R}}{\Gamma_1, \Gamma_2 \vdash (\nu xy)(\circ L \parallel \mathcal{D}) : R}$$

such that $\mathcal{D}$ is in tree canonical form. That $\Gamma_1, x : S \vdash \circ L : \circ$ follows by the definition of tree canonical forms, since $x \in \text{fv}(L)$.

By the IH, $\Gamma_2, y : \bar{S} \vdash \mathcal{D} : R$ in GV.

Thus, we can write:

$$\frac{\frac{\Gamma_1, x : S \vdash \circ L : \circ \quad \Gamma_2, y : \bar{S} \vdash \mathcal{D} : R}{\Gamma_1, \Gamma_2, \langle x, y \rangle : S^\sharp \vdash \circ L \parallel \mathcal{D} : R}}{\Gamma_1, \Gamma_2 \vdash (\nu xy)(\circ L \parallel \mathcal{D}) : R}$$

as required. $\square$

APPENDIX C. OMITTED PROOFS FOR SECTION 5: RELATION BETWEEN HGV AND CP

C.1. Full definition of HGV*.

Syntax.

Terms	$L, M, N ::= V \mid \mathbf{let} \ x = M \ \mathbf{in} \ N \mid V \ W$
	$\mid \mathbf{let} \ () = V \ \mathbf{in} \ M \mid \mathbf{let} \ (x, y) = V \ \mathbf{in} \ M$
	$\mid \mathbf{absurd} \ V \mid \mathbf{case} \ V \ \{\mathbf{inl} \ x \mapsto M; \ \mathbf{inr} \ y \mapsto N\}$
Values	$V, W ::= x \mid K \mid \lambda x. M \mid () \mid (V, W) \mid \mathbf{inl} \ V \mid \mathbf{inr} \ V$
Evaluation contexts	$E ::= \square \mid \mathbf{let} \ x = E \ \mathbf{in} \ M$
Thread contexts	$F ::= \phi \ E$

Typing rules for values

$$\boxed{\Gamma \vdash V : T}$$

$\frac{\text{TV}^*\text{-VAR}}{x : T \vdash x : T}$	$\frac{\text{TV}^*\text{-CONST}}{\cdot \vdash K : T}$	$\frac{\text{TV}^*\text{-LAM}}{\Gamma, x : T \vdash M : U}$	$\frac{\text{TV}^*\text{-UNIT}}{\cdot \vdash () : \mathbf{1}}$
$\frac{\text{TV}^*\text{-PAIR}}{\Gamma \vdash V : T \quad \Delta \vdash W : U}$	$\frac{\text{TV}^*\text{-ABSURD}}{\Gamma \vdash V : \mathbf{0}}$	$\frac{\text{TV}^*\text{-INL}}{\Gamma \vdash V : T}$	$\frac{\text{TV}^*\text{-INR}}{\Gamma \vdash V : U}$
$\frac{}{\Gamma, \Delta \vdash (V, W) : T \times U}$	$\frac{}{\Gamma \vdash \text{absurd } V : T}$	$\frac{}{\Gamma \vdash \text{inl } V : T + U}$	$\frac{}{\Gamma \vdash \text{inr } V : T + U}$
Typing rules for terms			
$\frac{\text{TM}^*\text{-APP}}{\Gamma \vdash V : T \multimap U \quad \Delta \vdash W : T}$		$\frac{\text{TM}^*\text{-LET}}{\Gamma \vdash M : T \quad \Delta, x : T \vdash N : U}$	
$\frac{}{\Gamma, \Delta \vdash V W : U}$		$\frac{}{\Gamma, \Delta \vdash \text{let } x = M \text{ in } N : U}$	
$\frac{\text{TM}^*\text{-LETUNIT}}{\Gamma \vdash V : \mathbf{1} \quad \Delta \vdash M : T}$		$\frac{\text{TM}^*\text{-LETPAIR}}{\Gamma \vdash V : T \times T' \quad \Delta, x : T, y : T' \vdash M : U}$	
$\frac{}{\Gamma, \Delta \vdash \text{let } () = V \text{ in } M : T}$		$\frac{}{\Gamma, \Delta \vdash \text{let } (x, y) = V \text{ in } M : U}$	
$\frac{\text{TM}^*\text{-CASESUM}}{\Gamma \vdash V : T + T'}$			
$\frac{\Delta, x : T \vdash M : U \quad \Delta, y : T' \vdash N : U}{\Gamma, \Delta \vdash \text{case } V \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} : U}$			

The typing of constants is the same as for HGV.

Operational Semantics. The operational semantics for HGV* is the same as for HGV (Figure 4), with the addition of the following explicit rule for **let**:

$$\text{E-LET} \quad \mathbf{let} x = V \mathbf{in} M \longrightarrow M\{V/x\}$$

Similarly, HGV* directly inherits HGV's runtime typing.

C.2. Translating HGV* to HCP. The translation is guaranteed to have only internal (i.e., α or β) transitions and transitions on the dedicated output channel. More specifically:

Lemma C.1.

- If $\llbracket C \rrbracket_r^c \xrightarrow{\ell}$, then $\ell \in \{\alpha, \beta\}$ or $\ell = \ell_r$.
- If $\llbracket M \rrbracket_r^m \xrightarrow{\ell}$ and M is a non-value, then $\ell \in \{\alpha, \beta\}$.
- If V is a value, then $\llbracket V \rrbracket_r^m \xrightarrow{\ell_r}$.
- If $\llbracket V \rrbracket_r^v \xrightarrow{\ell}$ then $\ell \in \{\alpha, \beta\}$.

Proof. By induction on the structure of M . □

We do not use the above lemma directly, but it is a useful sanity check.

Definition C.2 (process contexts). A process context $P[\]$ is a process with a single hole, denoted $\square$. We extend the typing rules, LTS and typing rules to process contexts. We write $P[\] \vdash \mathcal{G}/\mathcal{H}$ to mean that $P[\]$ is typed under hyper-environment $\mathcal{H}$ expecting a process typed under $\mathcal{G}$, i.e., if $Q \vdash \mathcal{G}$ then $P[Q] \vdash \mathcal{H}$.

Definition C.3. A process P is *blocked on* x if it only has transitions $P \xrightarrow{\ell_x}$.

Lemma C.4. If $P[\]$ is a process context with $z, w, w' \notin \text{cn}(P[\])$, and Q is a process blocked on w' , then $(\nu ww')(P[z \leftrightarrow w] \parallel Q) \approx_\alpha P[Q\{z/w'\}]$.

Proof. By induction on the process context $P[\]$.

Case $(\square)$.

$$\begin{aligned} & (\nu ww')(z \leftrightarrow w \parallel Q) \\ & \xrightarrow{\alpha} Q\{z/w'\} \\ & \sim Q\{z/w'\} \quad (\text{by reflexivity}) \end{aligned}$$

Case $((\nu xy)P[\])$.

$$\begin{aligned} & (\nu ww')((\nu xy)(P[z \leftrightarrow w] \parallel Q)) \\ & \sim (\nu xy)(\nu ww')(P[z \leftrightarrow w] \parallel Q) \quad (\text{by Lemma 5.6}) \\ & \approx_{\alpha} (\nu xy)(P[Q\{z/w'\}]) \quad (\text{by Lemma 5.6 and IH}) \end{aligned}$$

Case $(P[\] \parallel R)$.

$$\begin{aligned} & (\nu ww')(P[z \leftrightarrow w] \parallel R \parallel Q) \\ & \sim (\nu ww')(P[z \leftrightarrow w] \parallel Q) \parallel R \quad (\text{by Lemma 5.6}) \\ & \approx_{\alpha} P[Q\{z/w'\}] \parallel R \quad (\text{by Lemma 5.6 and IH}) \end{aligned}$$

Case $(R \parallel P[\])$.

$$\begin{aligned} & (\nu ww')(R \parallel P[z \leftrightarrow w] \parallel Q) \\ & \sim R \parallel (\nu ww')(P[z \leftrightarrow w] \parallel Q) \quad (\text{by Lemma 5.6}) \\ & \approx_{\alpha} R \parallel P[Q\{z/w'\}] \quad (\text{by Lemma 5.6 and IH}) \end{aligned}$$

Case $(\pi.P[\])$. Since Q is blocked on w' , the process $(\nu ww')(\pi.P[z \leftrightarrow w] \parallel Q)$ has only one transition,

$$(\nu ww')(\pi.P[z \leftrightarrow w] \parallel Q) \xrightarrow{\pi} (\nu ww')(P[z \leftrightarrow w] \parallel Q).$$

The process $\pi.P[Q\{z/w'\}]$ has only one transition, also with label π ,

$$\pi.P[Q\{z/w'\}] \xrightarrow{\pi} P[Q\{z/w'\}].$$

The resulting processes are bisimilar by the induction hypothesis.

Case $(x \triangleright \{\text{inl} : P[\]; \text{inr} : P'[\]\})$. Since Q is blocked on w' , the process $(\nu ww')(x \triangleright \{\text{inl} : P[z \leftrightarrow w]; \text{inr} : P'[z \leftrightarrow w]\} \parallel Q)$ has only two transitions,

$$(\nu ww')(x \triangleright \{\text{inl} : P[z \leftrightarrow w]; \text{inr} : P'[z \leftrightarrow w]\} \parallel Q) \xrightarrow{x \triangleright \text{inl}} (\nu ww')(P[z \leftrightarrow w] \parallel Q)$$

and

$$(\nu ww')(x \triangleright \{\text{inl} : P[z \leftrightarrow w]; \text{inr} : P'[z \leftrightarrow w]\} \parallel Q) \xrightarrow{x \triangleright \text{inr}} (\nu ww')(P'[z \leftrightarrow w] \parallel Q).$$

The process $x \triangleright \{\text{inl} : P[Q\{z/w'\}]; \text{inr} : P'[Q\{z/w'\}]\}$ has only two transitions, also with labels $x \triangleright \text{inl}$ and $x \triangleright \text{inr}$,

$$x \triangleright \{\text{inl} : P[Q\{z/w'\}]; \text{inr} : P'[Q\{z/w'\}]\} \xrightarrow{x \triangleright \text{inl}} P[Q\{z/w'\}]$$

and

$$x \triangleright \{\text{inl} : P[Q\{z/w'\}]; \text{inr} : P'[Q\{z/w'\}]\} \xrightarrow{x \triangleright \text{inr}} P'[Q\{z/w'\}].$$

The resulting processes are bisimilar by the induction hypothesis. $\square$

Lemma 5.10 (Substitution). *If M is a well-typed term with $w \in \text{fv}(M)$, and V is a well-typed value, then $(\nu ww')(\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{w'}^v) \approx_{\alpha} \llbracket M\{V/w\} \rrbracket_r^m$.*

Proof. Immediately from Lemma C.4. $\square$

Lemma 5.9 (Type Preservation).

- (1) If $\Gamma \vdash V : T$, then $\llbracket V \rrbracket_r^\vee \vdash \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket^\perp$.
- (2) If $\Gamma \vdash M : T$, then $\llbracket M \rrbracket_r^m \vdash \llbracket \Gamma \rrbracket, r : \mathbf{1} \otimes \llbracket T \rrbracket^\perp$.
- (3) If $\mathcal{G} \parallel \Gamma \vdash \mathcal{C} : T$, where Γ is the type environment for the main thread in $\mathcal{C}$, then $\llbracket \mathcal{C} \rrbracket_r^c \vdash \llbracket \mathcal{G} \rrbracket \parallel \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket^\perp$.

Proof. Part 1.

- **Case** (x). We have $\overline{x : T \vdash x : T}$ and $\llbracket x \rrbracket_r^\vee = r \leftrightarrow x$. We can derive:

$$\overline{x \leftrightarrow r \vdash x : \llbracket A \rrbracket, r : \llbracket A \rrbracket^\perp}$$

- **Case** (K). We have one case for each communication primitive.

– Subcase **link**. We have $\mathbf{link} : S \times \bar{S} \multimap \mathbf{end}_!$, where

$$\begin{aligned} \llbracket S \times S^\perp \multimap \mathbf{end}_! \rrbracket^\perp &= \llbracket S \times S^\perp \multimap \mathbf{end}_! \rrbracket \\ &= \llbracket S \times S^\perp \rrbracket^\perp \wp (\mathbf{1} \otimes \llbracket \mathbf{end}_! \rrbracket) \\ &= (\llbracket S \rrbracket^\perp \wp \llbracket S \rrbracket) \wp (\mathbf{1} \otimes \perp) \end{aligned}$$

and $\llbracket \mathbf{link} \rrbracket_r^\vee = r(y).y(x).\bar{r}.r().x \leftrightarrow y$. We can derive:

$$\begin{aligned} &\overline{x \leftrightarrow y \vdash x : \llbracket S \rrbracket^\perp, y : \llbracket S \rrbracket} \\ &\overline{r().x \leftrightarrow y \vdash x : \llbracket S \rrbracket^\perp, y : \llbracket S \rrbracket, r : \perp} \\ &\overline{\bar{r}.r().x \leftrightarrow y \vdash x : \llbracket S \rrbracket^\perp, y : \llbracket S \rrbracket, r : \mathbf{1} \otimes \perp} \\ &\overline{y(x).\bar{r}.r().x \leftrightarrow y \vdash y : \llbracket S \rrbracket^\perp \wp \llbracket S \rrbracket, r : \mathbf{1} \otimes \perp} \\ &\overline{r(y).y(x).\bar{r}.r().x \leftrightarrow y \vdash r : (\llbracket S \rrbracket^\perp \wp \llbracket S \rrbracket) \wp (\mathbf{1} \otimes \perp)} \end{aligned}$$

– Subcase **fork**. We have $\mathbf{fork} : (S \multimap \mathbf{end}_!) \multimap \bar{S}$ where

$$\begin{aligned} \llbracket (S \multimap \mathbf{end}_!) \multimap S^\perp \rrbracket^\perp &= \llbracket (S \multimap \mathbf{end}_!) \multimap S^\perp \rrbracket \\ &= \llbracket S \multimap \mathbf{end}_! \rrbracket^\perp \wp (\mathbf{1} \otimes \llbracket S^\perp \rrbracket) \\ &= (\llbracket S \rrbracket^\perp \wp (\mathbf{1} \otimes \llbracket \mathbf{end}_! \rrbracket))^\perp \wp (\mathbf{1} \otimes \llbracket S \rrbracket) \\ &= (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp (\mathbf{1} \otimes \llbracket S \rrbracket) \end{aligned}$$

and $\llbracket \mathbf{fork} \rrbracket_r^\vee = (\nu y y')(r(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \parallel y'(x).x\langle y' \rangle.x.x[].\mathbf{0})$. We derive:

$$\begin{aligned} &\overline{r \leftrightarrow y \vdash y : \llbracket S \rrbracket^\perp, r : \llbracket S \rrbracket} \\ &\overline{\bar{r}.r \leftrightarrow y \vdash y : \llbracket S \rrbracket^\perp, r : \mathbf{1} \otimes \llbracket S \rrbracket} \\ &\overline{y\langle x \rangle.\bar{r}.r \leftrightarrow y \vdash y : (\llbracket S \rrbracket \wp (\mathbf{1} \otimes \perp)) \otimes \llbracket S \rrbracket^\perp, x : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})), r : (\mathbf{1} \otimes \llbracket S \rrbracket)} \\ &\overline{(r(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y) \vdash y : (\llbracket S \rrbracket \wp (\mathbf{1} \otimes \perp)) \otimes \llbracket S \rrbracket^\perp, r : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp (\mathbf{1} \otimes \llbracket S \rrbracket)} \quad \mathcal{D} \\ &\overline{(r(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \parallel y'(x).x\langle y' \rangle.x.x[].\mathbf{0}) \vdash y : T, r : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp (\mathbf{1} \otimes \llbracket S \rrbracket) \parallel y' : T^\perp} \\ &\overline{(\nu y y')(r(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \parallel y'(x).x\langle y' \rangle.x.x[].\mathbf{0}) \vdash r : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp (\mathbf{1} \otimes \llbracket S \rrbracket)} \end{aligned}$$

where

$$\begin{aligned} T &= (\llbracket S \rrbracket \wp (\mathbf{1} \otimes \perp)) \otimes \llbracket S \rrbracket^\perp \\ T^\perp &= (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp \llbracket S \rrbracket \end{aligned}$$

and $\mathcal{D}$ is the derivation

$$\frac{\frac{\frac{\overline{\mathbf{0} \vdash \emptyset}}{x[].\mathbf{0} \vdash x : \mathbf{1}}}{x.x[],\mathbf{0} \vdash x : (\perp \wp \mathbf{1})}}{\frac{x\langle y' \rangle.x.x[],\mathbf{0} \vdash x : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})), y' : \llbracket S \rrbracket}{y'(x).x\langle y' \rangle.x.x[],\mathbf{0} \vdash y' : (\llbracket S \rrbracket^\perp \otimes (\perp \wp \mathbf{1})) \wp \llbracket S \rrbracket}}$$

– Subcase **send**: We have **send** : $T \times !T.S \multimap S$ where

$$\begin{aligned} \llbracket T \times !T.S \multimap S \rrbracket^\perp &= \llbracket T \times !T.S \multimap S \rrbracket \\ &= \llbracket T \times !T.S \rrbracket^\perp \wp (1 \otimes \llbracket S \rrbracket) \\ &= (\llbracket T \rrbracket^\perp \wp \llbracket !T.S \rrbracket) \wp (1 \otimes \llbracket S \rrbracket) \\ &= (\llbracket T \rrbracket^\perp \wp (\llbracket T \rrbracket \otimes \llbracket S \rrbracket)) \wp (1 \otimes \llbracket S \rrbracket^\perp) \end{aligned}$$

and $\llbracket \text{send} \rrbracket_r^\vee = r(y).y(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y$. We derive:

$$\frac{\frac{\frac{\overline{r \leftrightarrow y \vdash y : \llbracket S \rrbracket, r : \llbracket S \rrbracket^\perp}}{\bar{r}.r \leftrightarrow y \vdash y : \llbracket S \rrbracket, r : (1 \otimes \llbracket S \rrbracket^\perp)}}{y\langle x \rangle.\bar{r}.r \leftrightarrow y \vdash x : \llbracket T \rrbracket^\perp, y : (\llbracket T \rrbracket \otimes \llbracket S \rrbracket), r : (1 \otimes \llbracket S \rrbracket^\perp)}}{y(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \vdash y : (\llbracket T \rrbracket^\perp \wp (\llbracket T \rrbracket \otimes \llbracket S \rrbracket)), r : (1 \otimes \llbracket S \rrbracket^\perp)}}{r(y).y(x).y\langle x \rangle.\bar{r}.r \leftrightarrow y \vdash r : (\llbracket T \rrbracket^\perp \wp (\llbracket T \rrbracket \otimes \llbracket S \rrbracket)) \wp (1 \otimes \llbracket S \rrbracket^\perp)}$$

– Subcase **recv**: We have **recv** : $?T.S \multimap T \times S$ where

$$\begin{aligned} \llbracket ?T.S \multimap T \times S \rrbracket^\perp &= \llbracket ?T.S \multimap T \times S \rrbracket \\ &= \llbracket ?T.S \rrbracket^\perp \wp (1 \otimes \llbracket T \times S \rrbracket) \\ &= (\llbracket T \rrbracket^\perp \wp \llbracket S \rrbracket^\perp) \wp (1 \otimes (\llbracket T \rrbracket \otimes \llbracket S \rrbracket)) \end{aligned}$$

and $\llbracket \text{recv} \rrbracket_r^\vee = r(x).x(y).\bar{r}.r\langle y \rangle.r \leftrightarrow x$. We derive:

$$\frac{\frac{\frac{\overline{r \leftrightarrow x \vdash x : \llbracket S \rrbracket^\perp, r : \llbracket S \rrbracket}}{r\langle y \rangle.r \leftrightarrow x \vdash y : \llbracket T \rrbracket^\perp, x : \llbracket S \rrbracket^\perp, r : \llbracket T \rrbracket \otimes \llbracket S \rrbracket}}{\bar{r}.r\langle y \rangle.r \leftrightarrow x \vdash y : \llbracket T \rrbracket^\perp, x : \llbracket S \rrbracket^\perp, r : (1 \otimes (\llbracket T \rrbracket \otimes \llbracket S \rrbracket))}}{x(y).\bar{r}.r\langle y \rangle.r \leftrightarrow x \vdash x : (\llbracket T \rrbracket^\perp \wp \llbracket S \rrbracket^\perp), r : (1 \otimes (\llbracket T \rrbracket \otimes \llbracket S \rrbracket))}}{r(x).x(y).\bar{r}.r\langle y \rangle.r \leftrightarrow x \vdash r : (\llbracket T \rrbracket^\perp \wp \llbracket S \rrbracket^\perp) \wp (1 \otimes (\llbracket T \rrbracket \otimes \llbracket S \rrbracket))}$$

– Subcase **wait**: We have **wait** : $\text{end}_? \multimap \mathbf{1}$ where

$$\begin{aligned} \llbracket \text{end}_? \multimap \mathbf{1} \rrbracket^\perp &= \llbracket \text{end}_? \multimap \mathbf{1} \rrbracket \\ &= \llbracket \text{end}_? \rrbracket^\perp \wp (1 \otimes \llbracket \mathbf{1} \rrbracket) \\ &= \perp \wp (1 \otimes \mathbf{1}) \end{aligned}$$

- **Case (absurd V).** We assume $\llbracket V \rrbracket_{x'}^v : \llbracket \Gamma \rrbracket, x' : \mathbf{0}$, and derive:

$$\frac{\frac{x \triangleright \{\} \vdash r : \mathbf{1} \otimes \llbracket T \rrbracket^\perp, x : \top \quad \llbracket V \rrbracket_{x'}^v \vdash \llbracket \Gamma \rrbracket, x' : \mathbf{0}}{x \triangleright \{\} \parallel \llbracket V \rrbracket_{x'}^v \vdash r : \mathbf{1} \otimes \llbracket T \rrbracket^\perp, x : \top \parallel \llbracket \Gamma \rrbracket, x' : \mathbf{0}}}{(\nu x x')(x \triangleright \{\} \parallel \llbracket V \rrbracket_{x'}^v) \vdash \llbracket \Gamma \rrbracket, r : \mathbf{1} \otimes \llbracket T \rrbracket^\perp}$$

- **Case (let $x = M$ in N).**

We assume $\llbracket M \rrbracket_{x'}^m : \llbracket \Gamma \rrbracket, x' : \mathbf{1} \otimes \llbracket T \rrbracket$ and $\llbracket N \rrbracket_r^m : \llbracket \Delta \rrbracket, x : \llbracket T \rrbracket^\perp, r : \llbracket U \rrbracket$ and derive

$$\frac{\frac{\llbracket N \rrbracket_r^m \vdash \llbracket \Delta \rrbracket, x : \llbracket T \rrbracket^\perp, r : \llbracket U \rrbracket}{x. \llbracket N \rrbracket_r^m \vdash \llbracket \Delta \rrbracket, x : \perp \wp \llbracket T \rrbracket^\perp, r : \llbracket U \rrbracket} \quad \llbracket M \rrbracket_{x'}^m \vdash \llbracket \Gamma \rrbracket, x' : \mathbf{1} \otimes \llbracket T \rrbracket}{\frac{x. \llbracket N \rrbracket_r^m \parallel \llbracket M \rrbracket_{x'}^m \vdash \llbracket \Delta \rrbracket, x : \perp \wp \llbracket T \rrbracket^\perp, r : \llbracket U \rrbracket \parallel \llbracket \Gamma \rrbracket, x' : \mathbf{1} \otimes \llbracket T \rrbracket}{(\nu x x')(x. \llbracket N \rrbracket_r^m \parallel \llbracket M \rrbracket_{x'}^m) \vdash \llbracket \Gamma \rrbracket, \llbracket \Delta \rrbracket, r : \llbracket U \rrbracket}}$$

- **Case (V).** We assume $\llbracket V \rrbracket_r^v : \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket$ and derive

$$\frac{\llbracket V \rrbracket_r^v \vdash \llbracket \Gamma \rrbracket, r : \llbracket T \rrbracket}{\bar{r}. \llbracket V \rrbracket_r^v \vdash \llbracket \Gamma \rrbracket, r : \mathbf{1} \otimes \llbracket T \rrbracket}$$

Part 3. The cases are all by immediate induction. $\square$

Lemma C.5. Let F be an HGV^* evaluation context and r a result endpoint. Then there exists a process context $\llbracket F \rrbracket_r^f$ and a result endpoint $v = \text{hr}(F, r)$ for the hole such that for all M we have that $\llbracket F[M] \rrbracket_r^c = \llbracket F \rrbracket_r^f[\llbracket M \rrbracket_v^m]$.

Proof. By induction on the structure of F . $\square$

In the above lemma, if F is the empty context then $v = r$. Otherwise v is a variable bound by the process context $\llbracket F \rrbracket_r^f$.

Lemma C.6 (Operational Correspondence, Terms). *If M is a well-typed term:*

- (1) If $M \rightarrow_M M'$, then there exists a P such that $\llbracket M \rrbracket_r^m \xrightarrow{\beta^+}_\alpha P$ and $P \approx_\alpha \llbracket M' \rrbracket_r^m$; and
- (2) if $\llbracket M \rrbracket_r^m \xrightarrow{\beta}_P$, then there exists an M' and a P' such that $M \rightarrow_M M'$ and $P \xrightarrow{\beta^*}_\alpha P'$ and $P' \approx_\alpha \llbracket M' \rrbracket_r^m$.

Proof.

- (1) By induction on the reduction $M \rightarrow_M M'$.

Case (E-LAM).

$$\begin{array}{ccc} (\lambda x. M) V & \xrightarrow{\rightarrow_M} & M\{V/x\} \\ \downarrow \llbracket \cdot \rrbracket_r^m & & \downarrow \llbracket \cdot \rrbracket_r^m \\ (\nu x x')(\nu y y')(y(x). r \leftrightarrow y \parallel y'(x). \llbracket M \rrbracket_{y'}^m \parallel \llbracket V \rrbracket_{x'}^v) & & \\ \downarrow \xrightarrow{\beta} \alpha & & \\ (\nu x x')(\nu y y')(r \leftrightarrow y \parallel \llbracket M \rrbracket_{y'}^m \parallel \llbracket V \rrbracket_{x'}^v) & & \\ \downarrow \alpha & & \\ (\nu x x')(\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\} \rrbracket_r^m \end{array}$$

Case (E-UNIT).

$$\begin{array}{ccc}
 \text{let } () = () \text{ in } M & \xrightarrow{\rightarrow_M} & M \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x(). \llbracket M \rrbracket_r^m \parallel x'(). \mathbf{0}) & & \\
 \downarrow \xrightarrow{\beta} & & \\
 \llbracket M \rrbracket_r \parallel \mathbf{0} & \xrightarrow{\sim} & \llbracket M \rrbracket_r^m
 \end{array}$$

Case (E-PAIR).

$$\begin{array}{ccc}
 \text{let } (x, y) = (V, W) \text{ in } M & \xrightarrow{\rightarrow_M} & M\{V/x\}\{W/y\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu y y')(y(x). \llbracket M \rrbracket_r^m \parallel y'[x']. (\llbracket V \rrbracket_{x'}^v \parallel \llbracket W \rrbracket_{y'}^v)) & & \\
 \downarrow \xrightarrow{\beta} & & \downarrow \xrightarrow{\beta} \\
 (\nu y y')(\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v \parallel \llbracket W \rrbracket_{y'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\}\{W/y\} \rrbracket_r^m
 \end{array}$$

Case (E-INL).

$$\begin{array}{ccc}
 \text{case inl } V \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} & \xrightarrow{\rightarrow_M} & M\{V/x\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x \triangleright \{ \text{inl} : \llbracket M \rrbracket_r^m; \text{inr} : \llbracket N\{x/y\} \rrbracket_r^m \} \parallel x' \triangleleft \text{inl}. \llbracket V \rrbracket_{x'}^v) & & \\
 \downarrow \xrightarrow{\beta} & & \downarrow \xrightarrow{\beta} \\
 (\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\} \rrbracket_r^m
 \end{array}$$

Case (E-INR). As E-INL.

Case (E-LET).

$$\begin{array}{ccc}
 \text{let } x = V \text{ in } M & \xrightarrow{\rightarrow_M} & M\{V/x\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x. \llbracket M \rrbracket_r^m \parallel \bar{x}'. \llbracket V \rrbracket_{x'}^v) & & \\
 \downarrow \xrightarrow{\beta} \xrightarrow{\beta} & & \downarrow \xrightarrow{\beta} \\
 (\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\} \rrbracket_r^m
 \end{array}$$

Case (E-LIFT). The induction hypothesis yields the reasoning steps depicted by the first diagram, which we use, together with HGV's E-LIFT and HCP's STR-RES and STR-PAR2,

to justify the second diagram:

$$\begin{array}{ccc}
 M & \xrightarrow{\rightarrow_M} & M' \\
 \downarrow \llbracket \cdot \rrbracket_r^m & & \downarrow \llbracket \cdot \rrbracket_r^m \\
 \llbracket M \rrbracket_r^m & \xrightarrow[\alpha \approx \alpha]{\beta^+} & \llbracket M' \rrbracket_r^m
 \end{array}$$

$$\begin{array}{ccc}
 \text{let } x = E[M] \text{ in } N & \xrightarrow{\rightarrow_M} & \text{let } x = E[M'] \text{ in } N \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x. \llbracket N \rrbracket_r^m \parallel \llbracket M \rrbracket_{x'}^m) & \xrightarrow[\alpha \approx \alpha]{\beta^+} & (\nu x x')(x. \llbracket N \rrbracket_r^m \parallel \llbracket M' \rrbracket_{x'}^m)
 \end{array}$$

(2) By induction on M .

Case ($U \ V$). There are two well-typed cases for U : either $U = z$ for some z ; or $U = \lambda x.M$ for some x and M . If $U = z$, we have $(\nu x x')(\nu y y')(y \langle x \rangle. r \leftrightarrow y \parallel z \leftrightarrow y' \parallel \llbracket V \rrbracket_{x'}^v) \not\rightarrow_\beta$, which contradicts our premise. Therefore, $U = \lambda x.M$. The only possible β -transition is the one in the following diagram:

$$\begin{array}{ccc}
 (\lambda x.M) V & \xrightarrow{\rightarrow_M} & M\{V/x\} \\
 \downarrow \llbracket \cdot \rrbracket_r^m & & \downarrow \llbracket \cdot \rrbracket_r^m \\
 (\nu x x')(\nu y y')(y \langle x \rangle. r \leftrightarrow y \parallel y'(x). \llbracket M \rrbracket_{y'}^m \parallel \llbracket V \rrbracket_{x'}^v) & & \\
 \downarrow \xrightarrow[\alpha]{\beta} & & \\
 (\nu x x')(\nu y y')(r \leftrightarrow y \parallel \llbracket M \rrbracket_{y'}^m \parallel \llbracket V \rrbracket_{x'}^v) & & \\
 \downarrow \xrightarrow[\alpha]{} & & \\
 (\nu x x')(\llbracket M \rrbracket_r^m \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow[\approx_\alpha \text{ (by Lemma 5.10)}]{} & \llbracket M\{V/x\} \rrbracket_r^m
 \end{array}$$

Hence, $M' = M\{V/x\}$.

Case ($\text{let } () = U \text{ in } M$). There are two well-typed cases for U : either $U = z$ for some z ; or $U = ()$. If $U = z$, we have $(\nu x x')(x(). \llbracket M \rrbracket_r^m \parallel x' \leftrightarrow z) \not\rightarrow_\beta$, which contradicts our premise. Therefore, $U = ()$. The only possible β -transition is the one in the following diagram:

$$\begin{array}{ccc}
 \text{let } () = () \text{ in } M & \xrightarrow{\rightarrow_M} & M \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x(). \llbracket M \rrbracket_r^m \parallel x'(). 0) & & \\
 \downarrow \xrightarrow[\beta]{} & & \\
 \llbracket M \rrbracket_r \parallel 0 & \xrightarrow{\sim} & \llbracket M \rrbracket_r^m
 \end{array}$$

Hence, $M' = M$.

Case ($\text{let } (x, y) = U \text{ in } M$). There are two well-typed cases for U : either $U = z$ for some z , or $U = (V, W)$. If $U = z$, we have $(\nu y y')(y(x). \llbracket M \rrbracket_r^m \parallel y' \leftrightarrow z) \not\rightarrow_\beta$, which contradicts our premise. Therefore, $U = (V, W)$. The only possible β -transition is the one in the

following diagram:

$$\begin{array}{ccc}
 \text{let } (x, y) = (V, W) \text{ in } M & \xrightarrow{\rightarrow_M} & M\{V/x\}\{W/y\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu y y')(y(x). \llbracket M \rrbracket_r^m \parallel y'[x']. (\llbracket V \rrbracket_{x'}^v \parallel \llbracket W \rrbracket_{y'}^v)) & & \\
 \downarrow \xrightarrow{\beta} & & \\
 (\nu y y')(\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v \parallel \llbracket W \rrbracket_{y'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\}\{W/y\} \rrbracket_r^m
 \end{array}$$

Case (case $U \{ \text{inl } x \mapsto M; \text{inr } x \mapsto N \}$). There are two well-typed cases for U : either $U = z$ for some z ; or $U = \text{inl } V$. If $U = z$, we have $(\nu x x')(x \triangleright \{ \text{inl} : \llbracket M \rrbracket_r^m; \text{inr} : \llbracket N\{x/y\} \rrbracket_r^m \parallel x' \leftrightarrow z \}) \xrightarrow{\beta}$, which contradicts our premise. Therefore, $U = \text{inl } V$. The only possible β -transition is the one in the following diagram:

$$\begin{array}{ccc}
 \text{case inl } V \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} & \xrightarrow{\rightarrow_M} & M\{V/x\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x \triangleright \{ \text{inl} : \llbracket M \rrbracket_r^m; \text{inr} : \llbracket N\{x/y\} \rrbracket_r^m \parallel x' \triangleleft \text{inl}. \llbracket V \rrbracket_{x'}^v \}) & & \\
 \downarrow \xrightarrow{\beta} & & \\
 (\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\} \rrbracket_r^m
 \end{array}$$

Case (absurd U). There is only one well-typed case for U : $U = z$ for some z . However, $(\nu x x')(x \triangleright \{ \} \parallel x' \leftrightarrow z) \xrightarrow{\beta}$, which contradicts our premise.

Case (let $x = M$ in N). There are two possible cases: either $M = V$; or $\llbracket M \rrbracket_{x'}^m \xrightarrow{\beta} P$ for some P . If M is a value, the only possible β -transition is the one in the following diagram:

$$\begin{array}{ccc}
 \text{let } x = V \text{ in } M & \xrightarrow{\rightarrow_M} & M\{V/x\} \\
 \downarrow \llbracket \cdot \rrbracket_r & & \downarrow \llbracket \cdot \rrbracket_r \\
 (\nu x x')(x. \llbracket M \rrbracket_r^m \parallel \bar{x}'. \llbracket V \rrbracket_{x'}^v) & & \\
 \downarrow \xrightarrow{\beta} \xrightarrow{\beta} & & \\
 (\nu x x')(\llbracket M \rrbracket_r \parallel \llbracket V \rrbracket_{x'}^v) & \xrightarrow{\approx_\alpha \text{ (by Lemma 5.10)}} & \llbracket M\{V/x\} \rrbracket_r^m
 \end{array}$$

Otherwise, if $\llbracket M \rrbracket_{x'}^m \xrightarrow{\beta} P$ for some P , the induction hypothesis gives us an M' such that $M \xrightarrow{\rightarrow_M} M'$ and $P \approx \llbracket M' \rrbracket_r^m$. We apply HGV's E-LIFT and HCP's STR-RES and STR-PAR2.

Case (V). We have $\bar{r}. \llbracket V \rrbracket_r^v \xrightarrow{\beta}$, which contradicts our premise. $\square$

Theorem 5.11 (Operational Correspondence). *Suppose $\mathcal{C}$ is a well-typed configuration.*

- (1) (Preservation of reductions) If $\mathcal{C} \longrightarrow \mathcal{C}'$, then there exists a P such that $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\beta^+}_\alpha P$ and $P \approx_\alpha \llbracket \mathcal{C}' \rrbracket_r^c$; and
- (2) (Reflection of transitions)
 - if $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\alpha} P$, then $P \approx_\alpha \llbracket \mathcal{C} \rrbracket_r^c$; and

- if $\llbracket \mathcal{C} \rrbracket_r^c \xrightarrow{\beta} P$, then there exists a $\mathcal{C}'$ and a P' such that $\mathcal{C} \longrightarrow \mathcal{C}'$ and $P \xrightarrow{\beta^*}_{\alpha} P'$ and $P' \approx_{\alpha} \llbracket \mathcal{C}' \rrbracket_r^c$. Furthermore, $\mathcal{C}'$ is unique up to structural congruence.

Proof.

- (1) By induction on the reduction $\mathcal{C} \longrightarrow \mathcal{C}'$. We implicitly make use of Lemma C.5 throughout the proof in order to recast the translation of a plugged evaluation context $\llbracket F[M] \rrbracket_r^c$ into the plugging of the translated evaluation context with the translation of the plugged term $\llbracket F \rrbracket_r^f[\llbracket M \rrbracket_v^m]$ where $v = \text{hr}(F, r)$.

Case (E-REIFY-FORK).

$$\begin{array}{ccc}
F[\text{fork } V] & \xrightarrow{\quad} & (\nu x x')(F[x] \parallel \circ V x') \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket F \rrbracket_r^f[(\nu z z')(\nu y y') \left(\begin{array}{c} y\langle z \rangle.v \leftrightarrow y \parallel \\ (\nu x x') \left(\begin{array}{c} y'(w).x\langle w \rangle.\bar{y}'.y' \leftrightarrow x \parallel \\ x'(w).w\langle x' \rangle.w.w \parallel \mathbf{0} \end{array} \right) \parallel \end{array} \right) \parallel \llbracket V \rrbracket_{z'}^v] & & \\
\downarrow \xrightarrow{\beta} \xrightarrow{\beta} \xrightarrow{\alpha} & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket F \rrbracket_r^f[(\nu z z') \left(\begin{array}{c} (\nu x x') \left(\begin{array}{c} \bar{v}.v \leftrightarrow x \parallel \\ z\langle x' \rangle.z.z \parallel \mathbf{0} \end{array} \right) \parallel \\ \llbracket V \rrbracket_{z'}^v \end{array} \right) \parallel] & \xrightarrow{\approx_{\alpha}} & (\nu x x') \left(\begin{array}{c} \llbracket F \rrbracket_r^f[\bar{v}.v \leftrightarrow x] \parallel \\ (\nu y y') \left(\begin{array}{c} (\nu w w')(\nu z z') \left(\begin{array}{c} z\langle w \rangle.y \leftrightarrow z \parallel \\ w' \leftrightarrow x' \end{array} \right) \parallel \\ y'.y' \parallel \mathbf{0} \end{array} \right) \parallel \end{array} \right)
\end{array}$$

The endpoint $v = \text{hr}(F, r)$. The final two terms are bisimilar by Lemma C.4.

Case (E-REIFY-LINK).

$$\begin{array}{ccc}
\circ E[\text{link } (x, y)] & \xrightarrow{\quad} & (\nu z z')(x \xleftrightarrow{z} y \parallel \circ E[z']) \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu a a')(\llbracket E \rrbracket_r^m[(\nu z z')(\nu w w')(w\langle z \rangle.v \leftrightarrow w \parallel w'(t).t(s).\bar{w}'.w'().s \leftrightarrow t \parallel z'\langle x \rangle.y \leftrightarrow z' \parallel \bar{a}'.a' \parallel \mathbf{0})]) & & \\
\downarrow \xrightarrow{\beta} \xrightarrow{\beta} \xrightarrow{\alpha} \xrightarrow{\alpha} & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu a a')(\llbracket E \rrbracket_r^m[\bar{v}.v().x \leftrightarrow y \parallel \bar{a}'.a' \parallel \mathbf{0}]) & \xrightarrow{\approx_{\alpha}} & (\nu z z')(\bar{z}.z().x \leftrightarrow y \parallel (\nu a a')(\llbracket E[v \leftrightarrow z'] \rrbracket_a^m \parallel \bar{a}'.a' \parallel \mathbf{0}))
\end{array}$$

The endpoint $v = \text{hr}(F, r)$.

Case (E-COMM-LINK).

$$\begin{array}{ccc}
(\nu z z')(\nu x x')(x \xleftrightarrow{z} y \parallel \circ z' \parallel \phi M) & \xrightarrow{\quad} & \phi(M\{y/x'\}) \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu z z')(\nu x x')(\bar{z}.z().x \leftrightarrow y \parallel (\nu w w')(z' \leftrightarrow w \parallel w'.w' \parallel \mathbf{0}) \parallel \llbracket \phi M \rrbracket_r^c) & & \\
\downarrow \xrightarrow{\alpha} \xrightarrow{\beta} \times 3 & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket \phi M \rrbracket_r^c\{y/x'\} & \xrightarrow{\approx_{\alpha}} & \llbracket \phi M\{y/x'\} \rrbracket_r^c
\end{array}$$

Case (E-COMM-SEND).

$$\begin{array}{ccc}
(\nu xx')(F[\mathbf{send}(V, x)] \parallel F'[\mathbf{recv} x']) & \xrightarrow{\quad} & (\nu xx')(F[x] \parallel F'[(V, x')]) \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu xx') \left(\begin{array}{l} \llbracket F \rrbracket_r^f[(\nu yy')(\nu zz')(z\langle y \rangle.u \leftrightarrow z \parallel z'(t).t\langle s \rangle.\bar{z}'.z' \leftrightarrow t \parallel y'[w].(\llbracket V \rrbracket_w^v \parallel x \leftrightarrow y')) \parallel \\ \llbracket F' \rrbracket_r^c[(\nu yy')(\nu zz')(z\langle y \rangle.v \leftrightarrow z \parallel z'(s).s\langle t \rangle.\bar{z}'.z' \leftrightarrow s \parallel x' \leftrightarrow y')] \end{array} \right) & & \\
\downarrow \xrightarrow{\beta \times 5, \alpha \times 2} & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu xx')(\llbracket F \rrbracket_r^f[(x\langle w \rangle.\bar{u}.x \leftrightarrow u \parallel \llbracket V \rrbracket_w^v)] \parallel \llbracket F' \rrbracket_r^c[x'(t).\bar{v}.v \langle t \rangle.v \leftrightarrow x']) & & (\nu xx')(\llbracket F \rrbracket_r^f[\bar{u}.x \leftrightarrow u] \parallel \llbracket F' \rrbracket_r^c[\bar{v}.v \langle w \rangle.v \leftrightarrow x']) \\
\downarrow \beta & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu xx')(\llbracket F \rrbracket_r^f[(\bar{u}.u \leftrightarrow x \parallel \llbracket V \rrbracket_w^v)] \parallel \llbracket F' \rrbracket_r^c[\bar{v}.v \langle w \rangle.v \leftrightarrow x']) & \xrightarrow{\approx_\alpha} & (\nu xx')(\llbracket F \rrbracket_r^f[\bar{u}.u \leftrightarrow x \parallel \llbracket V \rrbracket_w^v] \parallel \llbracket F' \rrbracket_r^c[\bar{v}.v \langle w \rangle.v \leftrightarrow x'])
\end{array}$$

The endpoint $u = \mathbf{hr}(F, r)$ and the endpoint $v = \mathbf{hr}(F', r)$.

Case (E-COMM-CLOSE).

$$\begin{array}{ccc}
(\nu xx)(o \ x \parallel F[\mathbf{wait} x']) & \xrightarrow{\quad} & F[()] \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu xx) \left(\begin{array}{l} (\nu yy')(\bar{y}.x \leftrightarrow y \parallel y'.y' \square. \mathbf{0}) \parallel \\ \llbracket F \rrbracket_r^f[(\nu zz')(\nu ww')(w\langle z \rangle.v \leftrightarrow w \parallel w'(s).s\langle \rangle.\bar{w}'.w' \square. \mathbf{0} \parallel x' \leftrightarrow z')] \end{array} \right) & & \\
\downarrow \xrightarrow{\beta \times 3, \alpha \times 3, \beta} & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket F \rrbracket_r^f[\bar{v}.v \square. \mathbf{0}] & \xrightarrow{=} & \llbracket F \rrbracket_r^f[\bar{v}.v \square. \mathbf{0}]
\end{array}$$

The endpoint $v = \mathbf{hr}(F, r)$.

Case (E-RES).

$$\begin{array}{ccc}
(\nu xy)\mathcal{C} & \xrightarrow{\quad} & (\nu xy)\mathcal{C}' \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
(\nu xy)\llbracket \mathcal{C} \rrbracket_r^c & \xrightarrow{\beta^+ \approx_\alpha \approx_\alpha (\text{IH})} & (\nu xy)\llbracket \mathcal{C}' \rrbracket_r^c
\end{array}$$

Case (E-PAR).

$$\begin{array}{ccc}
\mathcal{C} \parallel \mathcal{D} & \xrightarrow{\quad} & \mathcal{C}' \parallel \mathcal{D} \\
\downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket \mathcal{C} \rrbracket_r^c \parallel \llbracket \mathcal{D} \rrbracket_r^c & & \llbracket \mathcal{C}' \rrbracket_r^c \parallel \llbracket \mathcal{D} \rrbracket_r^c \\
\downarrow \xrightarrow{\beta^+ \approx_\alpha \approx_\alpha (\text{IH})} & & \downarrow \llbracket \cdot \rrbracket_r^c \\
\llbracket \mathcal{C}' \rrbracket_r^c \parallel \llbracket \mathcal{D} \rrbracket_r^c & \xrightarrow{=} & \llbracket \mathcal{C}' \rrbracket_r^c \parallel \llbracket \mathcal{D} \rrbracket_r^c
\end{array}$$

Case (E-EQUIV).

$$\begin{array}{ccccccc}
 \mathcal{C} & \xrightarrow{\equiv} & \mathcal{C}' & \xrightarrow{\longrightarrow} & \mathcal{D}' & \xrightarrow{\equiv} & \mathcal{E} \\
 \downarrow \llbracket \cdot \rrbracket_r^c & & & & & & \downarrow \llbracket \cdot \rrbracket_r^c \\
 \llbracket \mathcal{C} \rrbracket_r^c & & & & & & \\
 \downarrow \approx_\alpha \text{ (Lemma 5.6)} & & & & & & \\
 \llbracket \mathcal{C}' \rrbracket_r^c & & & & & & \\
 \downarrow \xRightarrow{\beta^+} \alpha \approx_\alpha \text{ (IH)} & & & & & & \\
 \llbracket \mathcal{D}' \rrbracket_r^c & \xrightarrow{\approx_\alpha \text{ (Lemma 5.6)}} & & & & & \llbracket \mathcal{E} \rrbracket_r^c
 \end{array}$$

Case (E-LIFT-M). The cases for $\phi = \bullet$ and $\phi = \circ$ are similar; here we show the case for $\bullet$.

$$\begin{array}{ccc}
 \bullet M & \xrightarrow{\longrightarrow} & \bullet N \\
 \downarrow \llbracket \cdot \rrbracket_r^c & & \downarrow \llbracket \cdot \rrbracket_r^c \\
 \llbracket M \rrbracket_r^m & \xrightarrow{\xRightarrow{\beta^+} \alpha \approx_\alpha \text{ (Lemma C.6)}} & \llbracket N \rrbracket_r^m
 \end{array}$$

- (2) Reflection of α -transitions is trivial as α -transition is included in α -bisimulation. Reflection of β -transitions is by induction on $\mathcal{C}$; as with Lemma C.6, the only well-typed β -transitions that can occur for each case are those specified in the simulation case. $\square$

4.3 Discussion

This section proceeds as follows:

- In § 4.3.1, I present an alternative formulation of GV that uses cuts, rather than lock typing, and argue this significantly simplifies the metatheory of GV.
- In § 4.3.2, I discuss the relation between GV with cuts, lock typing, and hypersequents.

4.3.1 Cooking GV With Cut

In Lindley and Morris' GV, name restriction and parallel composition are separate configuration constructs. However, as discussed, this breaks type preservation for the structural congruence, which significantly complicates the metatheory.

This section presents an alternative formulation of GV, which we refer to as “GV with Cut” or GV_{Cut} , which more closely corresponds to CP. While this version is not preferable to Hypersequent GV, presented in this chapter, it represents my view of how GV should have been formulated prior to the introduction of hypersequents, and is helpful when examining Lindley and Morris' lock typing.

In this section, the terms and types of GV_{Cut} are printed in *pink* and *green*, respectively, and both are rendered in a sans-serif font, and any relations, such as typing and reduction, are marked by a subscript “GVC”.

To construct GV_{Cut} , we replace name restriction and parallel composition with a cut. For configuration typing, we replace the rules TG-NEW, TG-CONNECT₁, and TG-CONNECT₂ with the rule TG-CUT.

$$\begin{array}{c}
 C, D ::= (\nu x \bar{x})C \\
 \quad | C \parallel D \\
 \quad | (\nu x \bar{x})(C \parallel D)
 \end{array}
 \quad
 \text{TG-CUT}
 \quad
 \frac{\Gamma, x : S \vdash_{\text{GVC}} C : R \quad \Delta, \bar{x} : \bar{S} \vdash_{\text{GVC}} D : R'}{\Gamma, \Delta \vdash_{\text{GVC}} (\nu x \bar{x})(C \parallel D) : R \sqcap R'}$$

For the structural congruence, we replace the rules SC-PARASSOC, SC-PARCOMM, SC-NEWASSOC, SC-NEUSWAP, and SC-SCOPEEXT with the rules SC-CUTCOMM and SC-CUTASSOC.

$$\begin{array}{ll}
 (\nu x \bar{x})(C_1 \parallel C_2) & \equiv_{\text{GVC}} (\nu \bar{x} x)(C_2 \parallel C_1) & \text{SC-CUTCOMM} \\
 (\nu x \bar{x})(\nu y \bar{y})(C_1 \parallel C_2) \parallel C_3 & \equiv_{\text{GVC}} (\nu y \bar{y})(\nu x \bar{x})(C_1 \parallel C_3) \parallel C_2 & \text{SC-CUTASSOC} \\
 & \text{where } x \notin C_2 \text{ and } y \notin C_3 &
 \end{array}$$

In GV_{Cut} , the structural congruence preserves types. Much like the rules themselves, the proof is nearly identical to the proof of preservation for CP's structural congruence, and is left as an exercise.

Lemma 4.1. *If $\Gamma \vdash_{\text{GVC}} C : R$ and $C \equiv_{\text{GVC}} D$, then $\Gamma \vdash_{\text{GVC}} D : R$.*

Since GV_{cut} 's structural congruence preserves types, it can be harmlessly embedded into the reduction relation, as per EG-EQUIV.

$$\text{EG-EQUIV} \quad \frac{C \equiv_{\text{GVC}} C' \quad C' \longrightarrow_{\text{GVC}} D' \quad D' \equiv_{\text{GVC}} D}{C \longrightarrow_{\text{GVC}} D}$$

This small change tightens GV_{cut} 's correspondence with CP and considerably simplifies its metatheory, as there is no longer any need to prove preservation up to equivalence [e.g. Fowler, 2019, Theorem 2, pp. 40–43].

4.3.2 Lock Types, Cuts, and Hypersequents

This section discusses the differences between GV with lock types, cuts, and hypersequents, and argues that, while lock types and hypersequents both permit the syntax for configurations to separate name restriction from parallel composition, lock types are, in effect, more complicated cuts, whereas hypersequents capture parallelism.

We compare the configuration typing rules for GV (Figure I.5), GV_{cut} (§ 4.3.1), and HGV (Figure I.3). In this section, the terms and types of both GV and GV_{cut} are printed in *pink* and *green*, respectively, both are rendered in a sans-serif font, and any relations, such as typing and reduction, are marked by subscript “GV” and “GVC”, respectively.

Let us begin with an observation. Both GV and HGV use the exact same syntax for configurations, but the translation from HGV to GV (in § I.4, under “Translating HGV to GV”) shows a gap between the two systems. Every GV configuration is typeable in HGV, but there are HGV configurations that are not typeable in GV.

What is the relation between GV's configurations and GV_{cut} 's configurations? As we will see, GV's parallel composition is, in essence, a cut, and therefore the two are identical.

Let us begin by discussing GV's lock types in slightly more detail. GV splits cut into name restriction and parallel composition:

$$\begin{array}{lcl} C, D ::= & \dots & \\ & | (\nu x\bar{x})(C \parallel D) & \text{cut} \\ & | (\nu x\bar{x})C & \text{name restriction} \\ & | C \parallel D & \text{parallel composition} \end{array}$$

As naively decomposing cuts would break GV's tree connection structure and, more importantly, deadlock freedom, GV ensures these properties by lock typing. It extends runtime typing environments with locked channel

type assignments.³

$$\Gamma ::= \dots \mid \Gamma, \langle x, \bar{x} \rangle : S^\#$$

The special type assignment “ $\langle x, \bar{x} \rangle : S^\#$ ” represents a channel with endpoints x and $\bar{x}$ which has been created, but has not yet been split across a parallel composition. Hence, it is *locked*.

When a locked channel is split across a parallel composition, it becomes unlocked, and the endpoints become available for use. Each parallel composition must split exactly one locked channel. This is guaranteed by the typing rules for name restriction and parallel composition, reproduced below from Figure I.5 (eliding the superfluous TG-CONNECT₂).

$$\begin{array}{c} \text{TG-NEW} \\ \frac{\Gamma, \langle x, \bar{x} \rangle : S^\# \vdash_{\text{GV}} C : R}{\Gamma \vdash_{\text{GV}} (\nu x \bar{x})C : R} \end{array} \quad \begin{array}{c} \text{TG-CONNECT}_1 \\ \frac{\Gamma, x : S \vdash_{\text{GV}} C : R \quad \Delta, \bar{x} : \bar{S} \vdash_{\text{GV}} D : R'}{\Gamma, \Delta, \langle x, \bar{x} \rangle : S^\# \vdash_{\text{GV}} C \parallel D : R \sqcap R'} \end{array}$$

The rule TG-NEW creates a locked channel. The rule TG-CONNECT₁ splits a locked channel across a parallel composition. As the typing environments for term typing cannot contain locked channels, every locked channel must be split across some parallel composition.

There is a tension between the structural congruence and the typing rules. The typing rules require that each parallel composition splits exactly one locked channel, but the structural congruence—specifically, SC-PARASSOC—does not respect this invariant. Hence, GV’s structural congruence does not preserve types. For instance, the following configuration is well-typed:

$$\frac{\frac{\frac{\Gamma_1, x : S \vdash_{\text{GV}} C_1 : R_1 \quad \frac{\frac{\Gamma_2, y : S' \vdash_{\text{GV}} C_2 : R_2 \quad \Gamma_3, \bar{x} : \bar{S}, \bar{y} : \bar{S}' \vdash_{\text{GV}} C_3 : R_3}{\Gamma_2, \Gamma_3, \bar{x} : \bar{S}, \langle y, \bar{y} \rangle : S' \vdash_{\text{GV}} C_2 \parallel C_3 : R_2 \sqcap R_3}}{\Gamma_1, \Gamma_2, \Gamma_3, \langle x, \bar{x} \rangle : S, \langle y, \bar{y} \rangle : S' \vdash_{\text{GV}} C_1 \parallel (C_2 \parallel C_3) : R_1 \sqcap R_2 \sqcap R_3}}{\Gamma_1, \Gamma_2, \Gamma_3, \langle x, \bar{x} \rangle : S \vdash_{\text{GV}} (\nu y \bar{y})(C_1 \parallel (C_2 \parallel C_3)) : R_1 \sqcap R_2 \sqcap R_3}}{\Gamma_1, \Gamma_2, \Gamma_3 \vdash_{\text{GV}} (\nu x \bar{x})(\nu y \bar{y})(C_1 \parallel (C_2 \parallel C_3)) : R_1 \sqcap R_2 \sqcap R_3}$$

However, by SC-PARASSOC

$$(\nu x \bar{x})(\nu y \bar{y})(C_1 \parallel (C_2 \parallel C_3)) \equiv_{\text{GV}} (\nu x \bar{x})(\nu y \bar{y})((C_1 \parallel C_2) \parallel C_3)$$

The configuration on the right-hand side is not typeable, as the left-most parallel composition splits *no* locked channels, and the right-most parallel composition splits *two* locked channels.

Consequently, GV, as well as a significant portion of work based on Lindley and Morris’ GV, proves some variation of the following

³Technically, the presentation of GV in § 4.2 adds $S^\#$ as a runtime session type, rather than as part of a special type assignment. Nonetheless, the two presentations are equivalent. Lock types cannot occur in user programs, and so cannot occur as part of any other type. Hence, they already occur only as the top-most connective in a type assignment.

proposition, which states that if we break typing before some reduction, we can restore it afterwards.

Proposition 4.2. *If $\Gamma \vdash_{GV} C : R$ and $C \equiv_{GV} C'$ and $C' \longrightarrow_{GV} D'$, then there exists some D such that $D' \equiv_{GV} D$ and $\Gamma \vdash_{GV} D : R$.*

Proof. See Fowler [2019, pp. 41-43, Theorem 2]. Fowler’s proof is for a variant of GV *without* link or link threads, but the only rule of structural congruence that breaks typing is SC-PARASSOC. Hence, the proof of Theorem 2 trivially extends to the system *with* link and link threads. Fowler’s proof is stronger, as it restores typing before and after the reduction. □

GV’s parallel composition is *covertly* a cut, albeit one that leaves the channel *implicit*.⁴ Like cut, parallel composition must split exactly one channel, and, while cut restricts a channel by removing its endpoints from the typing environment, parallel composition restricts a channel by locking its endpoints. This is just as effective, since there is no way to use the endpoints of a locked channel. The locking mechanism keeps the endpoints around for the sole purpose of being removed by name restriction, which only adds unnecessary complexity.

To formalise this fact, we define a translation on configuration typing derivations from GV to GV_{cut} , written $\llbracket \cdot \rrbracket_{GV_{cut}}$, which removes all lock types from typing environments, removes all name restrictions, and replaces parallel compositions with cuts. All uses of the rule TG-NEW are removed:

$$\left\llbracket \frac{\delta}{\frac{\Gamma, \langle x, \bar{x} \rangle : S^\# \vdash_{GV} C : R}{\Gamma \vdash_{GV} (vx\bar{x})C : R} \text{TG-NEW}} \right\llbracket_{GV_{cut}} \triangleq \frac{\llbracket \delta \rrbracket_{GV_{cut}}}{\Gamma \vdash_{GV_{cut}} C : R}$$

Each use of the rule TG-CONNECT₁ is replaced with a use of the rule TG-CUT:

$$\begin{aligned} & \left\llbracket \frac{\delta_1 \quad \delta_2}{\frac{\Gamma, x : S \vdash_{GV} C : R \quad \Delta, \bar{x} : \bar{S} \vdash_{GV} D : R'}{\Gamma, \Delta, \langle x, \bar{x} \rangle : S^\# \vdash_{GV} C \parallel D : R \sqcap R'} \text{TG-CONNECT}_1} \right\llbracket_{GV_{cut}} \\ & \quad \mapsto \\ & \frac{\llbracket \delta_1 \rrbracket_{GV_{cut}} \quad \llbracket \delta_2 \rrbracket_{GV_{cut}}}{\frac{\Gamma, x : S \vdash_{GV_{cut}} C : R \quad \Delta, \bar{x} : \bar{S} \vdash_{GV_{cut}} D : R'}{\Gamma, \Delta \vdash_{GV_{cut}} (vx\bar{x})(C \parallel D) : R \sqcap R'} \text{TG-CUT}} \end{aligned}$$

The translation acts as the identity on threads, terms, and types. We believe that the translation preserves reduction, but proving this matter

⁴Tellingly, Lindley and Morris [2015, § 3.3] write “ $\parallel_x$ ” when they need the channel name.

would be tedious and not worthwhile. The intuition is that, under the translation $\llbracket \cdot \rrbracket_{\text{HVC}}$, any use of SC-LINKCOMM is preserved, any use of SC-PARCOMM becomes SC-CUTCOMM, and all uses of SC-NEWASSOC, SC-NEWSWAP, or SC-SCOPEEXT become reflexivity. The tedious part is the proof that all type breaking uses of SC-PARASSOC are repaired in such a way that the source and target can be proven equivalent using SC-CUTASSOC. The structure of this proof depends heavily on the computational content of the proof that repairs ill-typed configurations.

4.4 Conclusion

In this chapter, we introduced Hypersequent GV, as well as GV, with their typing rules, reduction semantics, and their metatheory. Hypersequent GV rectifies a design flaw that complicates all previous metatheory of GV by dropping lock typing in favour of hypersequents. For HGV, we proved *preservation* (Theorem I.3.3), the *tree-structure* of connections in configurations (Theorem I.3.14), *global progress* (Theorem I.3.20), the *diamond property* (Theorem I.3.21), and *termination* (Theorem I.3.22). We related HGV to GV by means of a translation from GV to HGV (Theorem I.4.3) and a translation from HGV to GV (Corollary I.4.7). We related HGV to HCP by means of two translations, which translate HGV to HCP via fine-grain call-by-value HGV, and proved that the latter translation preserves types (Lemma I.5.9), and that it gives rise to a sound and complete operational correspondence (Theorem I.5.11). Finally, I introduced a variant of GV which uses cuts, instead of lock typing, and compared Hypersequent GV to GV with lock typing and GV with cuts.

In the future, it would be interesting to extend Hypersequent GV with the extensions for GV described in previous work, such as unlimited types [Lindley and Morris, 2015], fixed points [Lindley and Morris, 2016b], polymorphism [Lindley and Morris, 2017], and exceptions [Fowler et al., 2019]. Furthermore, it would be interesting to describe the construction of GV compositionally, as the combination of a process calculus and a λ -calculus.

Part III

Prioritise The Best Variation

Chapter 5

Priority Classical Processes & Priority Good Variation

This chapter presents both Priority CP (PCP) and Priority GV (PGV), variants of CP and GV that use priorities [Kobayashi, 2006, Padovani, 2014]. PCP is a session-typed process calculus, based on CP, which was introduced by Dardha and Gay [2018]. PGV is a session-typed functional language, based on GV, which was introduced by Kokke and Dardha [2021a].

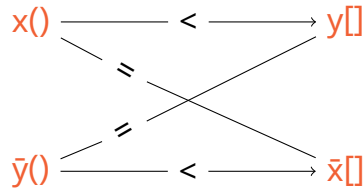
Priorities are an extra-logical mechanism for guaranteeing deadlock freedom, which arose from research on typed π -calculus rather than logic. In a system using priorities, each session type constructor—or, equivalently each communication action—is annotated with some priority, which the type system uses to ensure deadlock freedom. Let us consider two intuitions for the mechanism by which priorities work.

Concretely, *priorities* are natural numbers that represent the time at which some communication happens, i.e. for the process $x[y].z(w).P$, we might assign the action $x[y]$ priority 1, because it happens at time 1, and $z(w)$ priority 2, since it necessarily happens at some later time. Type checking requires that the priority assignment reflects this dependency, and that dual actions take place at the same time. In a session-typed system, this suffices to rule out deadlocks, as there are no valid priority assignments for deadlocking processes. For instance, for the process

$$(vx\bar{x})(vy\bar{y})(x().y[].0 \parallel \bar{y}().\bar{x}[],0)$$

there exists no valid priority assignment, since, by duality, $x()$ must happen at the same time as $\bar{x}[]$, and $y()$ must happen at the same time as $\bar{y}[]$, but, by dependency, $x()$ must happen before $y[]$, and $\bar{y}()$ must happen before $\bar{x}[]$. Represented visually, we can see that these requirements

are essentially cyclic:



Programs are annotated with a lower and an upper *priority bound*, which represent the time at which the program begins to communicate, and the time it finishes. Each priority bound is either a concrete priority, or $\top$, or $\perp$, where the latter two are the supremum and infimum, respectively, i.e. $\top$ is greater than any priority, and $\perp$ is smaller than any priority. The supremum and infimum are useful in a variety of circumstances. For instance, a program that does not communicate is assigned the lower bound $\top$ and the upper bound $\perp$, i.e. it never begins communicating, and has already finished. Alternatively, $\perp$ and $\top$ may be used as lower and upper bounds, respectively, to represent unknowns. For instance, a program that loops might not have a concrete upper bound, and can be assigned upper bound $\top$, to mean that we do not know if or when it finishes communicating.

Abstractly, we can view priority metavariables as names for actions. (Here, we say priority metavariables, rather than priorities, to emphasize that we are referring to the names, rather than the concrete numbers with which they will be instantiated.) Type checking imposes equality and inequality constraints on priority metavariables, which define the deep dependency graph by analogy to the shallow dependency graph as defined for, e.g. HCP in § 3.2.4. For instance, the visual representation of constraints, shown above, is isomorphic to the dependency graph for the same process. If the deep dependency graph is essentially acyclic, there exists a topological sort of the priority metavariables, which we can use to instantiate them with concrete priorities. This is the mechanism behind priority inference, which we discuss in § 5.3.2.

Generally, priority type systems place the priority annotations on session type connectives, rather than actions, since it makes it easier to integrate priorities into type checking. (This is what the type systems for PCP and PGV, presented in § 5.2, do as well.) Explicitly tracking both the priority lower and upper bounds can be quite syntactically burdensome, but there are several tricks to ease this:

- In general, the lower bound can be under-approximated by taking the smallest priority of all endpoints in the typing environment, since a program cannot start communicating earlier than its endpoints permit.
- In systems such as the π -calculus, we do not need to track the upper bound. When composing two programs in sequence, we must check

that the first program finishes communicating before the second program starts. However, the only kind of sequential composition in the π -calculus is prefixing a process with an action. In this case, we must check that the action finishes before the process starts. Hence, we only need to check that the priority of the action is smaller than the lower bound of the process.

- Finally, if we must track both the lower and the upper bound, we can slightly ease the burden of syntax by tracking pairs of lower and upper bounds instead. Since pairs of priority bounds form a lattice, this lets us track both bounds while only manipulating a single value.

One of the benefits of GV is that it closely resembles the manner in which concurrency is exposed to the user in programming languages, which makes it easy to implement the deadlock-free session-typed communication offered by GV as a library, especially in a host language that supports linear types. Less so for Priority GV, as presented in § 5.2, since priority typing requires more of the type system. Fortunately, session-typed communication can be reflected into a monad, in the sense of the *monadic reflection* of an effect, following Filinski [1994], and priority typing can be reflected into a parametrised monad, parametrised by pairs of lower and upper priority bounds, following Atkey [2009]. We explore the monadic reflection of priorities in § 6.2.

The version of PCP introduced in this chapter is different from the version introduced by Dardha and Gay [2018].

- We drop the commuting conversions. Hence, it bears the same relation to Dardha and Gay’s PCP as the version of CP in Chapter 2 does to Wadler’s CP. Notably, Dardha and Gay’s PCP is non-confluent for the same reason Wadler’s CP is non-confluent (see § 2.3).
- We simplify variant select and offer to binary select and offer, and reintroduce a separate construct for the absurd offer.
- We allow arbitrary process continuations in the close construct, rather than forcing the process to terminate.
- We restrict ourselves to studying the multiplicative additive core, and omit the exponentials, as is done in most of this thesis. (Neither publication includes the second order quantifiers.)
- We add the additive units, which are omitted in Dardha and Gay [2018], and use the ‘inert’ semantics for the absurd offer, as discussed in § 2.1.6. (Since the priority constraints encode the blocking behaviour of actions, using the exceptional semantics, as discussed in § 3.3.1, requires a different typing rule for the absurd offer.)

While Dardha and Gay [2018] introduce PCP as an *extension* of CP, it does not formally extend CP, as there are CP processes that are not typeable in PCP. We discuss the relation between CP and PCP in detail in § 5.3.1.

The bulk of the chapter consists of the paper *Prioritise the Best Variation* by Kokke and Dardha [2023], hereafter referred to as Paper II. References made from the main body of this thesis into Paper II will be prefixed by an “II”, e.g. “Theorem II.3.13”. This chapter proceeds as follows:

- In § 5.1, we provide a legend and an errata for Paper II.
- In § 5.2, we present Priority GV and Priority CP, their metatheories, and the correspondence between Priority GV and Priority CP. This section consists entirely of Paper II, and proceeds as follows:

- In § II.2, we introduce PGV.
- In § II.3, we present the metatheory for PGV.

Notably, we prove *preservation* (referred to as *subject reduction*, § II.3.1, Theorem II.3.5) and *global progress* (§ II.3.2, Theorem II.3.14).

- In § II.4, we revisit PCP and its metatheory, and prove a sound and complete operational correspondence between PCP and PGV.

The version of PCP presented in this section differs from Dardha and Gay’s PCP, as discussed previously, as it drops commuting conversions, allows arbitrary process continuations after a close action, and does not include the exponentials. Notably, we prove *progress* (§ II.4.5, Theorem II.4.4), define a translation from PCP into PGV (§ II.4.6), and prove that the translation preserves types (§ II.4.6, Theorem II.4.6), and gives rise to a complete (§ II.4.6, Theorem II.4.7) and sound (§ II.4.6, Theorem II.4.10) operational correspondence.

- In § II.5, we discuss an example program.
- In § II.6, we discuss related work.
- In § 5.3, we discuss PCP in further detail:
 - In § 5.3.1, we discuss the relation between PCP and CP.
 - In § 5.3.2, we introduce priority inference for PCP.

5.1 Legend and Errata

The conventions and terminology in Paper II are different from those used in the rest of this thesis.

- The terms, types, and priorities of *both* Priority CP and Priority GV are printed in **red**, **blue**, and **green**, respectively, and are rendered in an italicised or bolded font with serif.

- The names for the rules of structural congruence, reduction, and the type system differ slightly from those used for HGV in Chapter 4 and for HCP in Chapter 3.
- The following changes relate to the presentation of PCP:
 - Duality is written as $A^\perp$ instead of $\bar{A}$.
 - The absurd offer is written as $x \triangleright \{\}$ instead of $x \not\leq N$.
 - Reduction is not definitionally closed under evaluation contexts, as by E-CONG. Instead, there are separate congruence rules for name restriction and parallel composition (E-LIFTRES and E-LIFTPAR).
 - A *ready* process is called an “action” and a process *ready to act on* some endpoint is said to “act on” that endpoint (§ II.4.5, Definition II.4.1).
 - There is no proof of Progress analogous to Proposition 3.35. Instead, the paper proves Closed Progress (§ II.4.5, Theorem II.4.4).
 - There is no definition of *canonical form* analogous to Definition 3.33, which includes both normal and neutral forms. (The latter is not needed for Closed Progress. Hence, $\mathbf{0}$ suffices.)
Unfortunately, the paper uses the name “canonical form” to refer to a form that serves the same purpose as the *maximum configuration context* or *right-branching form*, i.e. the form

$$(vx_1 \bar{x}_1) \cdots (vx_n \bar{x}_n)(P_1 \parallel \cdots \parallel P_m)$$

where each process P_i (for $1 \leq i \leq m$) is ready (§ II.4.5, Definition II.4.2).

- The following changes relate to the presentation of PGV:
 - Configuration typing does not use configuration types, as in Chapter 4, but omits the type of the value returned by the configuration, as, e.g. in Lindley and Morris [2015] (see Figure II.2).

There are minor errors in Paper II:

- The text states that “Priority GV is more flexible than GV” and “more expressive than GV” (§ II.2), which is not true, as we discuss in § 5.3.1.
- The text states that “PLL [Priority Linear Logic] is an extension of CLL” (§ II.4.4), which—when taken as a statement about proofs—is not true, as we discuss in § 5.3.1. The next sentence, which states that “PCP is not in correspondence with CLL” implies the previous statement concerns provability, rather than proofs, and while this is not known to be false, it remains unproven.
- The definition of *normal form* (§ II.3.2, Definition I.3.11) reads “a configuration $\mathcal{C}$ is in normal form if it is of the form

$$(vx_1 x'_1) \cdots (vx_n x'_n)(\circ M_1 \parallel \cdots \parallel \circ M_m \parallel \bullet V)$$

where each M_i is ready to act on x_i .” Instead, the final part of the phrase should read “where each M_i (for $1 \leq i \leq m$) is ready to act on some x_j (for $1 \leq j \leq n$) or some free channel endpoint.”

- *Closed Progress* (Theorem II.4.4) is stated using equality rather than structural congruence. The correct statement is:

If $P \vdash \emptyset$, then either $P \equiv Q$ or there exists a Q such that $P \Longrightarrow Q$

- In the proof for Lemma II.4.8, the cases for $x \triangleleft \text{inl}.P$ and $x \triangleleft \text{inr}.P$ erroneously list the premise and conclusion of the reduction as $\langle x[y].P \rangle_M$ and $\langle x[y].P \rangle_C$, rather than $\langle x \triangleleft \text{inl}.P \rangle_M$ and $\langle x \triangleleft \text{inl}.P \rangle_C$, and $\langle x \triangleleft \text{inl}.P \rangle_M$ and $\langle x \triangleleft \text{inr}.P \rangle_C$, respectively. The intermediate steps are correct.

5.2 Paper II: Prioritise the Best Variation

This section contains the paper with the same title, written in collaboration with Ornela Dardha, which was originally published in the journal Logical Methods in Computer Science, Volume 19, Issue 4, 2023, which is an extended journal version of the paper with the same title and authors originally published in the proceedings for the 41st International Conference on Formal Techniques for Distributed Objects, Components, and Systems (FORTE 2021) as part of the Lecture Notes in Computer Science (LNCS) series.

The work presented in the paper was conceived of by all the authors. I co-developed Priority GV and was primarily responsible for the initial draft of its theory and metatheory. I adapted the changes to the reduction semantics of Priority CP from my previous work on CP.

PRIORITISE THE BEST VARIATION

WEN KOKKE  AND ORNELA DARDHA 

^a University of Strathclyde, Glasgow, UK
e-mail address: wen.kokke@strath.ac.uk

^b School of Computing Science, University of Glasgow, UK
e-mail address: ornela.dardha@glasgow.ac.uk

ABSTRACT. Binary session types guarantee communication safety and session fidelity, but *alone* they cannot rule out deadlocks arising from the interleaving of different sessions.

In Classical Processes (CP) [Wad14]—a process calculus based on classical linear logic—deadlock freedom is guaranteed by combining channel creation and parallel composition under the same logical cut rule. Similarly, in Good Variation (GV) [Wad15, LM15]—a linear concurrent λ -calculus—deadlock freedom is guaranteed by combining channel creation and thread spawning under the same operation, called fork.

In both CP and GV, deadlock freedom is achieved at the expense of expressivity, as the only processes allowed are tree-structured. This is true more broadly than for CP and GV, and it holds for all works in the research line based on Curry-Howard correspondences between linear logic and session types, starting with Caires and Pfenning [CP10]. To overcome the limitations of tree-structured processes, Dardha and Gay [DG18a] define Priority CP (PCP), which allows cyclic-structured processes and restores deadlock freedom by using *priorities*, in line with Kobayashi and Padovani [Kob06, Pad14].

Following PCP, we present Priority GV (PGV), a variant of GV which decouples channel creation from thread spawning. Consequently, we type cyclic-structured processes and restore deadlock freedom by using priorities. We show that our type system is sound by proving subject reduction and progress. We define an encoding from PCP to PGV and prove that the encoding preserves typing and is sound and complete with respect to the operational semantics.

1. INTRODUCTION

Session types [Hon93, THK94, HVK98] are a type formalism that ensures *communication channels* are used according to their protocols, much like, *e.g.*, data types ensure that functions are used according to their signature. Session types have been studied in many settings. Most notably, they have been defined for the π -calculus [Hon93, THK94, HVK98], a foundational calculus for communication and concurrency, and the concurrent λ -calculi [GV10], including the main focus of our paper: Good Variation [Wad15, LM15, GV].

Key words and phrases: session types, π -calculus, functional programming, deadlock freedom, GV, CP.

Supported by the EU HORIZON 2020 MSCA RISE project 778233 “Behavioural Application Program Interfaces” (BehAPI).

GV is a concurrent λ -calculus with *binary* session types, where each channel is shared between exactly two processes. Binary session types guarantee two crucial properties: *communication safety*—*e.g.*, if the protocol says to transmit an integer, you transmit an integer—and *session fidelity*—*e.g.*, if the protocol says send, you send. A third crucial property is *deadlock freedom*, which ensures that processes do not have cyclic dependencies—*e.g.*, when two processes wait for each other to send a value. Binary session types alone are insufficient to rule out deadlocks arising from interleaved sessions, but several additional techniques have been developed to guarantee deadlock freedom in session-typed π -calculus and concurrent λ -calculus.

In the π -calculus literature, there have been several developments of Curry-Howard correspondences between session-typed π -calculus and linear logic [Gir87]: Caires and Pfenning’s π DILL [CP10] corresponds to dual intuitionistic linear logic [Bar96], and Wadler’s Classical Processes [Wad14, CP] corresponds to classical linear logic [Gir87, CLL]. Both calculi guarantee deadlock freedom, which they achieve by the combination of binary session types with a restriction on the way processes can be connected by channels: two processes can share at most one channel, and, more generally, information transmitted between any two processes must pass through one unique series of channels and intermediate processes. We refer to such processes as *tree-structured*, because the *communication graph*—where the vertices are ready processes and two vertices are connected by an edge if and only if the corresponding processes share a channel—is a tree. For π DILL, CP, and GV, tree-structure follows from a syntactic restriction: the combination of name restriction and parallel composition into a single syntactic construct, corresponding to the logical cut.

There are many downsides to combining name restriction and parallel composition, such as lack of modularity, difficulty typing structural congruence and formulating label-transition semantics. GV, specifically, struggles with a complicated metatheory due to the mismatch between its term language—where restriction and parallel composition are combined—and its configuration language—where they are not. There have been various approaches to decoupling restriction and parallel composition. Hypersequent CP [MP18, KMP19a, KMP19b, HCP], Hypersequent GV [FKD⁺21], and Linear Compositional Choreographies [CMS18] decouple them, but maintain the tree-structure of processes and a correspondence to linear logic, *e.g.*, while the typing rules for HCP are no longer exactly the proof rules for CLL, every typing derivations in HCP is isomorphic a proof in CLL, and vice versa. Priority CP [DG18b, PCP] weakens the correspondence to CLL. PCP is a non-conservative extension of CLL: every proof in CLL can be translated to a typing derivation in PCP [DP22], but PCP can prove strictly more theorems, including (partial) bijections between several CLL connectives. In exchange, PCP has a much more expressive language which allows cyclic-structured processes. PCP decouples CP’s cut rule into two separate constructs: one for parallel composition via a mix rule, and one for name restriction via a cycle rule. To restore deadlock freedom, PCP uses *priorities* [Kob06, Pad14]. Priorities encode the *order of actions* and rule out bad cyclic interleavings. Dardha and Gay [DG18b] prove cycle-elimination for PCP, adapting the cut-elimination proof for classical linear logic, and deadlock freedom follows as a corollary.

CP and GV are related via a pair of translations which satisfy simulation [LM16], and which can be tweaked to satisfy operational correspondence. The two calculi share the same strong guarantees. GV achieves deadlock freedom via a similar syntactic restriction: it combines channel creation and thread spawning into a single operation, called “fork”, which

is related to the cut construct in CP. Unfortunately, as with CP, this syntactic restriction has its downsides.

Our aim is to develop a more expressive version of GV while maintaining deadlock freedom. While process calculi have their advantages, *e.g.*, their succinctness compared to concurrent λ -calculi, we chose to work with GV for several reasons. In general, concurrent λ -calculi support higher-order functions, and have a capability for abstraction not usually present in process calculi. Within a concurrent λ -calculus, one can derive extensions of the communication capabilities of the language via well-understood extensions of the functional fragment, *e.g.*, we can derive internal/external choice from sum types. Concurrent λ -calculi, among other languages, maintain a clear separation between the program which the user writes and the configurations which represent the state of the system as it evaluates the program. However, our main motivation is that results obtained for λ -calculi transfer more easily to real-world functional programming languages. Case in point: we easily adapted the type system of PGV to Linear Haskell [BBN⁺18], which gives us a library for deadlock-free session-typed programming [KD21a]. The benefit of working specifically with GV, as opposed to other concurrent λ -calculi, is its relation to CP [Wad14], and its formal properties, including deadlock freedom.

We thus pose our research question for GV:

RQ: *Can we design a more expressive GV which guarantees deadlock freedom for cyclic-structured processes?*

We follow the line of work from CP to Priority CP, and present Priority GV (PGV), a variant of GV which decouples channel creation from thread spawning, thus allowing cyclic-structured processes, but which nonetheless guarantees deadlock freedom via priorities. This closes the circle of the connection between CP and GV [Wad14], and their priority-based versions, PCP [DG18b] and PGV. We cannot straightforwardly adapt the priority typing from PCP to PGV, as PGV adds higher-order functions. Instead, the priority typing for PGV follows the work by Padovani and Novara [PN15].

We make the following main contributions:

- (1) **Priority GV.** We present Priority GV (§.2, PGV), a session-typed functional language with priorities, and prove *subject reduction* (Theorem 3.5) and *progress* (Theorem 3.13).

We address several problems in the original GV language, most notably:

- (a) PGV does not require the pseudo-type $S^\#$;
- (b) Structural congruence is type preserving.

PGV answers our research question positively as it allows cyclic-structured binary session-typed processes that are deadlock free.

- (2) **Translation from PCP to PGV.** We present a *sound and complete encoding* of PCP [DG18b] in PGV (§.4). We prove the encoding preserves typing (Theorem 4.6) and satisfies operational correspondence (Theorem 4.7 and Theorem 4.10).

To obtain a tight correspondence, we update PCP, moving away from commuting conversions and reduction as cut elimination towards reduction based on structural congruence, as it is standard in process calculi.

This paper is an improved and extended version of a paper published at FORTE 2021 international conference [KD21b]. We present detailed examples and complete proofs of our technical results.

Example 2.1.

```

let (x, x') = new () in
let (y, y') = new () in
spawn ( λ(). let y = send (( ), y) in
        wait x'; close y );
let x = send (( ), x) in
let (( ), y') = recv y' in
close x; wait y'

```

Example 2.2.

```

let (x, x') = new () in
let (y, y') = new () in
spawn ( λ(). let y = send (( ), y) in
        wait x'; close y );
let (( ), y') = recv y' in
let x = send (( ), x) in
close x; wait y'

```

2. PRIORITY GV

We present Priority GV (PGV), a session-typed functional language based on GV [Wad15, LM15] which uses priorities à la Kobayashi and Padovani [Kob06, PN15] to enforce deadlock freedom. Priority GV is more flexible than GV because it allows processes to share more than one communication channel.

We illustrate this with two programs in PGV, Example 2.1 and Example 2.2. Each program contains two processes—the main process, and the child process created by **spawn**—which communicate using *two* channels. The child process receives a unit over the channel x/x' , and then sends a unit over the channel y/y' . The main process does one of two things:

- (a) in Example 2.1, it sends a unit over the channel x/x' , and then waits to receive a unit over the channel y/y' ;
- (b) in Example 2.2, it does these in the opposite order, which results in a deadlock.

PGV is more expressive than GV: Example 2.1 is typeable and guaranteed to be deadlock-free in PGV, but is not typeable in GV [Wad14] and not guaranteed deadlock-free in GV's predecessor [GV10]. We believe PGV is a non-conservative *extension* of GV, as CP can be embedded in a Kobayashi-style system [DP18].

2.1. Syntax of Types and Terms.

Session types. Session types (S) are defined by the following grammar:

$$S ::= !^o T.S \mid ?^o T.S \mid \mathbf{end}_!^o \mid \mathbf{end}_?^o$$

Session types $!^o T.S$ and $?^o T.S$ describe the endpoints of a channel over which we send or receive a value of type T , and then proceed as S . Types $\mathbf{end}_!^o$ and $\mathbf{end}_?^o$ describe endpoints of a channel whose communication has finished, and over which we must synchronise before closing the channel. Each connective in a session type is annotated with a *priority* $o \in \mathbb{N}$.

Types. Types (T, U) are defined by the following grammar:

$$T, U ::= T \times U \mid \mathbf{1} \mid T + U \mid \mathbf{0} \mid T \multimap^{p,q} U \mid S$$

Types $T \times U$, $\mathbf{1}$, $T + U$, and $\mathbf{0}$ are the standard linear λ -calculus product type, unit type, sum type, and empty type. Type $T \multimap^{p,q} U$ is the standard linear function type, annotated with *priority bounds* $p, q \in \mathbb{N} \cup \{\perp, \top\}$. Every session type is also a type. Given a function with type $T \multimap^{p,q} U$, p is a *lower bound* on the priorities of the endpoints captured by the body of the function, and q is an *upper bound* on the priority of the communications

that take place as a result of applying the function. The type of pure functions $T \multimap U$, *i.e.*, those which perform no communications, is syntactic sugar for $T \multimap^{\top, \perp} U$. The lower bound for a pure function is $\top$ as pure functions never start communicating. For similar reasons, the upper bound for a pure function is $\perp$.

We postulate that the only function types—and, consequently, sequents—that are inhabited in PGV are pure functions and functions $T \multimap^{p,q} U$ for which $p < q$.

Typing Environments. Typing environments Γ, Δ, Θ associate types to names. Environments are linear, so two environments can only be combined as Γ, Δ if their names are distinct, *i.e.*, $\text{fv}(\Gamma) \cap \text{fv}(\Delta) = \emptyset$.

$$\Gamma, \Delta ::= \emptyset \mid \Gamma, x : T$$

Type Duality. Duality plays a crucial role in session types. The two endpoints of a channel are assigned dual types, ensuring that, for instance, whenever one program sends a value on a channel, the program on the other end is waiting to receive. Each session type S has a dual, written $\overline{S}$. Duality is an involutive function which preserves priorities:

$$\overline{!^o T.S} = ?^o \overline{T.S} \quad \overline{?^o T.S} = !^o \overline{T.S} \quad \overline{\text{end}_!^o} = \text{end}_?^o \quad \overline{\text{end}_?^o} = \text{end}_!^o$$

Priorities. Function $\text{pr}(\cdot)$ returns the smallest priority of a session type. The type system guarantees that the top-most connective always holds the smallest priority, so we simply return the priority of the top-most connective:

$$\text{pr}(!^o T.S) = o \quad \text{pr}(?^o T.S) = o \quad \text{pr}(\text{end}_!^o) = o \quad \text{pr}(\text{end}_?^o) = o$$

We extend the function $\text{pr}(\cdot)$ to types and typing contexts by returning the smallest priority in the type or context, or $\top$ if there is no priority. We use $\sqcap$ and $\sqcup$ to denote the minimum and maximum, respectively:

$$\begin{array}{lll} \text{pr}(T \times U) & = & \text{pr}(T) \sqcap \text{pr}(U) & \text{pr}(\mathbf{1}) & = & \top \\ \text{pr}(T + U) & = & \text{pr}(T) \sqcap \text{pr}(U) & \text{pr}(\mathbf{0}) & = & \top \\ \text{pr}(T \multimap^{p,q} U) & = & p & \text{pr}(S) & = & \text{pr}(S) \\ \text{pr}(\Gamma, x : A) & = & \text{pr}(\Gamma) \sqcap \text{pr}(A) & \text{pr}(\emptyset) & = & \top \end{array}$$

Terms. Terms (L, M, N) are defined by the following grammar:

$$\begin{array}{l} L, M, N ::= x \mid K \mid \lambda x. M \mid M N \\ \quad \mid () \mid M; N \mid (M, N) \mid \text{let } (x, y) = M \text{ in } N \\ \quad \mid \text{inl } M \mid \text{inr } M \mid \text{case } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} \mid \text{absurd } M \\ K ::= \text{link} \mid \text{new} \mid \text{spawn} \mid \text{send} \mid \text{recv} \mid \text{close} \mid \text{wait} \end{array}$$

Let x, y, z , and w range over variable names. Occasionally, we use a, b, c , and d . The term language is the standard linear λ -calculus with products, sums, and their units, extended with constants K for the communication primitives.

Constants are best understood in conjunction with their typing and reduction rules in Figs. 1 and 2.

Briefly, **link** links two endpoints together, forwarding messages from one to the other, **new** creates a new channel and returns a pair of its endpoints, and **spawn** spawns off its argument as a new thread.

The **send** and **recv** functions send and receive values on a channel. However, since the typing rules for PGV ensure the linear usage of endpoints, they also return a new copy of the endpoint to continue the session.

The **close** and **wait** functions close a channel.

We use syntactic sugar to make terms more readable: we write **let** $x = M$ **in** N in place of $(\lambda x.N) M$, $\lambda().M$ in place of $\lambda z.z; M$, and $\lambda(x,y).M$ in place of $\lambda z.\mathbf{let} (x,y) = z \mathbf{in} M$. We can recover GV's **fork** as $\lambda x.\mathbf{let} (y,z) = \mathbf{new} () \mathbf{in} \mathbf{spawn} (\lambda().x y); z$.

Internal and External Choice. Typically, session-typed languages feature constructs for internal and external choice. In GV, these can be defined in terms of the core language, by sending or receiving a value of a sum type [LM15]. We use the following syntactic sugar for internal $(S \oplus^o S')$ and external $(S \&^o S')$ choice and their units:

$$\begin{aligned} S \oplus^o S' &\triangleq !^o(\overline{S} + \overline{S}').\mathbf{end}_!^{o+1} & \oplus^o\{\} &\triangleq !^o\mathbf{0}.\mathbf{end}_!^{o+1} \\ S \&^o S' &\triangleq ?^o(S + S').\mathbf{end}_?^{o+1} & \&^o\{\} &\triangleq ?^o\mathbf{0}.\mathbf{end}_?^{o+1} \end{aligned}$$

As the syntax for units suggests, these are the binary and nullary forms of the more common n-ary choice constructs $\oplus^o\{l_i : S_i\}_{i \in I}$ and $\&^o\{l_i : S_i\}_{i \in I}$, which one may obtain generalising the sum types to variant types. For simplicity, we present only the binary and nullary forms.

Similarly, we use syntactic sugar for the term forms of choice, which combine sending and receiving with the introduction and elimination forms for the sum and empty types. There are two constructs for binary internal choice, expressed using the meta-variable ℓ which ranges over $\{\mathbf{inl}, \mathbf{inr}\}$. As there is no introduction for the empty type, there is no construct for nullary internal choice:

$$\begin{aligned} \mathbf{select} \ell &\triangleq \lambda x.\mathbf{let} (y,z) = \mathbf{new} \mathbf{in} \mathbf{close} (\mathbf{send} (\ell y, x)); z \\ \mathbf{offer} L \{\mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N\} &\triangleq \\ &\mathbf{let} (z,w) = \mathbf{recv} L \mathbf{in} \mathbf{wait} w; \mathbf{case} z \{\mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N\} \\ \mathbf{offer} L \{\} &\triangleq \mathbf{let} (z,w) = \mathbf{recv} L \mathbf{in} \mathbf{wait} w; \mathbf{absurd} z \end{aligned}$$

2.2. Operational Semantics.

Configurations. Priority GV terms are evaluated as part of a configuration of processes. Configurations are defined by the following grammar:

$$\phi ::= \bullet \mid \circ \quad \mathcal{C}, \mathcal{D}, \mathcal{E} ::= \phi M \mid \mathcal{C} \parallel \mathcal{D} \mid (\nu xx')\mathcal{C}$$

Configurations $(\mathcal{C}, \mathcal{D}, \mathcal{E})$ consist of threads ϕM , parallel compositions $\mathcal{C} \parallel \mathcal{D}$, and name restrictions $(\nu xx')\mathcal{C}$. To preserve the functional nature of PGV, where programs return a single value, we use flags (ϕ) to differentiate between the main thread, marked $\bullet$, and child threads created by **spawn**, marked $\circ$. Only the main thread returns a value. We determine the flag of a configuration by combining the flags of all threads in that configuration:

$$\bullet + \circ = \bullet \quad \circ + \bullet = \bullet \quad \circ + \circ = \circ \quad (\bullet + \bullet \text{ is undefined})$$

To distinguish between the use of $\circ$ to mark child threads [LM15] and the use of the meta-variable $\circ$ for priorities [DG18b], they are typeset in a different font and colour.

Values. Values (V, W), evaluation contexts (E), thread evaluation contexts ($\mathcal{F}$), and configuration contexts ($\mathcal{G}$) are defined by the following grammars:

$$\begin{aligned}
V, W &::= x \mid K \mid \lambda x. M \mid () \mid (V, W) \mid \mathbf{inl} V \mid \mathbf{inr} V \\
E &::= \square \mid E M \mid V E \\
&\quad \mid E; N \mid (E, M) \mid (V, E) \mid \mathbf{let} (x, y) = E \mathbf{in} M \\
&\quad \mid \mathbf{inl} E \mid \mathbf{inr} E \mid \mathbf{case} E \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} \mid \mathbf{absurd} E \\
\mathcal{F} &::= \phi E \\
\mathcal{G} &::= \square \mid \mathcal{G} \parallel \mathcal{C} \mid (\nu xy) \mathcal{G}
\end{aligned}$$

Values are the subset of terms which cannot reduce further. Evaluation contexts are one-hole term contexts, *i.e.*, terms with exactly one hole, written $\square$. We write $E[M]$ for the evaluation context E with its hole replaced by the term M . Evaluation contexts are specifically those one-hole term contexts under which term reduction can take place. Thread contexts are a convenient way to lift the notion of evaluation contexts to threads. We write $\mathcal{F}[M]$ for the thread context $\mathcal{F}$ with its hole replaced by the term M . Configuration contexts are one-hole configuration contexts, *i.e.*, configurations with exactly one hole, written $\square$. Specifically, configuration contexts are those one-hole term contexts under which configuration reduction can take place. The definition for $\mathcal{G}$ only gives the case in which the hole is in the left-most parallel process, *i.e.*, it only defines $\mathcal{G} \parallel \mathcal{C}$ and not $\mathcal{C} \parallel \mathcal{G}$. The latter is not needed, as $\parallel$ is symmetric under structural congruence, though it would be harmless to add. We write $\mathcal{G}[\mathcal{C}]$ for the evaluation context $\mathcal{G}$ with its hole replaced by the term $\mathcal{C}$.

Reduction Relation. We factor the reduction relation of PGV into a deterministic reduction on terms ($\rightarrow_M$) and a non-deterministic reduction on configurations ($\rightarrow_C$), see Fig. 1. We write $\rightarrow_M^+$ and $\rightarrow_C^+$ for the transitive closures, and $\rightarrow_M^*$ and $\rightarrow_C^*$ for the reflexive-transitive closures.

Term reduction is the standard call-by-value, left-to-right evaluation for GV, and only deviates from reduction for the linear λ -calculus in that it reduces terms to values or ready terms waiting to perform a communication action.

Configuration reduction resembles evaluation for a process calculus: E-LINK, E-SEND, and E-CLOSE perform communications, E-LIFTC allows reduction under configuration contexts, and E-LIFTSC embeds a structural congruence $\equiv$. The remaining rules mediate between the process calculus and the functional language: E-NEW and E-SPAWN evaluate the **new** and **spawn** constructs, creating the equivalent configuration constructs, and E-LIFTM embeds term reduction.

Structural congruence satisfies the following axioms: SC-LINKSWAP allows swapping channels in the link process. SC-RESLINK allows restriction to be applied to link which is structurally equivalent to the terminated process, thus allowing elimination of unnecessary restrictions. SC-RESSWAP allows swapping channels and SC-RESCOMM states that restriction is commutative. SC-RESEXT is the standard scope extrusion rule. Rules SC-PARNIL, SC-PARCOMM and SC-PARASSOC state that parallel composition uses the terminated process as the neutral element; it is commutative and associative.

While our configuration reduction is based on the standard evaluation for GV, the increased expressiveness of PGV allows us to simplify the relation on two counts.

- (i) *We decompose the **fork** construct.* In GV, **fork** creates a new channel, spawns a child thread, and, when the child thread finishes, it closes the channel to its parent. In PGV, these are three separate operations: **new**, **spawn**, and **close**. We no longer require

Term reduction.

E-LAM	$(\lambda x.M) V$	$\longrightarrow_M M\{V/x\}$
E-UNIT	$() ; M$	$\longrightarrow_M M$
E-PAIR	let $(x, y) = (V, W)$ in M	$\longrightarrow_M M\{V/x\}\{W/y\}$
E-INL	case inl V {inl $x \mapsto M$; inr $y \mapsto N$ }	$\longrightarrow_M M\{V/x\}$
E-INR	case inr V {inl $x \mapsto M$; inr $y \mapsto N$ }	$\longrightarrow_M N\{V/y\}$

$$\frac{\text{E-LIFT} \quad M \longrightarrow_M M'}{E[M] \longrightarrow_M E[M']}$$

Structural congruence.

SC-LINKSWAP	$\mathcal{F}[\mathbf{link}(x, y)]$	$\equiv \mathcal{F}[\mathbf{link}(y, x)]$
SC-RESLINK	$(\nu xy)(\phi \mathbf{link}(x, y))$	$\equiv \phi()$
SC-RESWAP	$(\nu xy)\mathcal{C}$	$\equiv (\nu yx)\mathcal{C}$
SC-RESCOMM	$(\nu xy)(\nu zw)\mathcal{C}$	$\equiv (\nu zw)(\nu xy)\mathcal{C}$, if $\{x, y\} \cap \{z, w\} = \emptyset$
SC-RESEXT	$(\nu xy)(\mathcal{C} \parallel \mathcal{D})$	$\equiv \mathcal{C} \parallel (\nu xy)\mathcal{D}$, if $x, y \notin \text{fv}(\mathcal{C})$
SC-PARNIL	$\mathcal{C} \parallel \circ()$	$\equiv \mathcal{C}$
SC-PARCOMM	$\mathcal{C} \parallel \mathcal{D}$	$\equiv \mathcal{D} \parallel \mathcal{C}$
SC-PARASSOC	$\mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E})$	$\equiv (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}$

Configuration reduction.

E-LINK	$(\nu xy)(\mathcal{F}[\mathbf{link}(w, x)] \parallel \mathcal{C})$	$\longrightarrow_{\mathcal{C}} \mathcal{F}[] \parallel \mathcal{C}\{w/y\}$
E-NEW	$\mathcal{F}[\mathbf{new}()]$	$\longrightarrow_{\mathcal{C}} (\nu xy)(\mathcal{F}[(x, y)])$, if $x, y \notin \text{fv}(\mathcal{F})$
E-SPAWN	$\mathcal{F}[\mathbf{spawn} V]$	$\longrightarrow_{\mathcal{C}} \mathcal{F}[] \parallel \circ V()$
E-SEND	$(\nu xy)(\mathcal{F}[\mathbf{send}(V, x)] \parallel \mathcal{F}'[\mathbf{recv} y])$	$\longrightarrow_{\mathcal{C}} (\nu xy)(\mathcal{F}[x] \parallel \mathcal{F}'[(V, y)])$
E-CLOSE	$(\nu xy)(\mathcal{F}[\mathbf{wait} x] \parallel \mathcal{F}'[\mathbf{close} y])$	$\longrightarrow_{\mathcal{C}} \mathcal{F}[] \parallel \mathcal{F}'[]$

$$\frac{\text{E-LIFT}_{\mathcal{C}} \quad \mathcal{C} \longrightarrow_{\mathcal{C}} \mathcal{C}'}{\mathcal{G}[\mathcal{C}] \longrightarrow_{\mathcal{C}} \mathcal{G}[\mathcal{C}']} \quad \frac{\text{E-LIFT}_M \quad M \longrightarrow_M M'}{\mathcal{F}[M] \longrightarrow_M \mathcal{F}[M']} \quad \frac{\text{E-LIFT}_{\mathcal{SC}} \quad \mathcal{C} \equiv \mathcal{C}' \quad \mathcal{C}' \longrightarrow_{\mathcal{C}} \mathcal{D}' \quad \mathcal{D}' \equiv \mathcal{D}}{\mathcal{C} \longrightarrow_{\mathcal{C}} \mathcal{D}}$$

Figure 1: Operational Semantics for PGV.

that every child thread finishes by returning a terminated channel. Consequently, we also simplify the evaluation of the **link** construct.

Intuitively, evaluating **link** causes a substitution: if we have a channel bound as (νxy) , then **link** (w, x) replaces all occurrences of y by w . However, in GV, **link** is required to return a terminated channel, which means that the semantics for **link** must create a fresh channel of type **end_l/end_r**. The endpoint of type **end_l** is returned by the **link** construct, and a **wait** on the other endpoint guards the actual substitution. In PGV, evaluating **link** simply causes a substitution.

- (ii) *Our structural congruence is type preserving.* Consequently, we can embed it directly into the reduction relation. In GV, this is not the case, and subject reduction relies on proving that if the result of rewriting via $\equiv$ followed by reducing via $\longrightarrow_{\mathcal{C}}$ is an ill-typed configuration, we can rewrite it to a well-typed configuration via $\equiv$.

Static Typing Rules.

$\frac{}{x : T \vdash^\perp x : T}$	$\frac{}{\emptyset \vdash^\perp K : T}$	$\frac{\text{T-LAM} \quad \Gamma, x : T \vdash^q M : U}{\Gamma \vdash^\perp \lambda x.M : T \multimap^{\text{pr}(\Gamma), q} U}$
$\frac{\text{T-APP} \quad \Gamma \vdash^p M : T \multimap^{p', q'} U \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta) \quad q < p'}{\Gamma, \Delta \vdash^{p \sqcup q \sqcup q'} M N : U}$	$\frac{\text{T-UNIT} \quad \emptyset \vdash^\perp () : \mathbf{1}}{\emptyset \vdash^\perp () : \mathbf{1}}$	$\frac{\text{T-LETUNIT} \quad \Gamma \vdash^p M : \mathbf{1} \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{p \sqcup q} M; N : T}$
$\frac{\text{T-PAIR} \quad \Gamma \vdash^p M : T \quad \Delta \vdash^q N : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{p \sqcup q} (M, N) : T \times U}$	$\frac{\text{T-LETPAIR} \quad \Gamma \vdash^p M : T \times T' \quad \Delta, x : T, y : T' \vdash^q N : U \quad p < \text{pr}(\Delta, T, T')}{\Gamma, \Delta \vdash^{p \sqcup q} \text{let } (x, y) = M \text{ in } N : U}$	$\frac{\text{T-INL} \quad \Gamma \vdash^p M : T \quad \text{pr}(T) = \text{pr}(U)}{\Gamma \vdash^p \text{inl } M : T + U}$
$\frac{\text{T-INR} \quad \Gamma \vdash^p M : U \quad \text{pr}(T) = \text{pr}(U)}{\Gamma \vdash^p \text{inr } M : T + U}$	$\frac{\text{T-CASESUM} \quad \Gamma \vdash^p L : T + T' \quad \Delta, x : T \vdash^q M : U \quad \Delta, y : T' \vdash^q N : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{p \sqcup q} \text{case } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} : U}$	$\frac{\text{T-ABSURD} \quad \Gamma \vdash^p M : \mathbf{0}}{\Gamma, \Delta \vdash^p \text{absurd } M : T}$

Type Schemas for Constants.

$\text{link} : S \times \bar{S} \multimap \mathbf{1}$	$\text{new} : \mathbf{1} \multimap S \times \bar{S}$	$\text{spawn} : (\mathbf{1} \multimap^{p, q} \mathbf{1}) \multimap \mathbf{1}$
$\text{send} : T \times !^o T.S \multimap^{\top, o} S$	$\text{recv} : ?^o T.S \multimap^{\top, o} T \times S$	
$\text{close} : \text{end}_!^o \multimap^{\top, o} \mathbf{1}$	$\text{wait} : \text{end}_?^o \multimap^{\top, o} \mathbf{1}$	

Runtime Typing Rules.

$\frac{\text{T-MAIN} \quad \Gamma \vdash^p M : T}{\Gamma \vdash^\bullet \bullet M}$	$\frac{\text{T-CHILD} \quad \Gamma \vdash^p M : \mathbf{1}}{\Gamma \vdash^\circ \circ M}$	$\frac{\text{T-RES} \quad \Gamma, x : S, y : \bar{S} \vdash^\phi \mathcal{C}}{\Gamma \vdash^\phi (\nu xy)\mathcal{C}}$	$\frac{\text{T-PAR} \quad \Gamma \vdash^\phi \mathcal{C} \quad \Delta \vdash^{\phi'} \mathcal{D}}{\Gamma, \Delta \vdash^{\phi + \phi'} \mathcal{C} \parallel \mathcal{D}}$
---	---	--	--

Figure 2: Typing Rules for PGV.

2.3. Typing Rules.

Terms Typing. Typing rules for terms are at the top of Fig. 2. Terms are typed by a judgement $\Gamma \vdash^p M : T$ stating that “a term M has type T and an upper bound on its priority p under the typing environment Γ ”. Typing for the linear λ -calculus is standard. Linearity is ensured by splitting environments on branching rules, requiring that the environment in the variable rule consists of just the variable, and the environment in the constant and unit rules are empty. Constants K are typed using type schemas, which hold for any concrete assignment of types and priorities to their meta-variables. Instantiated type schemas are embedded into typing derivations using T-CONST in Fig. 2, *e.g.*, the type schema for **send** can be instantiated with $o = 2$, $T = \mathbf{1}$, and $S = \mathbf{0}$, and embedded using T-CONST to give the following typing derivation:

$$\text{send} : \mathbf{1} \times !^2 \mathbf{1.0} \multimap^{\top, 2} \mathbf{0}$$

The typing rules treat all variables as linear resources, even those of non-linear types such as $\mathbf{1}$, though they can easily be extended to allow values with unrestricted usage [Wad14].

The only non-standard feature of the typing rules is the priority annotations. Priorities are based on obligations/capabilities used by Kobayashi [Kob06], and simplified to single priorities following Padovani [Pad14]. The integration of priorities into GV is adapted from Padovani and Novara [PN15]. Paraphrasing Dardha and Gay [DG18b], priorities obey the following two laws:

- (i) an action with lower priority happens before an action with higher priority; and
- (ii) communication requires *equal* priorities for dual actions.

In PGV, we keep track of a lower and upper bound on the priorities of a term, *i.e.*, while evaluating the term, when it starts communicating, and when it finishes, respectively. The upper bound is written on the sequent and the lower bound is approximated from the typing environment, *e.g.*, for $\Gamma \vdash^p M : T$ the upper bound is p and the lower bound is at least $\text{pr}(\Gamma)$. The latter is correct because a term cannot communicate at a priority earlier than the earliest priority amongst the channels it has access to. It is an *approximation* on function terms, as these can “skip” communication by returning the corresponding channel unused. However, linearity prevents such functions from being used in well typed configurations: once the unused channel’s priority has passed, it can no longer be used.

Typing rules for sequential constructs enforce sequentiality, *e.g.*, the typing for $M; N$ has a side condition which requires that the upper bound of M is smaller than the lower bound of N , *i.e.*, M finishes before N starts. The typing rule for **new** ensures that both endpoints of a channel share the same priorities. Together, these two constraints guarantee deadlock freedom.

To illustrate this, let’s go back to the deadlocked program in Example 2.2. Crucially, it composes the terms below in parallel. While each of these terms itself is well typed, they impose opposite conditions on the priorities, so connecting x to x' and y to y' using T-RES is ill-typed, as there is no assignment to o and o' that can satisfy both $o < o'$ and $o' < o$.

(We omit the priorities on $\mathbf{end}_l$ and $\mathbf{end}_r$.)

$$\begin{array}{c}
\frac{y' : ?^{o'} \mathbf{1.end}_r \vdash^{o'} \mathbf{recv } y' : \mathbf{1} \times \mathbf{end}_r}{x : !^o \mathbf{1.end}_l, y' : \mathbf{end}_r \vdash^p \mathbf{let } x = \mathbf{send } ((), x) \mathbf{ in } \dots : \mathbf{1} \quad o' < o} \\
x : !^o \mathbf{1.end}_l, y' : ?^{o'} \mathbf{1.end}_r \vdash^p \mathbf{let } ((), y') = \mathbf{recv } y' \mathbf{ in let } x = \mathbf{send } ((), x) \mathbf{ in } \dots : \mathbf{1}
\end{array}$$

$$\begin{array}{c}
\frac{x' : ?^o \mathbf{1.end}_r \vdash^o \mathbf{recv } x' : \mathbf{1} \times \mathbf{end}_r}{y : !^{o'} \mathbf{1.end}_l, x' : \mathbf{end}_r \vdash^q \mathbf{let } y = \mathbf{send } ((), y) \mathbf{ in } \dots : \mathbf{1} \quad o < o'} \\
y : !^{o'} \mathbf{1.end}_l, x' : ?^o \mathbf{1.end}_r \vdash^q \mathbf{let } ((), x') = \mathbf{recv } x' \mathbf{ in let } y = \mathbf{send } ((), y) \mathbf{ in } \dots : \mathbf{1}
\end{array}$$

Closures suspend communication, so T-LAM stores the priority bounds of the function body on the function type, and T-APP restores them. For instance, $\lambda x. \mathbf{send } (x, y)$ is assigned the type $A \multimap^{o,o} S$, *i.e.*, a function which, when applied, starts and finishes communicating at priority o .

$$\begin{array}{c}
\frac{\mathbf{send} : T \times !^o T.S \multimap^{\top, o} S \quad \frac{x : T \vdash^\perp x : T \quad x : T, y : !^o T.S \vdash^\perp y : !^o T.S}{x : T, y : !^o T.S \vdash^\perp (x, y) : T \times !^o T.S}}{x : T, y : !^o T.S \vdash^o \mathbf{send } (x, y) : S} \\
y : !^o T.S \vdash^\perp \lambda x. \mathbf{send } (x, y) : T \multimap^{o,o} S
\end{array}$$

For simplicity, we assume priority annotations are not inferred, but provided as an input to type checking. However, for any term, priorities can be inferred, *e.g.*, by using the topological ordering of the directed graph where the vertices are the priority meta-variables and the edges are the inequality constraints between the priority meta-variables in the typing derivation.

Configurations Typing. Typing rules for configurations are at the bottom of Fig. 2. Configurations are typed by a judgement $\Gamma \vdash^\phi \mathcal{C}$ stating that “a configuration $\mathcal{C}$ with flag ϕ is well typed under typing environment Γ ”. Configuration typing is based on the standard typing for GV. Terms are embedded either as main or as child threads. The priority bound from the term typing is discarded, as configurations contain no further blocking actions. Main threads are allowed to return a value, whereas child threads are required to return the unit value. Sequents are annotated with a flag ϕ , which ensures that there is at most one main thread.

While our configuration typing is based on the standard typing for GV, it differs on two counts:

- (i) *we require that child threads return the unit value*, as opposed to a terminated channel; and
- (ii) *we simplify typing for parallel composition*.

In order to guarantee deadlock freedom, in GV each parallel composition must split *exactly one* channel of the channel pseudo-type $S^\sharp$ into two endpoints of type S and $\bar{S}$. Consequently, associativity of parallel composition does not preserve typing. In PGV, we guarantee deadlock freedom using priorities, which removes the need for the channel pseudo-type $S^\sharp$, and simplifies typing for parallel composition, while restoring type preservation for the structural congruence.

$$\begin{array}{c}
\text{T-LAMUNIT} \\
\frac{\Gamma \vdash^q M : T}{\Gamma \vdash^\perp \lambda().M : \mathbf{1} \multimap^{\text{pr}(\Gamma),q} T} \triangleq \frac{\frac{z : \mathbf{1} \vdash^\perp z : \mathbf{1} \quad \Gamma \vdash^q M : T}{\Gamma, z : \mathbf{1} \vdash^q z; M : T}}{\Gamma \vdash^\perp \lambda z.z; M : \mathbf{1} \multimap^{\text{pr}(\Gamma),q} T}
\\[20pt]
\text{T-LAMPAIR} \\
\frac{\Gamma, x : T, y : T' \vdash^q M : U}{\Gamma \vdash^\perp \lambda(x, y).M : T \times T' \multimap^{\text{pr}(\Gamma),q} U} \triangleq \\
\frac{\frac{z : T \times T' \vdash^\perp z : T \times T' \quad \Gamma, x : T, y : T' \vdash^q M : U}{\Gamma, z : T \times T' \vdash^q \text{let } (x, y) = z \text{ in } M : T}}{\Gamma \vdash^\perp \lambda z.\text{let } (x, y) = z \text{ in } M : T \times T' \multimap^{\text{pr}(\Gamma),q} U}
\\[20pt]
\text{T-LET} \\
\frac{\Gamma \vdash^p M : T \quad \Delta, x : T \vdash^q N : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{p \sqcup q} \text{let } x = M \text{ in } N : U} \triangleq \\
\frac{\frac{\Delta, x : T \vdash^q N : U}{\Delta \vdash^\perp \lambda x.N : T \multimap^{\text{pr}(\Delta),q} U} \quad \Gamma \vdash^p M : T \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{q \sqcup p} (\lambda x.N) M : U}
\\[20pt]
\text{T-FORK} \\
\frac{}{\emptyset \vdash^\perp \text{fork} : (S \multimap^{p,q} \mathbf{1}) \multimap \overline{S}} \triangleq \\
\frac{\text{(A)} \quad \frac{}{\emptyset \vdash^\perp () : \mathbf{1}}}{\emptyset \vdash^\perp \text{new } () : S \times \overline{S}} \\
\frac{\frac{x : S \multimap^{p,q} \mathbf{1} \vdash^\perp x : S \multimap^{p,q} \mathbf{1} \quad y : S \vdash^\perp y : S}{x : S \multimap^{p,q} \mathbf{1}, y : S \vdash^q x y : \mathbf{1}}}{\text{(B)} \quad \frac{x : S \multimap^{p,q} \mathbf{1}, y : S \vdash^\perp \lambda().x y : \mathbf{1} \multimap^{p,q} \mathbf{1}}{x : S \multimap^{p,q} \mathbf{1}, y : S \vdash^\perp \text{spawn } (\lambda().x y) : \mathbf{1} \quad z : \overline{S} \vdash^\perp z : \overline{S}}} \\
\frac{x : S \multimap^{p,q} \mathbf{1}, y : S, z : \overline{S} \vdash^\perp \text{spawn } (\lambda().x y); z : \overline{S}}{x : S \multimap^{p,q} \mathbf{1} \vdash^\perp \text{let } (y, z) = \text{new } () \text{ in spawn } (\lambda().x y); z : \overline{S}} \\
\frac{}{\emptyset \vdash^\perp \lambda x.\text{let } (y, z) = \text{new } () \text{ in spawn } (\lambda().x y); z : (S \multimap^{p,q} \mathbf{1}) \multimap \overline{S}}
\\[20pt]
\text{(A)} = \text{new} : \mathbf{1} \multimap S \times \overline{S} \qquad \text{(B)} = \text{spawn} : (\mathbf{1} \multimap^{p,q} \mathbf{1}) \multimap \mathbf{1}
\end{array}$$

Figure 3: Typing Rules for Syntactic Sugar for PGV (T-LAMUNIT, T-LAMPAIR, T-LET, and T-FORK).

$$\begin{array}{c}
\text{T-SELECT-INL} \\
\frac{\text{pr}(S) = \text{pr}(S')}{\emptyset \vdash \perp \text{select inl} : S \oplus^o S' \multimap^{\top, o} S} \triangleq \\
\\
\begin{array}{c}
\text{(A)} \quad \frac{\overline{\emptyset \vdash () : \mathbf{1}}}{\emptyset \vdash \text{new } () : \overline{S} \times S} \\
\\
\text{(c)} \quad \frac{\frac{\overline{y : \overline{S} \vdash y : \overline{S}}}{y : \overline{S} \vdash \text{inl } y : \overline{S} + \overline{S'}} \quad \overline{x : S \oplus^o S' \vdash x : S \oplus^o S'}}{x : S \oplus^o S', y : \overline{S} \vdash (\text{inl } y, x) : (\overline{S} + \overline{S'}) \times (S \oplus^o S')} \\
\\
\text{(B)} \quad \frac{x : S \oplus^o S', y : \overline{S} \vdash \text{send } (\text{inl } y, x) : \text{end}_!^{o+1}}{x : S \oplus^o S', y : \overline{S} \vdash \text{close } (\text{send } (\text{inl } y, x)) : \mathbf{1}} \\
\\
\frac{z : S \vdash z : S}{x : S \oplus^o S', y : \overline{S}, z : S \vdash \text{close } (\text{send } (\text{inl } y, x)); z : S} \\
\\
\frac{x : S \oplus^o S' \vdash \text{let } (y, z) = \text{new } () \text{ in close } (\text{send } (\text{inl } y, x)); z : S}{\emptyset \vdash \lambda x. \text{let } (y, z) = \text{new } () \text{ in close } (\text{send } (\text{inl } y, x)); z : S \oplus^o S' \multimap^{\top, o} S}
\end{array} \\
\\
\begin{array}{c}
\text{(A)} = \text{new} : \mathbf{1} \multimap^{\top, o} \overline{S} \times S \quad \text{(B)} = \text{close} : \text{end}_!^{o+1} \multimap^{\top, o+1} \mathbf{1} \\
\\
\text{(c)} = \text{send} : (\overline{S} + \overline{S'}) \times (S \oplus^o S') \multimap^{\top, o} \text{end}_!^{o+1}
\end{array}
\end{array}$$

$$\begin{array}{c}
\text{T-SELECT-INR} \\
\frac{\text{pr}(S) = \text{pr}(S')}{\emptyset \vdash \perp \text{select inr} : S \oplus^o S' \multimap^{\top, o} S'} \triangleq \\
\\
\begin{array}{c}
\text{(A)} \quad \frac{\overline{\emptyset \vdash () : \mathbf{1}}}{\emptyset \vdash \text{new } () : \overline{S'} \times S'} \\
\\
\text{(c)} \quad \frac{\frac{\overline{y : \overline{S'} \vdash y : \overline{S'}}}{y : \overline{S'} \vdash \text{inr } y : \overline{S} + \overline{S'}} \quad \overline{x : S \oplus^o S' \vdash x : S \oplus^o S'}}{x : S \oplus^o S', y : \overline{S'} \vdash (\text{inr } y, x) : (\overline{S} + \overline{S'}) \times (S \oplus^o S')} \\
\\
\text{(B)} \quad \frac{x : S \oplus^o S', y : \overline{S'} \vdash \text{send } (\text{inr } y, x) : \text{end}_!^{o+1}}{x : S \oplus^o S', y : \overline{S'} \vdash \text{close } (\text{send } (\text{inr } y, x)) : \mathbf{1}} \\
\\
\frac{z : S' \vdash z : S'}{x : S \oplus^o S', y : \overline{S'}, z : S' \vdash \text{close } (\text{send } (\text{inr } y, x)); z : S'} \\
\\
\frac{x : S \oplus^o S' \vdash \text{let } (y, z) = \text{new } () \text{ in close } (\text{send } (\text{inr } y, x)); z : S'}{\emptyset \vdash \lambda x. \text{let } (y, z) = \text{new } () \text{ in close } (\text{send } (\text{inr } y, x)); z : S \oplus^o S' \multimap^{\top, o} S'}
\end{array} \\
\\
\begin{array}{c}
\text{(A)} = \text{new} : \mathbf{1} \multimap^{\top, o} \overline{S'} \times S' \quad \text{(B)} = \text{close} : \text{end}_!^{o+1} \multimap^{\top, o+1} \mathbf{1} \\
\\
\text{(c)} = \text{send} : (\overline{S} + \overline{S'}) \times (S \oplus^o S') \multimap^{\top, o} \text{end}_!^{o+1}
\end{array}
\end{array}$$

Figure 4: Typing Rules for Syntactic Sugar for PGV (T-SELECT-INL and T-SELECT-INR).

$$\begin{array}{c}
\text{T-OFFER} \\
\frac{\Gamma \vdash^p L : S \&^o S' \quad \Delta, x : S \vdash^q M : T \quad \Delta, y : S' \vdash^q N : T \quad o \sqcup p < \text{pr}(\Delta, S, S')}{\Gamma, \Delta \vdash^{o \sqcup p \sqcup q} \mathbf{offer} L \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : T} \triangleq \\
\\
\text{(c)} \quad \frac{\frac{\text{(B)} \quad \overline{w : \mathbf{end}_?^{o+1} \vdash^\perp w : \mathbf{end}_?^{o+1}}}{w : \mathbf{end}_?^{o+1} \vdash^o \mathbf{wait} w : \mathbf{1}}}{\frac{\frac{z : S + S' \vdash^\perp z : S + S' \quad \Delta, x : S \vdash^q M : T \quad \Delta, y : S' \vdash^q N : T}{\Delta, z : S + S' \vdash^q \mathbf{case} z \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : T} \quad o < \text{pr}(\Delta, S + S')}{\Delta, z : S + S', w : \mathbf{end}_?^{o+1} \vdash^{o \sqcup q} \mathbf{wait} w; \mathbf{case} z \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : T}} \\
\\
\frac{\text{(A)} \quad \Gamma \vdash^p L : ?^o(S + S').\mathbf{end}_?^{o+1}}{\Gamma \vdash^{o \sqcup p} \mathbf{recv} L : (S + S') \times \mathbf{end}_?^{o+1}} \quad \text{(c)} \quad o \sqcup p < \text{pr}(\Delta) \\
\hline
\Gamma, \Delta \vdash^{o \sqcup p \sqcup q} \mathbf{let} (z, w) = \mathbf{recv} L \mathbf{in} \mathbf{wait} w; \mathbf{case} z \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : T \\
\\
\text{(A)} = \mathbf{recv} : ?^o(S + S').\mathbf{end}_?^{o+1} \multimap^{\top, o} (S + S') \times \mathbf{end}_?^{o+1} \quad \text{(B)} = \mathbf{wait} : \mathbf{end}_?^{o+1} \multimap^{\top, o} \mathbf{1} \\
\\
\text{T-OFFER-ABSURD} \\
\frac{\Gamma \vdash^p L : \&^o \{ \} \quad o \sqcup p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{o \sqcup p} \mathbf{offer} L \{ \} : T} \triangleq \\
\\
\frac{\text{(A)} \quad \Gamma \vdash^p L : ?^o \mathbf{0}.\mathbf{end}_?^{o+1}}{\Gamma \vdash^{o \sqcup p} \mathbf{recv} L : \mathbf{0} \times \mathbf{end}_?^{o+1}} \\
\\
\text{(B)} \quad \frac{\overline{w : \mathbf{end}_?^{o+1} \vdash^\perp w : \mathbf{end}_?^{o+1}} \quad \overline{z : \mathbf{0} \vdash^\perp z : \mathbf{0}}}{\frac{w : \mathbf{end}_?^{o+1} \vdash^o \mathbf{wait} w : \mathbf{1} \quad \Delta, z : \mathbf{0} \vdash^\perp \mathbf{absurd} z : T \quad o < \text{pr}(\Delta)}{\Delta, z : \mathbf{0}, w : \mathbf{end}_?^{o+1} \vdash^o \mathbf{wait} w; \mathbf{absurd} z : T} \quad o \sqcup p < \text{pr}(\Delta)} \\
\hline
\Gamma, \Delta \vdash^{o \sqcup p} \mathbf{let} (z, w) = \mathbf{recv} L \mathbf{in} \mathbf{wait} w; \mathbf{absurd} z : T \\
\\
\text{(A)} = \mathbf{recv} : ?^o \mathbf{0}.\mathbf{end}_?^{o+1} \multimap^{\top, o} \mathbf{0} \times \mathbf{end}_?^{o+1} \quad \text{(B)} = \mathbf{wait} : \mathbf{end}_?^{o+1} \multimap^{\top, o} \mathbf{1}
\end{array}$$

Figure 5: Typing Rules for Syntactic Sugar for PGV (T-OFFER and T-OFFER-ABSURD).

Syntactic Sugar Typing. The following typing rules given in Figs. 3 to 5, cover syntactic sugar typing for PGV.

3. TECHNICAL DEVELOPMENTS

3.1. Subject Reduction. Unlike with previous versions of GV, structural congruence, term reduction, and configuration reduction are all type preserving.

We must show that substitution preserves priority constraints. For this, we prove Lemma 3.1, which shows that values have finished all their communication, and that any priorities in the type of the value come from the typing environment.

Lemma 3.1. *If $\Gamma \vdash^p V : T$, then $p = \perp$, and $\text{pr}(\Gamma) = \text{pr}(T)$.*

Proof. By induction on the derivation of $\Gamma \vdash^o V : T$.

Case T-VAR. *Immediately.*

$$\overline{x : T \vdash^{\perp} x : T}$$

Case T-CONST. *Immediately.*

$$\overline{\emptyset \vdash^{\perp} K : T}$$

Case T-LAM. *Immediately.*

$$\frac{\Gamma, x : T \vdash^q M : U}{\Gamma \vdash^{\perp} \lambda x. M : T \multimap^{\text{pr}(\Gamma), q} U}$$

Case T-UNIT. *Immediately.*

$$\overline{\emptyset \vdash^{\perp} () : \mathbf{1}}$$

Case T-PAIR. *The induction hypotheses give us $p = q = \perp$, hence $p \sqcup q = \perp$, and $\text{pr}(\Gamma) = \text{pr}(T)$ and $\text{pr}(\Delta) = \text{pr}(U)$, hence $\text{pr}(\Gamma, \Delta) = \text{pr}(\Gamma) \sqcap \text{pr}(\Delta) = \text{pr}(T) \sqcap \text{pr}(U) = \text{pr}(T \times U)$.*

$$\frac{\Gamma \vdash^p V : T \quad \Delta \vdash^q W : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta \vdash^{p \sqcup q} (V, W) : T \times U}$$

Case T-INL. *The induction hypothesis gives us $p = \perp$, and $\text{pr}(\Gamma) = \text{pr}(T)$. We know $\text{pr}(T) = \text{pr}(U)$, hence $\text{pr}(\Gamma) = \text{pr}(T + U)$.*

$$\frac{\Gamma \vdash^p V : T \quad \text{pr}(T) = \text{pr}(U)}{\Gamma \vdash^p \text{inl } V : T + U}$$

Case T-INR. *The induction hypothesis gives us $p = \perp$, and $\text{pr}(\Gamma) = \text{pr}(U)$. We know $\text{pr}(T) = \text{pr}(U)$, hence $\text{pr}(\Gamma) = \text{pr}(T + U)$.*

$$\frac{\Gamma \vdash^p V : U \quad \text{pr}(T) = \text{pr}(U)}{\Gamma \vdash^p \text{inr } V : T + U}$$

□

Lemma 3.2. *If $\Gamma, x : U' \vdash^p M : T$ and $\Theta \vdash^q V : U'$, then $\Gamma, \Theta \vdash^p M\{V/x\} : T$.*

Proof. By induction on the derivation of $\Gamma, x : U' \vdash^p M : T$.

Case T-VAR. *By Lemma 3.1, $q = \perp$.*

$$\overline{x : U' \vdash^{\perp} x : U'} \xrightarrow{\{V/x\}} \Theta \vdash^{\perp} V : U'$$

Case T-LAM. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Gamma, \Theta) = \text{pr}(\Gamma, U')$.

$$\frac{\Gamma, x : U', y : T \vdash^q M : U}{\Gamma, x : U' \vdash^\perp \lambda y. M : T \multimap^{\text{pr}(\Gamma, U'), q} U} \xrightarrow{\{V/x\}} \frac{\Gamma, \Theta, y : T \vdash^q M\{V/x\} : U}{\Gamma, \Theta \vdash^\perp \lambda y. M\{V/x\} : T \multimap^{\text{pr}(\Gamma, \Theta), q} U}$$

Case T-APP. There are two subcases:

Subcase $x \in M$. Immediately, from the induction hypothesis.

$$\frac{\Gamma, x : U' \vdash^p M : T \multimap^{p', q'} U \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta) \quad q < p'}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q \sqcup q'} M N : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma, \Theta \vdash^p M\{V/x\} : T \multimap^{p', q'} U \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta) \quad q < p'}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q \sqcup q'} (M\{V/x\}) N : U}$$

Subcase $x \in N$. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Delta, \Theta) = \text{pr}(\Delta, U')$.

$$\frac{\Gamma \vdash^p M : T \multimap^{p', q'} U \quad \Delta, x : U' \vdash^q N : T \quad p < \text{pr}(\Delta, U') \quad q < p'}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q \sqcup q'} M N : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma \vdash^p M : T \multimap^{p', q'} U \quad \Delta, \Theta \vdash^q N\{V/x\} : T \quad p < \text{pr}(\Delta, \Theta) \quad q < p'}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q \sqcup q'} M (N\{V/x\}) : U}$$

Case T-LETUNIT. There are two subcases:

Subcase $x \in M$. Immediately, from the induction hypothesis.

$$\frac{\Gamma, x : U' \vdash^p M : \mathbf{1} \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta)}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} M; N : T} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma, \Theta \vdash^p M\{V/x\} : \mathbf{1} \quad \Delta \vdash^q N : T \quad p < \text{pr}(\Delta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} M\{V/x\}; N : T}$$

Subcase $x \in N$. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Delta, \Theta) = \text{pr}(\Delta, U')$.

$$\frac{\Gamma \vdash^p M : \mathbf{1} \quad \Delta, x : U' \vdash^q N : T \quad p < \text{pr}(\Delta, U')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} M; N : T} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma \vdash^p M : \mathbf{1} \quad \Delta, \Theta \vdash^q N\{V/x\} : T \quad p < \text{pr}(\Delta, \Theta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} M; N\{V/x\} : T}$$

Case T-PAIR. There are two subcases:

Subcase $x \in M$. Immediately, from the induction hypothesis.

$$\frac{\Gamma, x : U' \vdash^p M : T \quad \Delta \vdash^q N : U \quad p < \text{pr}(\Delta, U')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} (M, N) : T \times U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma, \Theta \vdash^p M\{V/x\} : T \quad \Delta \vdash^q N : U \quad p < \text{pr}(\Delta, \Theta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} (M\{V/x\}, N) : T \times U}$$

Subcase $x \in N$. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Delta, \Theta) = \text{pr}(\Delta, U')$.

$$\frac{\Gamma \vdash^p M : T \quad \Delta, x : U' \vdash^q N : U \quad p < \text{pr}(\Delta, U')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} (M, N) : T \times U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma \vdash^p M : T \quad \Delta, \Theta \vdash^q N\{V/x\} : U \quad p < \text{pr}(\Delta, \Theta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} (M, N\{V/x\}) : T \times U}$$

Case T-LETPAIR. There are two subcases:

Subcase $x \in M$. Immediately, from the induction hypothesis.

$$\frac{\Gamma, x : U' \vdash^p M : T \times T' \quad \Delta, y : T, z : T' \vdash^q N : U \quad p < \text{pr}(\Delta, T, T')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} \text{let } (y, z) = M \text{ in } N : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma, \Theta \vdash^p M\{V/x\} : T \times T' \quad \Delta, y : T, z : T' \vdash^q N : U \quad p < \text{pr}(\Delta, T, T')}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} \text{let } (y, z) = M\{V/x\} \text{ in } N : U}$$

Subcase $x \in N$. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Delta, \Theta, T, T') = \text{pr}(\Delta, U', T, T')$.

$$\frac{\Gamma \vdash^p M : T \times T' \quad \Delta, x : U', y : T, z : T' \vdash^q N : U \quad p < \text{pr}(\Delta, U', T, T')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} \text{let } (y, z) = M \text{ in } N : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma \vdash^p M : T \times T' \quad \Delta, \Theta, y : T, z : T' \vdash^q N\{V/x\} : U \quad p < \text{pr}(\Delta, \Theta, T, T')}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} \text{let } (y, z) = M \text{ in } N\{V/x\} : U}$$

Case T-ABSURD.

$$\frac{\Gamma, x : U' \vdash^p M : \mathbf{0}}{\Gamma, \Delta, x : U' \vdash^p \text{absurd } M : T} \xrightarrow{\{V/x\}} \frac{\Gamma, \Theta \vdash^p M\{V/x\} : \mathbf{0}}{\Gamma, \Delta, \Theta \vdash^p \text{absurd } M\{V/x\} : T}$$

Case T-INL.

$$\frac{\Gamma, x : U' \vdash^p M : T \quad \text{pr}(T) = \text{pr}(U)}{\Gamma, x : U' \vdash^p \text{inl } M : T + U} \xrightarrow{\{V/x\}} \frac{\Gamma, \Theta \vdash^p M\{V/x\} : T \quad \text{pr}(T) = \text{pr}(U)}{\Gamma, \Theta \vdash^p \text{inl } M\{V/x\} : T + U}$$

Case T-INR.

$$\frac{\Gamma, x : U' \vdash^p M : U \quad \text{pr}(T) = \text{pr}(U)}{\Gamma, x : U' \vdash^p \text{inr } M : T + U} \xrightarrow{\{V/x\}} \frac{\Gamma, \Theta \vdash^p M\{V/x\} : U \quad \text{pr}(T) = \text{pr}(U)}{\Gamma, \Theta \vdash^p \text{inr } M\{V/x\} : T + U}$$

Case T-CASESUM. There are two subcases:

Subcase $x \in L$. Immediately, from the induction hypothesis.

$$\frac{\Gamma, x : U' \vdash^p L : T + T' \quad \Delta, y : T \vdash^q M : U \quad \Delta, z : T' \vdash^q N : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} \text{case } L \{ \text{inl } y \mapsto M; \text{inr } z \mapsto N \} : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma, \Theta \vdash^p L\{V/x\} : T + T' \quad \Delta, y : T \vdash^q M : U \quad \Delta, z : T' \vdash^q N : U \quad p < \text{pr}(\Delta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} \text{case } L\{V/x\} \{ \text{inl } y \mapsto M; \text{inr } z \mapsto N \} : U}$$

Subcase $x \in M$ and $x \in N$. By Lemma 3.1, $\text{pr}(\Theta) = \text{pr}(U')$, hence $\text{pr}(\Delta, \Theta, T) = \text{pr}(\Delta, U', T)$ and $\text{pr}(\Delta, \Theta, T') = \text{pr}(\Delta, U', T')$.

$$\frac{\Gamma \vdash^p L : T + T' \quad \Delta, x : U', y : T \vdash^q M : U \quad \Delta, x : U', z : T' \vdash^q N : U \quad p < \text{pr}(\Delta, U')}{\Gamma, \Delta, x : U' \vdash^{p \sqcup q} \text{case } L \{ \text{inl } y \mapsto M; \text{inr } z \mapsto N \} : U} \xrightarrow{\{V/x\}}$$

$$\frac{\Gamma \vdash^p L : T + T' \quad \Delta, \Theta, y : T \vdash^q M\{V/x\} : U \quad \Delta, \Theta, z : T' \vdash^q N\{V/x\} : U \quad p < \text{pr}(\Delta, \Theta)}{\Gamma, \Delta, \Theta \vdash^{p \sqcup q} \text{case } L \{ \text{inl } y \mapsto M\{V/x\}; \text{inr } z \mapsto N\{V/x\} \} : U}$$

We omit the cases where $x \notin M$, as they are straightforward. $\square$

Lemma 3.3. If $\Gamma \vdash^p M : T$ and $M \rightarrow_M M'$, then $\Gamma \vdash^p M' : T$.

Proof. The proof closely follows the standard proof of subject reduction for the simply-typed linear λ -calculus, as the constants are uninterpreted by the term reduction ($\rightarrow_M$) and priority constraints are maintained consequence of § 3.1.

By induction on the derivation of $M \rightarrow_M M'$.

Case E-LAM. By Lemma 3.2.

$$\frac{\frac{\Gamma, x : T \vdash^p M : U}{\Gamma \vdash^{\perp} \lambda x. M : T \multimap^{\text{pr}(\Gamma), p} U} \quad \Delta \vdash^{\perp} V : T}{\Gamma, \Delta \vdash^p (\lambda x. M) V : U} \rightarrow_M \Gamma, \Delta \vdash^p M\{V/x\} : U$$

Case E-UNIT. By Lemma 3.2.

$$\frac{\overline{\emptyset \vdash^{\perp} () : \mathbf{1}} \quad \Gamma \vdash^p M : T}{\Gamma \vdash^p (); M : T} \rightarrow_M \Gamma \vdash^p M : T$$

Case E-PAIR. By Lemma 3.2.

$$\frac{\frac{\Gamma \vdash^{\perp} V : T \quad \Delta \vdash^{\perp} W : T'}{\Gamma, \Delta \vdash^{\perp} (V, W) : T \times T'} \quad \Theta, x : T, y : T' \vdash^p M : U}{\Gamma, \Delta, \Theta \vdash \text{let } (x, y) = (V, W) \text{ in } M : U}$$

$\downarrow$

$$\Gamma, \Delta, \Theta \vdash^p M\{V/x\}\{W/y\} : U$$

Case E-INL. *By Lemma 3.2.*

$$\begin{array}{c}
 \frac{\Gamma \vdash^{\perp} V : T}{\Gamma \vdash^{\perp} \mathbf{inl} V : T + T'} \quad \Delta, x : T \vdash^p M : U \quad \Delta, y : T' \vdash^p N : U \\
 \hline
 \Gamma, \Delta \vdash^p \mathbf{case} \mathbf{inl} V \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : U \\
 \\
 \downarrow \mathcal{N} \\
 \Gamma, \Delta \vdash^p M\{V/x\} : U
 \end{array}$$

Case E-INR. *By Lemma 3.2.*

$$\begin{array}{c}
 \frac{\Gamma \vdash^{\perp} V : T'}{\Gamma \vdash^{\perp} \mathbf{inr} V : T + T'} \quad \Delta, x : T \vdash^p M : U \quad \Delta, y : T' \vdash^p N : U \\
 \hline
 \Gamma, \Delta \vdash^p \mathbf{case} \mathbf{inr} V \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} : U \\
 \\
 \downarrow \mathcal{N} \\
 \Gamma, \Delta \vdash^p N\{V/y\} : U
 \end{array}$$

Case E-LIFT. *Immediately by induction on the evaluation context E .*

□

Lemma 3.4. *If $\Gamma \vdash^{\phi} \mathcal{C}$ and $\mathcal{C} \equiv \mathcal{C}'$, then $\Gamma \vdash^{\phi} \mathcal{C}'$.*

Proof. By induction on the derivation of $\mathcal{C} \equiv \mathcal{C}'$.

Case SC-LINKSWAP.

$$\begin{array}{c}
 \frac{\mathbf{link} : S \times \bar{S} \multimap \mathbf{1} \quad \frac{\frac{x : S \vdash^{\perp} x : S}{x : S, y : \bar{S} \vdash^{\perp} (x, y) : S \times \bar{S}} \quad \frac{y : \bar{S} \vdash^{\perp} y : \bar{S}}{x : S, y : \bar{S} \vdash^{\perp} \mathbf{link} (x, y) : \mathbf{1}}}{x : S, y : \bar{S} \vdash^{\perp} \mathbf{link} (x, y) : \mathbf{1}} \\
 \vdots \\
 \hline
 \Gamma, x : S, y : \bar{S} \vdash^{\phi} \mathcal{F}[\mathbf{link} (x, y)] \\
 \\
 ||| \\
 \frac{\mathbf{link} : \bar{S} \times S \multimap \mathbf{1} \quad \frac{\frac{y : \bar{S} \vdash^{\perp} y : \bar{S}}{x : S, y : \bar{S} \vdash^{\perp} (y, x) : S \times \bar{S}} \quad \frac{x : S \vdash^{\perp} x : S}{x : S, y : \bar{S} \vdash^{\perp} \mathbf{link} (y, x) : \mathbf{1}}}{x : S, y : \bar{S} \vdash^{\perp} \mathbf{link} (y, x) : \mathbf{1}} \\
 \vdots \\
 \hline
 \Gamma, x : S, y : \bar{S} \vdash^{\phi} \mathcal{F}[\mathbf{link} (y, x)]
 \end{array}$$

Case SC-RESLINK.

$$\begin{array}{c}
\frac{\text{link} : S \times \bar{S} \multimap \mathbf{1} \quad \frac{\frac{x : S \vdash^\perp x : S \quad y : \bar{S} \vdash^\perp y : \bar{S}}{x : S, y : \bar{S} \vdash^\perp (x, y) : S \times \bar{S}}}{x : S, y : \bar{S} \vdash^\perp \text{link}(x, y) : \mathbf{1}}}{x : S, y : \bar{S} \vdash^\phi \phi \text{ link}(x, y)} \\
\hline
\frac{}{\emptyset \vdash^\phi (\nu xy)(\phi \text{ link}(x, y))} \equiv \frac{}{\emptyset \vdash^\phi \phi()} : \mathbf{1}
\end{array}$$

Case SC-RESSWAP.

$$\frac{\Gamma, x : S, y : \bar{S} \vdash^\phi \mathcal{C}}{\Gamma \vdash^\phi (\nu xy)\mathcal{C}} \equiv \frac{\Gamma, x : S, y : \bar{S} \vdash^\phi \mathcal{C}}{\Gamma \vdash^\phi (\nu yx)\mathcal{C}}$$

Case SC-RESCOMM.

$$\frac{\frac{\Gamma, x : S, y : \bar{S}, z : S', w : \bar{S}' \vdash^\phi \mathcal{C}}{\Gamma, x : S, y : \bar{S} \vdash^\phi (\nu zw)\mathcal{C}}}{\Gamma \vdash^\phi (\nu xy)(\nu zw)\mathcal{C}} \equiv \frac{\frac{\Gamma, x : S, y : \bar{S}, z : S', w : \bar{S}' \vdash^\phi \mathcal{C}}{\Gamma, z : S', w : \bar{S}' \vdash^\phi (\nu xy)\mathcal{C}}}{\Gamma \vdash^\phi (\nu zw)(\nu xy)\mathcal{C}}$$

Case SC-RESEXT.

$$\frac{\frac{\Gamma \vdash^\phi \mathcal{C} \quad \Delta, x : S, y : \bar{S} \vdash^\phi \mathcal{D}}{\Gamma, \Delta, x : S, y : \bar{S} \vdash^\phi (\mathcal{C} \parallel \mathcal{D})}}{\Gamma, \Delta \vdash^\phi (\nu xy)(\mathcal{C} \parallel \mathcal{D})} \equiv \frac{\Gamma \vdash^\phi \mathcal{C} \quad \frac{\Delta, x : S, y : \bar{S} \vdash^\phi \mathcal{D}}{\Delta \vdash^\phi (\nu xy)\mathcal{D}}}{\Gamma, \Delta \vdash^\phi \mathcal{C} \parallel (\nu xy)\mathcal{D}}$$

Case SC-PARNIL.

$$\frac{\Gamma \vdash^\phi \mathcal{C} \quad \frac{\frac{}{\emptyset \vdash^\perp () : \mathbf{1}}{\emptyset \vdash^\circ \circ()}}{\Gamma \vdash^\phi \mathcal{C} \parallel \circ()}}{\Gamma \vdash^\phi \mathcal{C}} \equiv \Gamma \vdash^\phi \mathcal{C}$$

Case SC-PARCOMM.

$$\frac{\Gamma \vdash^\phi \mathcal{C} \quad \Delta \vdash^{\phi'} \mathcal{D}}{\Gamma, \Delta \vdash^{\phi+\phi'} (\mathcal{C} \parallel \mathcal{D})} \equiv \frac{\Delta \vdash^{\phi'} \mathcal{D} \quad \Gamma \vdash^\phi \mathcal{C}}{\Gamma, \Delta \vdash^{\phi'+\phi} (\mathcal{D} \parallel \mathcal{C})}$$

Case SC-PARASSOC.

$$\frac{\Gamma \vdash^\phi \mathcal{C} \quad \frac{\frac{\Delta \vdash^{\phi'} \mathcal{D} \quad \Theta \vdash^{\phi''} \mathcal{E}}{\Delta, \Theta \vdash^{\phi'+\phi''} (\mathcal{D} \parallel \mathcal{E})}}{\Gamma, \Delta, \Theta \vdash^{\phi+\phi'+\phi''} \mathcal{C} \parallel (\mathcal{D} \parallel \mathcal{E})}}{\Gamma, \Delta, \Theta \vdash^{\phi+\phi'+\phi''} (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}} \equiv \frac{\frac{\Gamma \vdash^\phi \mathcal{C} \quad \Delta \vdash^{\phi'} \mathcal{D}}{\Gamma, \Delta \vdash^{\phi+\phi'} (\mathcal{C} \parallel \mathcal{D})} \quad \Theta \vdash^{\phi''} \mathcal{E}}{\Gamma, \Delta, \Theta \vdash^{\phi+\phi'+\phi''} (\mathcal{C} \parallel \mathcal{D}) \parallel \mathcal{E}}$$

□

Theorem 3.5. *If $\Gamma \vdash^\phi \mathcal{C}$ and $\mathcal{C} \longrightarrow_{\mathcal{C}} \mathcal{C}'$, then $\Gamma \vdash^\phi \mathcal{C}'$.*

Proof. By induction on the derivation of $\mathcal{C} \longrightarrow_{\mathcal{C}} \mathcal{C}'$.

Case E-NEW.

$$\begin{array}{c}
 \frac{\frac{\overline{\text{new} : \mathbf{1} \multimap S \times \bar{S}} \quad \overline{\emptyset \vdash^\perp () : \mathbf{1}}}{\emptyset \vdash^\perp \text{new} () : S \times \bar{S}} \quad \vdots}{\Gamma \vdash^\phi \mathcal{F}[\text{new} ()]} \quad \rightarrow_{\mathcal{C}} \quad \frac{\frac{\overline{x : S \vdash^\perp x : S} \quad \overline{y : \bar{S} \vdash^\perp y : \bar{S}}}{x : S, y : \bar{S} \vdash^\perp (x, y) : S \times \bar{S}} \quad \vdots}{\Gamma, x : S, y : \bar{S} \vdash^\phi \mathcal{F}[(x, y)]} \\
 \Gamma \vdash^\phi (\nu xy) \mathcal{F}[(x, y)]
 \end{array}$$

Case E-SPAWN.

$$\begin{array}{c}
 \frac{\overline{\text{spawn} : (\mathbf{1} \multimap^{p,q} \mathbf{1}) \multimap \mathbf{1} \quad \Delta \vdash^\perp V : \mathbf{1} \multimap^{p,q} \mathbf{1}}}{\Delta \vdash^\perp \text{spawn} V : \mathbf{1}} \quad \vdots \\
 \hline
 \Gamma, \Delta \vdash^\phi \mathcal{F}[\text{spawn} V] \\
 \downarrow \\
 \frac{\frac{\overline{\emptyset \vdash^\perp () : \mathbf{1}} \quad \vdots}{\Gamma \vdash^\phi \mathcal{F}[]]}{\Gamma, \Delta \vdash^\phi \mathcal{F}[] \parallel \circ (V ())} \quad \frac{\Delta \vdash^\perp V : \mathbf{1} \multimap^{p,q} \mathbf{1} \quad \overline{\emptyset \vdash^\perp () : \mathbf{1}}}{\Delta \vdash^q V () : \mathbf{1}} \\
 \Delta \vdash^\circ \circ (V ())
 \end{array}$$

Case E-SEND. *See Fig. 6.*

Case E-CLOSE.

$$\begin{array}{c}
 \frac{\overline{\text{close} : \text{end}_!^\circ \multimap^{\top, \circ} \mathbf{1}} \quad \overline{x : \text{end}_!^\circ \vdash^\perp x : \text{end}_!^\circ}}{x : \text{end}_!^\circ \vdash^\circ \text{close} x : \mathbf{1}} \quad \vdots \quad \frac{\overline{\text{wait} : \text{end}_?^\circ \multimap^{\top, \circ} \mathbf{1}} \quad \overline{y : \text{end}_?^\circ \vdash^\perp y : \text{end}_?^\circ}}{y : \text{end}_?^\circ \vdash^\circ \text{wait} y : \mathbf{1}} \quad \vdots \\
 \hline
 \Gamma, x : \text{end}_!^\circ \vdash^\phi \mathcal{F}[\text{close} x] \quad \Delta, y : \text{end}_?^\circ \vdash^{\phi'} \mathcal{F}'[\text{wait} y] \\
 \hline
 \Gamma, \Delta, x : \text{end}_!^\circ, y : \text{end}_?^\circ \vdash^{\phi+\phi'} \mathcal{F}[\text{close} x] \parallel \mathcal{F}'[\text{wait} y] \\
 \hline
 \Gamma, \Delta \vdash^{\phi+\phi'} (\nu xy) (\mathcal{F}[\text{close} x] \parallel \mathcal{F}'[\text{wait} y]) \\
 \downarrow \\
 \frac{\frac{\overline{\emptyset \vdash^\perp () : \mathbf{1}} \quad \vdots}{\Gamma \vdash^\phi \mathcal{F}[]]}{\Gamma, \Delta \vdash^{\phi+\phi'} \mathcal{F}[] \parallel \mathcal{F}'[]} \quad \frac{\frac{\overline{\emptyset \vdash^\perp () : \mathbf{1}} \quad \vdots}{\Delta \vdash^{\phi'} \mathcal{F}'[]}}{\Delta \vdash^{\phi'} \mathcal{F}'[]}
 \end{array}$$

Case E-LIFTC. *By induction on the evaluation context $\mathcal{G}$.*

Case E-LIFTM. *By Lemma 3.3.*

Case E-LIFTSC. *By Lemma 3.4.*

□

$$\begin{array}{c}
\frac{\frac{\frac{}{(A)} \quad \frac{}{(B)}}{\Gamma, \Delta, \Theta, x : !^o T.S, y : ?^o T.\bar{S} \vdash^{\phi+\phi'} \mathcal{F}[\text{send}(V, x)] \parallel \mathcal{F}'[\text{recv } y]}}{\Gamma, \Delta, \Theta \vdash^{\phi+\phi'} (\nu xy)(\mathcal{F}[\text{send}(V, x)] \parallel \mathcal{F}'[\text{recv } y])}} \\
\\
\frac{\frac{\text{send} : T \times !^o T.S \multimap^{\top, o} \bar{S} \quad \frac{\Delta \vdash^p V : T \quad \overline{x : !^o T.S \vdash^\perp x : !^o T.S}}{\Delta, x : !^o T.S \vdash^p (V, x) : T \times !^o T.S}}}{\Delta, x : !^o T.S \vdash^{p \sqcup o} \text{send}(V, x) : S} \\
\vdots \\
(A) \triangleq \frac{}{\Gamma, \Delta, x : !^o T.S \vdash^\phi \mathcal{F}[\text{send}(V, x)]} \\
\\
\frac{\frac{\text{recv} : ?^o T.\bar{S} \multimap^{\top, o} T \times \bar{S} \quad \overline{y : ?^o T.\bar{S} \vdash^\perp y : ?^o T.\bar{S}}}{y : ?^o T.\bar{S} \vdash^o \text{recv } y : T \times \bar{S}}}{\vdots} \\
(B) \triangleq \frac{}{\Theta, y : ?^o T.\bar{S} \vdash^{\phi'} \mathcal{F}'[\text{recv } y]} \\
\\
\downarrow \rightsquigarrow \\
\frac{\frac{\frac{x : S \vdash^\perp x : S}{\vdots}}{\Gamma, x : S \vdash^\phi \mathcal{F}[x]} \quad \frac{\frac{\Delta \vdash^p V : T \quad \overline{\Delta, y : \bar{S} \vdash^\perp y : \bar{S}}}{\Delta, y : \bar{S} \vdash^p (V, y) : T \times \bar{S}}}{\vdots}}{\frac{\Delta, \Theta, y : \bar{S} \vdash^{\phi'} \mathcal{F}'[(V, y)]}{\Gamma, \Delta, \Theta, x : S, y : \bar{S} \vdash^{\phi+\phi'} \mathcal{F}[x] \parallel \mathcal{F}'[(V, y)]}} \\
\frac{}{\Gamma, \Delta, \Theta \vdash^{\phi+\phi'} (\nu xy)(\mathcal{F}[x] \parallel \mathcal{F}'[(V, y)])}
\end{array}$$

Figure 6: Subject Reduction (E-SEND)

3.2. Progress and Deadlock Freedom. PGV satisfies progress, as PGV configurations either reduce or are in normal form. However, the normal forms may seem surprising at first, as evaluating a well-typed PGV term does not necessarily produce just a value. If a term returns an endpoint, then its normal form contains a thread which is ready to communicate on the dual of that endpoint. This behaviour is not new to PGV.

Let us consider an example, adapted from Lindley and Morris [LM15], in which a term returns an endpoint linked to an echo server. The echo server receives a value and sends it back unchanged. Consider the program which creates a new channel, with endpoints x and

x' , spawns off an echo server listening on x , and then returns x' :

- $\text{let } (x, x') = \text{new } () \text{ in } \text{spawn } (\lambda().\text{echo}_x); x'$ $\text{echo}_x \triangleq \text{let } (y, x) = \text{recv } x \text{ in } \text{let } x = \text{send } (y, x) \text{ in close } x$

If we reduce the above program, we get $(\nu x x')(\circ \text{echo}_x \parallel \bullet x')$. Clearly, no more evaluation is possible, even though the configuration contains the thread $\circ \text{echo}_x$, which is blocked on x . In Corollary 3.14 we will show that if a term does not return an endpoint, it must produce only a value.

Actions are terms which perform communication actions and which synchronise between two threads.

Definition 3.6. A term *acts on* an endpoint x if it is $\text{send } (V, x)$, $\text{recv } x$, $\text{close } x$, or $\text{wait } x$. A term is an *action* if it acts on some endpoint x .

Ready terms are terms which perform communication actions, either by themselves, *e.g.*, creating a new channel or thread, or with another thread, *e.g.*, sending or receiving. It is worth mentioning that the notion of readiness presented here is akin to live processes introduced by Caires and Pfenning [CP10, DP22], and poised processes introduced by Pfenning and Griffith [PG15] and later used by Balzer *et al.* [BP17, BTP19]. Ready processes like live/poised processes denote processes that are ready to communicate on their providing channel.

Definition 3.7. A term L is *ready* if it is of the form $E[M]$, where M is of the form new , $\text{spawn } N$, $\text{link } (x, y)$, or M acts on x . In the latter case, we say that L is *ready to act on* x or is *blocked on*.

Progress for the term language is standard for GV, and deviates from progress for linear λ -calculus only in that terms may reduce to values or *ready terms*, where the definition of ready terms encompasses all terms whose reduction is stuck on some constant K .

Lemma 3.8. If $\Gamma \vdash^P M : T$ and Γ contains only session types, then: (i) M is a value; (ii) $M \longrightarrow_M N$ for some N ; or (iii) M is ready.

With “ Γ contains only session types” we mean that for every $x : T \in \Gamma$, T is a session type, *i.e.*, is of the form S .

Canonical forms deviate from those for GV, in that we opt to move all ν -binders to the top. The standard GV canonical form, alternating ν -binders and their corresponding parallel compositions, does not work for PGV, since multiple channels may be split across a single parallel composition.

A configuration either reduces, or it is equivalent to configuration in normal form. Crucial to the normal form is that each term M_i is blocked on the corresponding channel x_i , and hence no two terms act on dual endpoints. Furthermore, no term M_i can perform a communication action by itself, since those are excluded by the definition of actions. Finally, as a corollary, we get that well-typed terms which do not return endpoints return just a value:

Definition 3.9. A configuration $\mathcal{C}$ is in canonical form if it is of the form $(\nu x_1 x'_1) \dots (\nu x_n x'_n)(\circ M_1 \parallel \dots \parallel \circ M_m \parallel \bullet N)$ where no term M_i is a value.

Lemma 3.10. If $\Gamma \vdash^\bullet \mathcal{C}$, there exists some $\mathcal{D}$ such that $\mathcal{C} \equiv \mathcal{D}$ and $\mathcal{D}$ is in canonical form.

Proof. We move any ν -binders to the top using SC-RESEXT, discard any superfluous occurrences of $\circ ()$ using SC-PARNIL, and move the main thread to the rightmost position using SC-PARCOMM and SC-PARASSOC. $\square$

Definition 3.11. A configuration $\mathcal{C}$ is in normal form if it is of the form $(\nu x_1 x'_1) \dots (\nu x_n x'_n)(\circ M_1 \parallel \dots \parallel \circ M_m \parallel \bullet V)$ where each M_i is ready to act on x_i .

Lemma 3.12. If $\Gamma \vdash^p L : T$ is ready to act on $x : S \in \Gamma$, then the priority bound p is some priority $\circ$, i.e., not $\perp$ or $\top$.

Proof. Let $L = E[M]$. By induction on the structure of E . M has priority $\text{pr}(S)$, and each constructor of the evaluation context E passes on the *maximum* of the priorities of its premises. No rule introduces the priority bound $\top$ on the sequent. $\square$

Theorem 3.13. If $\emptyset \vdash^\bullet \mathcal{C}$ and $\mathcal{C}$ is in canonical form, then either $\mathcal{C} \rightarrow_{\mathcal{C}} \mathcal{D}$ for some $\mathcal{D}$; or $\mathcal{C} \equiv \mathcal{D}$ for some $\mathcal{D}$ in normal form.

Proof. Let $\mathcal{C} = (\nu x_1 x'_1) \dots (\nu x_n x'_n)(\circ M_1 \parallel \dots \parallel \circ M_m \parallel \bullet N)$. We apply Lemma 3.8 to each M_i and N . If for any M_i or N we obtain a reduction $M_i \rightarrow_M M'_i$ or $N \rightarrow_M N'$, we apply E-LIFTM and E-LIFTC to obtain a reduction on $\mathcal{C}$. Otherwise, each term M_i is ready, and N is either ready or a value. Pick the *ready* term $L \in \{M_1, \dots, M_m, N\}$ with the smallest priority bound.

- (1) If L is a new $E[\text{new } ()]$, we apply E-NEW.
- (2) If L is a spawn $E[\text{spawn } M]$, we apply E-SPAWN.
- (3) If L is a link $E[\text{link } (y, z)]$ or $E[\text{link } (z, y)]$, we apply E-LINK.
- (4) Otherwise, L is ready to act on some endpoint $y : S$. Let $y' : \bar{S}$ be the dual endpoint of y . The typing rules enforce the linear use of endpoints, so there must be a term $L' \in \{M_1, \dots, M_m, N\}$ which uses y' . L' must be either a ready term or a value:
 - (a) L' is ready. By Lemma 3.12, the priority of L is $\text{pr}(S)$. By duality, $\text{pr}(\bar{S}) = \text{pr}(S)$. We cannot have $L = L'$, otherwise the action on y' would be guarded by the action on y , requiring $\text{pr}(\bar{S}) < \text{pr}(S)$.
The term L' must be ready to act on y' , otherwise the action y' would be guarded by another action with priority smaller than $\text{pr}(S)$, which contradicts our choice of L as having the smallest priority.
Therefore, we have two terms ready to act on dual endpoints. We apply the appropriate reduction rule, i.e., E-SEND or E-CLOSE.
 - (b) $L' = N$ and is a value. We rewrite $\mathcal{C}$ to put L in the position corresponding to the endpoint it is blocked on, using SC-PARCOMM, SC-PARASSOC, and optionally SC-RESSWAP. We then repeat the steps above with the term with the next smallest priority, until either we find a reduction, or the configuration has reached the desired normal form.
The argument based on the priority being the smallest continues to hold, since we know that neither L nor L' will be picked, and no other term uses y or y' .

$\square$

Corollary 3.14. If $\emptyset \vdash^\phi \mathcal{C}$, $\mathcal{C} \not\rightarrow_{\mathcal{C}}$, and $\mathcal{C}$ contains no endpoints, then $\mathcal{C} \equiv \phi V$ for some value V .

An immediate consequence of Theorem 3.13 and Corollary 3.14 is that a term which does not return an endpoint will complete all its communication actions, thus satisfying deadlock freedom.

4. RELATION TO PRIORITY CP

Thus far we have presented Priority GV (PGV) together with the relevant technical results. We remind the reader that this line of work of adding priorities, started with Priority CP (PCP) [DG18a] where priorities are integrated in Wadler's Classical Processes (CP), which is a π -calculus leveraging the correspondence of session types as linear logic propositions [Wad12]. In his work, Wadler presents a connection (via encoding) of CP and GV. Following that work, we sat out to understand the connection between the priority versions of CP and GV, thus comparing PGV and PCP. Before presenting our formal results, we will revisit PCP in the following section.

4.1. Revisiting Priority CP.

Types. Types (A, B) in PCP are based on classical linear logic propositions, and are defined by the following grammar:

$$A, B ::= A \otimes^o B \mid A \wp^o B \mid \mathbf{1}^o \mid \perp^o \mid A \oplus^o B \mid A \&^o B \mid \mathbf{0}^o \mid \top^o$$

Each connective is annotated with a priority $o \in \mathbb{N}$.

Types $A \otimes^o B$ and $A \wp^o B$ type the endpoints of a channel over which we send or receive a channel of type A , and then proceed as type B . Types $\mathbf{1}^o$ and $\perp^o$ type the endpoints of a channel whose session has terminated, and over which we send or receive a *ping* before closing the channel. These two types act as units for $A \otimes^o B$ and $A \wp^o B$, respectively.

Types $A \oplus^o B$ and $A \&^o B$ type the endpoints of a channel over which we can receive or send a choice between two branches A or B . We have opted for a simplified version of choice and followed the original Wadler's CP [Wad14], however types $\oplus$ and $\&$ can be trivially generalised to $\oplus^o \{l_i : A_i\}_{i \in I}$ and $\&^o \{l_i : A_i\}_{i \in I}$, respectively, as in the original PCP [DG18b].

Types $\mathbf{0}^o$ and $\top^o$ type the endpoints of a channel over which we can send or receive a choice between no options. These two types act as units for $A \oplus^o B$ and $A \&^o B$, respectively.

Typing Environments. Typing environments Γ, Δ associate names to types. Environments are linear, so two environments can only be combined as Γ, Δ if their names are distinct, i.e., $\text{fv}(\Gamma) \cap \text{fv}(\Delta) = \emptyset$.

$$\Gamma, \Delta ::= \emptyset \mid \Gamma, x : A$$

Type Duality. Duality is an involutive function on types which preserves priorities:

$$\begin{array}{llll} (\mathbf{1}^o)^\perp = \perp^o & (A \otimes^o B)^\perp = A^\perp \wp^o B^\perp & (\mathbf{0}^o)^\perp = \top^o & (A \oplus^o B)^\perp = A^\perp \&^o B^\perp \\ (\perp^o)^\perp = \mathbf{1}^o & (A \wp^o B)^\perp = A^\perp \otimes^o B^\perp & (\top^o)^\perp = \mathbf{0}^o & (A \&^o B)^\perp = A^\perp \oplus^o B^\perp \end{array}$$

Priorities. The function $\text{pr}(\cdot)$ returns smallest priority of a type. As with PGV, the type system guarantees that the top-most connective always holds the smallest priority. The function $\text{pr}(\cdot)$ returns the *minimum* priority of all types a typing context, or $\top$ if the context is empty:

$$\begin{aligned} \text{pr}(\mathbf{1}^o) &= o & \text{pr}(A \otimes^o B) &= o & \text{pr}(\mathbf{0}^o) &= o & \text{pr}(A \oplus^o B) &= o \\ \text{pr}(\perp^o) &= o & \text{pr}(A \wp^o B) &= o & \text{pr}(\top^o) &= o & \text{pr}(A \&^o B) &= o \\ \text{pr}(\emptyset) &= \top & \text{pr}(\Gamma, x : T) &= \text{pr}(\Gamma) \sqcap \text{pr}(T) \end{aligned}$$

Terms. Processes (P, Q) in PCP are defined by the following grammar.

$$\begin{aligned} P, Q ::= & x \leftrightarrow y \mid (\nu xy)P \mid (P \parallel Q) \mid \mathbf{0} \\ & \mid x[y].P \mid x\Box.P \mid x(y).P \mid x().P \\ & \mid x \triangleleft \text{inl}.P \mid x \triangleleft \text{inr}.P \mid x \triangleright \{\text{inl} : P; \text{inr} : Q\} \mid x \triangleright \{\} \end{aligned}$$

Process $x \leftrightarrow y$ links endpoints x and y and forwards communication from one to the other. $(\nu xy)P$, $(P \parallel Q)$ and $\mathbf{0}$ denote respectively the restriction processes where channel endpoints x and y are bound together and with scope P , the parallel composition of processes P and Q and the terminated process.

Processes $x[y].P$ and $x(y).P$ send or receive over channel x a value y and proceed as process P . Processes $x\Box.P$ and $x().P$ send and receive an empty value—denoting the closure of channel x , and continue as P .

Processes $x \triangleleft \text{inl}.P$ and $x \triangleleft \text{inr}.P$ make a left and right choice, respectively and proceed as process P . Dually, $x \triangleright \{\text{inl} : P; \text{inr} : Q\}$ offers both left and right branches, with continuations P and Q , and $x \triangleright \{\}$ is the empty offer.

We write *unbound* send as $x\langle y \rangle.P$, which is syntactic sugar for $x[z].(y \leftrightarrow z \parallel P)$. Alternatively, we could take $x\langle y \rangle.P$ as primitive, and let $x[y].P$ be syntactic sugar for $(\nu yz)(x\langle z \rangle.P)$. CP takes *bound* sending as primitive, as it is impossible to eliminate the top-level cut in terms such as $(\nu yz)(x\langle z \rangle.P)$, even with commuting conversions. In our setting without commuting conversions and with more permissive normal forms, this is no longer an issue, but, for simplicity, we keep bound sending as primitive.

On Commuting Conversions. The main change we make to PCP is removing commuting conversions. Commuting conversions are necessary if we want our reduction strategy to correspond exactly to cut (or cycle in [DG18b]) elimination. However, as Lindley and Morris [LM15] show, all communications that can be performed *with* the use of commuting conversions, can also be performed *without* them, but using structural congruence.

From the perspective of process calculi, commuting conversions behave strangely. Consider the commuting conversion $(\kappa_{\wp})$ for $x(y).P$:

$$(\kappa_{\wp}) \quad (\nu zz')(x(y).P \parallel Q) \Longrightarrow x(y).(\nu zz')(P \parallel Q)$$

As a result of $(\kappa_{\wp})$, Q becomes blocked on $x(y)$, and any actions Q was able to perform become unavailable. Consequently, CP is non-confluent:

$$\begin{array}{ccc} (\nu xx')(a(y).P \parallel (\nu zz')(z\Box.\mathbf{0} \parallel z'().Q)) & & (\nu xx')(a(y).P \parallel (\nu zz')(z\Box.\mathbf{0} \parallel z'().Q)) \\ \Downarrow & & \Downarrow_+ \\ a(y).(\nu xx')(P \parallel (\nu zz')(z\Box.\mathbf{0} \parallel z'().Q)) & & a(y).(\nu xx')(P \parallel Q) \end{array}$$

Structural congruence.

SC-LINKSWAP	$x \leftrightarrow y$	$\equiv$	$y \leftrightarrow x$
SC-RESLINK	$(\nu xy)x \leftrightarrow y$	$\equiv$	$\mathbf{0}$
SC-RESSWAP	$(\nu xy)P$	$\equiv$	$(\nu yx)P$
SC-RESCOMM	$(\nu xy)(\nu zw)P$	$\equiv$	$(\nu zw)(\nu xy)P$
SC-RESEXT	$(\nu xy)(P \parallel Q)$	$\equiv$	$P \parallel (\nu xy)Q$, if $x, y \notin \text{fv}(P)$
SC-PARNIL	$P \parallel \mathbf{0}$	$\equiv$	P
SC-PARCOMM	$P \parallel Q$	$\equiv$	$Q \parallel P$
SC-PARASSOC	$P \parallel (Q \parallel R)$	$\equiv$	$(P \parallel Q) \parallel R$

Reduction.

E-LINK	$(\nu xy)(w \leftrightarrow x \parallel P)$	$\Longrightarrow$	$P\{w/x\}$
E-SEND	$(\nu xy)(x[z].P \parallel x(w).Q)$	$\Longrightarrow$	$(\nu xy)(\nu zw)(P \parallel Q)$
E-CLOSE	$(\nu xy)(x[] . P \parallel y().Q)$	$\Longrightarrow$	$P \parallel Q$
E-SELECT-INL	$(\nu xy)(x \triangleleft \text{inl}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\})$	$\Longrightarrow$	$(\nu xy)(P \parallel Q)$
E-SELECT-INR	$(\nu xy)(x \triangleleft \text{inr}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\})$	$\Longrightarrow$	$(\nu xy)(P \parallel R)$
E-LIFTRES	$\frac{P \Longrightarrow P'}{(\nu xy)P \Longrightarrow (\nu xy)P'}$	E-LIFTPAR	$\frac{P \Longrightarrow P'}{P \parallel Q \Longrightarrow P' \parallel Q}$
E-LIFTSC	$\frac{P \equiv P' \quad P' \Longrightarrow Q' \quad Q' \equiv Q}{P \Longrightarrow Q}$		

Figure 7: Operational Semantic for PCP.

In PCP, commuting conversions break our intuition that an action with lower priority occurs before an action with higher priority. To cite Dardha and Gay [DG18b] “if a prefix on a channel endpoint x with priority o is pulled out at top level, then to preserve priority constraints in the typing rules $[\cdot]$, it is necessary to increase priorities of all actions after the prefix on x ” by $o + 1$.

4.2. Operational Semantics. The operational semantics for PCP, given in Fig. 7, is defined as a reduction relation $\Longrightarrow$ on processes (bottom) and uses structural congruence (top). Each of the axioms of structural congruence corresponds to the axiom of the same name for PGV. We write $\Longrightarrow^+$ for the transitive closures, and $\Longrightarrow^*$ for the reflexive-transitive closures.

The reduction relation is given by a set of axioms and inference rules for context closure. Reduction occurs under restriction. E-LINK reduces a parallel composition with a link into a substitution. E-SEND is the main communication rule, where send and receive processes synchronise and reduce to the corresponding continuations. E-CLOSE follows the previous rule and it closes the channel identified by endpoints x and y . E-SELECT-INL and E-SELECT-INR are generalised versions of E-SEND. They state respectively that a left and right selection synchronises with a choice offering and reduces to the corresponding continuations. The last three rules state that reduction is closed under restriction, parallel composition and structural congruence, respectively.

4.3. Typing Rules. Figure 8 gives the typing rules for our version of PCP. A typing judgement $P \vdash \Gamma$ states that “process P is well typed under the typing context Γ ”.

$$\begin{array}{c}
\text{T-LINK} \\
\frac{}{x \leftrightarrow^A y \vdash x : A, y : A^\perp} \\
\\
\text{T-RES} \\
\frac{P \vdash \Gamma, x : A, y : A^\perp}{(\nu xy)P \vdash \Gamma} \\
\\
\text{T-PAR} \\
\frac{P \vdash \Gamma \quad Q \vdash \Delta}{P \parallel Q \vdash \Gamma, \Delta} \\
\\
\text{T-HALT} \\
\frac{}{\mathbf{0} \vdash \emptyset} \\
\\
\text{T-SEND} \\
\frac{P \vdash \Gamma, y : A, x : B \quad o < \text{pr}(\Gamma, A, B)}{x[y].P \vdash \Gamma, x : A \otimes^o B} \\
\\
\text{T-CLOSE} \\
\frac{P \vdash \Gamma \quad o < \text{pr}(\Gamma)}{x[] . P \vdash \Gamma, x : \mathbf{1}^o} \\
\\
\text{T-RECV} \\
\frac{P \vdash \Gamma, y : A, x : B \quad o < \text{pr}(\Gamma, A, B)}{x(y).P \vdash \Gamma, x : A \wp^o B} \\
\\
\text{T-WAIT} \\
\frac{P \vdash \Gamma \quad o < \text{pr}(\Gamma)}{x().P \vdash \Gamma, x : \perp^o} \\
\\
\text{T-SELECT-INL} \\
\frac{P \vdash \Gamma, x : A \quad o < \text{pr}(\Gamma, A, B) \quad \text{pr}(A) = \text{pr}(B)}{x \triangleleft \text{inl}.P \vdash \Gamma, x : A \oplus^o B} \\
\\
\text{T-SELECT-INR} \\
\frac{P \vdash \Gamma, x : B \quad o < \text{pr}(\Gamma, A, B) \quad \text{pr}(A) = \text{pr}(B)}{x \triangleleft \text{inr}.P \vdash \Gamma, x : A \oplus^o B} \\
\\
\text{T-OFFER} \\
\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B \quad o < \text{pr}(\Gamma, A, B)}{x \triangleright \{\text{inl} : P; \text{inr} : Q\} \vdash \Gamma, x : A \&^o B} \\
\\
\text{T-OFFER-ABSURD} \\
\frac{o < \text{pr}(\Gamma)}{x \triangleright \{\} \vdash \Gamma, x : \top^o}
\end{array}$$

Figure 8: Typing Rules for PCP.

T-LINK states that the link process $x \leftrightarrow y$ is well typed under channels x and y having dual types, respectively A and $A^\perp$. T-RES states that the restriction process $(\nu xy)P$ is well typed under typing context Γ if process P is well typed in Γ augmented with channel endpoints x and y having dual types, respectively A and $A^\perp$. T-PAR states that the parallel composition of processes P and Q is well typed under the disjoint union of their respective typing contexts. T-HALT states that the terminated process $\mathbf{0}$ is well typed in the empty context.

T-SEND and T-RECV state that the sending and receiving of a bound name y over a channel x is well typed under Γ and x of type $A \otimes^o B$, respectively $A \wp^o B$. Priority o is the smallest among all priorities of the types used by the output or input process, captured by the side condition $o < \text{pr}(\Gamma, A, B)$.

Rules T-CLOSE and T-WAIT type the closure of channel x and are in the same lines as the previous two rules, requiring that the priority of channel x is the smallest among all priorities in Γ .

T-SELECT-INL and T-SELECT-INR type respectively the left $x \triangleleft \text{inl}.P$ and right $x \triangleleft \text{inr}.P$ choice performed on channel x . T-OFFER and T-OFFER-ABSURD type the offering of a choice, or empty choice, on channel x . In all the above rules the priority o of channel x is the smallest with respect to the typing context $o < \text{pr}(\Gamma)$ and types involved in the choice $o < \text{pr}(\Gamma, A, B)$.

Figure 9 shows how syntactic sugar in PCP is well typed.

$$\begin{array}{c}
\text{T-UNBOUNDSEND} \\
\frac{P \vdash \Gamma, x : B \quad o < \text{pr}(\Gamma, A, B)}{x\langle y \rangle.P \vdash \Gamma, x : A \otimes B, y : A^\perp} \triangleq \\
\\
\frac{\frac{z \leftrightarrow^A y \vdash y : A^\perp, z : A \quad P \vdash \Gamma, x : B}{z \leftrightarrow^A y \parallel P \vdash \Gamma, x : B, y : A^\perp, z : A} \quad o < \text{pr}(\Gamma, A, B)}{x[z].(z \leftrightarrow^A y \parallel P) \vdash \Gamma, x : A \otimes B, y : A^\perp}
\end{array}$$

Figure 9: Typing Rules for Syntactic Sugar for PCP.

Finally, since our reduction relation is a strict subset of the reduction relation in the original [DG18b], we defer to their proof of subject reduction (Theorem 2 in [DG18b]). We prove progress for our version of PCP, see §4.5.

4.4. PCP and PLL. In this subsection, we highlight the connection between PCP and linear logic. Dardha and Gay [DG18a] present PCP—consequently also PCP given in this paper—in a way which can be viewed both as Classical Processes with restriction (T-Res) and parallel composition (T-Par) typing rules, and as a new version of linear logic, which they call Priority Linear Logic (PLL). PLL builds on CLL by replacing the cut rule with two logical rules: a mix and a cycle rule—here corresponding to T-Par and T-Res, respectively. Dardha and Gay [DG18a, §4] prove cycle-elimination, in the same lines as cut-elimination for CLL. As a corollary of cycle-elimination for PLL, we obtain deadlock freedom for PCP (Theorem 3 in [DG18a, §4]). In summary, PLL is an extension of CLL and the authors show the correspondence of PCP and PLL. Notice however, that PCP is not in correspondence with CLL itself, since processes in PCP are graphs, whether CLL induces trees.

4.5. Technical Developments.

Definition 4.1. A process acts on an endpoint x if it is $x \leftrightarrow y$, $y \leftrightarrow x$, $x[y].P$, $x(y).P$, $x[] .P$, $x().P$, $x \triangleleft \text{inl}.P$, $x \triangleleft \text{inr}.P$, $x \triangleright \{\text{inl} : P; \text{inr} : Q\}$, or $x \triangleright \{\}$. A process is an action if it acts on some endpoint x .

Definition 4.2. A process P is in canonical form if it is either $\mathbf{0}$ or of the form $(\nu x_1 x'_1) \dots (\nu x_n x'_n)(P_1 \parallel \dots \parallel P_m)$ where $m > 0$ and each P_j is an action.

Lemma 4.3. If $P \vdash \Gamma$, there exists some Q such that $P \equiv Q$ and Q is in canonical form.

Proof. If $P = \mathbf{0}$, we are done. Otherwise, we move any ν -binders to the top using SC-RESEXT, and discard any superfluous occurrences of $\mathbf{0}$ using SC-PARNIL. $\square$

The proof for progress (below) follows the same reasoning by Kobayashi [Kob06] used in the proof of deadlock freedom for closed processes (Theorem 2).

Theorem 4.4. If $P \vdash \emptyset$, then either $P = \mathbf{0}$ or there exists a Q such that $P \Longrightarrow Q$.

Proof (Sketch). This proof follows the exact same reasoning and proof sketch given by Kobayashi in [Kob06] and later adopted by Dardha and Gay for PCP in their technical report [DG18b].

By Lemma 4.3, we rewrite P to canonical form. If the resulting process is $\mathbf{0}$, we are done. Otherwise, it is of the form

$$(\nu x_1 x'_1) \dots (\nu x_n x'_n) (P_1 \parallel \dots \parallel P_m) \vdash \emptyset$$

where $m > 0$ and each $P_i \vdash \Gamma_i$ is an action.

Consider processes $P_1 \parallel \dots \parallel P_m$. Among them, we pick the process with the smallest priority $\text{pr}(\Gamma_i)$ for all i . Let this process be P_i and the priority of the top prefix be o . P_i acts on some endpoint $y : A \in \Gamma_i$. We must have $\text{pr}(\Gamma_i) = \text{pr}(A) = o$, since the other actions in P_i are guarded by the action on $y : A$, thus satisfying law (i) of priorities.

If P_i is a link $y \leftrightarrow z$ or $z \leftrightarrow y$, we apply E-LINK.

Otherwise, P_i is an input/branching or output/selection action on endpoint y of type A with priority o . Since process P is closed and consequently it respects law (ii) of priorities, there must be a co-action y' of type $A^\perp$ where y and y' are dual endpoints of the same channel (by application of rule T-RES). By duality, $\text{pr}(A) = \text{pr}(A^\perp) = o$. In the following we show that: y' is the subject of a top level action of a process P_j with $i \neq j$. This allows for the communication among P_i and P_j to happen immediately over channel endpoints y and y' .

Suppose that y' is an action not in a different parallel process P_j but rather of P_i itself. That means that the action on y' must be prefixed by the action on y , which is top level in P_i . To respect law (i) of priorities we must have $o < o$, which is absurd. This means that y' is in another parallel process P_j for $i \neq j$.

Suppose that y' in P_j is not at top level. In order to respect law (i) of priorities, it means that y' is prefixed by actions that are smaller than its priority o . This leads to a contradiction because stated that o is the smallest priority. Hence, y' must be the subject of a top level action.

We have two processes, acting on dual endpoints. We apply the appropriate reduction rule, *i.e.*, E-SEND, E-CLOSE, E-SELECT-INL, or E-SELECT-INR. $\square$

4.6. Correspondence between PGV and PCP. We illustrate the relation between PCP and PGV by defining a translation $(\cdot)_c$ from PCP processes to PGV configurations which satisfies operational correspondence.

The translation $(\cdot)_c$ translates PCP processes to PGV configurations by translating as much as possible of the π -calculus constructs from PCP to the identical configuration constructs in PGV, *i.e.*, all top-level ν and $\parallel$ constructs. When it encounters the first action, it translates the remainder of the process to a term using $(\cdot)_M$. The translation $(\cdot)_M$ translates PCP processes to PGV terms by mapping the π -calculus constructs from PCP (*e.g.*, ν , $\parallel$, $\cdot[\cdot]\cdot$, $\cdot[\cdot]\cdot$, $\cdot(\cdot)\cdot$, *etc.*) to the corresponding constants in PGV (*e.g.*, **new**, **spawn**, **send**, **recv**, *etc.*). Translating a process with $(\cdot)_c$ is the same as translating that process with $(\cdot)_M$ followed by several steps of configuration reduction. The translation $(\cdot)$ translates session types from PCP to session types in PGV matching the translation on processes.

The translation on types is defined as follows:

$$\begin{aligned} (A \otimes^o B) &= !^o \overline{(A)}.(B) & (\mathbf{1}^o) &= \mathbf{end}^o \\ (A \oplus^o B) &= (A) \oplus^o (B) & (\mathbf{0}^o) &= \oplus^o \{\} \\ (A \wp^o B) &= ?^o (A).(B) & (\perp^o) &= \mathbf{end}^o \\ (A \&^o B) &= (A) \&^o (B) & (\top^o) &= \&^o \{\} \end{aligned}$$

The translation $\langle \cdot \rangle_M$ translates processes to *terms* and maps the π -calculus constructs from PCP to the corresponding constants in PGV:

$$\begin{aligned}
\langle x \leftrightarrow y \rangle_M &= \mathbf{link} \ (x, y) \\
\langle (\nu xy) P \rangle_M &= \mathbf{let} \ (x, y) = \mathbf{new} \ () \ \mathbf{in} \ \langle P \rangle_M \\
\langle P \parallel Q \rangle_M &= \mathbf{spawn} \ (\lambda(). \langle P \rangle_M); \langle Q \rangle_M \\
\langle \mathbf{0} \rangle_M &= () \\
\langle x[] . P \rangle_M &= \mathbf{close} \ x; \langle P \rangle_M \\
\langle x(). P \rangle_M &= \mathbf{wait} \ x; \langle P \rangle_M \\
\langle x[y]. P \rangle_M &= \mathbf{let} \ (y, z) = \mathbf{new} \ () \ \mathbf{in} \ \mathbf{let} \ x = \mathbf{send} \ (z, x) \ \mathbf{in} \ \langle P \rangle_M \\
\langle x(y). P \rangle_M &= \mathbf{let} \ (y, x) = \mathbf{recv} \ x \ \mathbf{in} \ \langle P \rangle_M \\
\langle x \triangleleft \mathbf{inl}. P \rangle_M &= \mathbf{let} \ x = \mathbf{select} \ \mathbf{inl} \ x \ \mathbf{in} \ \langle P \rangle_M \\
\langle x \triangleleft \mathbf{inr}. P \rangle_M &= \mathbf{let} \ x = \mathbf{select} \ \mathbf{inr} \ x \ \mathbf{in} \ \langle P \rangle_M \\
\langle x \triangleright \{ \mathbf{inl} : P; \mathbf{inr} : Q \} \rangle_M &= \mathbf{offer} \ x \ \{ \mathbf{inl} \ x \mapsto \langle P \rangle_M; \mathbf{inr} \ x \mapsto \langle Q \rangle_M \} \\
\langle x \triangleright \{ \} \rangle_M &= \mathbf{offer} \ x \ \{ \}
\end{aligned}$$

Unfortunately, the operational correspondence along $\langle \cdot \rangle_M$ is unsound, as it translates ν -binders and parallel compositions to **new** and **spawn**, which can reduce to their equivalent configuration constructs using E-NEW and E-SPAWN. The same goes for ν -binders which are inserted when translating bound send to unbound send. For instance, the process $x[y].P$ is blocked, but its translation uses **new** and can reduce. To address this issue, we introduce a second translation, $\langle \cdot \rangle_c$, which is equivalent to translating with $\langle \cdot \rangle_M$ then reducing with E-NEW and E-SPAWN:

$$\begin{aligned}
\langle (\nu xy) P \rangle_c &= (\nu xy) \langle P \rangle_c \\
\langle P \parallel Q \rangle_c &= \langle P \rangle_c \parallel \langle Q \rangle_c \\
\langle x[y]. P \rangle_c &= (\nu yz) (\circ \mathbf{let} \ x = \mathbf{send} \ (z, x) \ \mathbf{in} \ \langle P \rangle_M) \\
\langle x \triangleleft \mathbf{inl}. P \rangle_c &= (\nu yz) (\circ \mathbf{let} \ x = \mathbf{close} \ (\mathbf{send} \ (\mathbf{inl} \ y, x)); z \ \mathbf{in} \ \langle P \rangle_M) \\
\langle x \triangleleft \mathbf{inr}. P \rangle_c &= (\nu yz) (\circ \mathbf{let} \ x = \mathbf{close} \ (\mathbf{send} \ (\mathbf{inr} \ y, x)); z \ \mathbf{in} \ \langle P \rangle_M) \\
\langle P \rangle_c &= \circ \langle P \rangle_M, \quad \text{if none of the above apply}
\end{aligned}$$

Typing environments are translated pointwise, and sequents $P \vdash \Gamma$ are translated as $\langle \Gamma \rangle \vdash^\circ \langle P \rangle_c$, where $\circ$ indicates a child thread, since translated processes do not have a main thread. The translations $\langle \cdot \rangle_M$ and $\langle \cdot \rangle_c$ preserve typing, and the latter induces a sound and complete operational correspondence.

Lemma 4.5. *If $P \vdash \Gamma$, then $\langle \Gamma \rangle \vdash^P \langle P \rangle_M : \mathbf{1}$.*

Proof. By induction on the derivation of $P \vdash \Gamma$.

Case T-LINK, T-RES, T-PAR, and T-HALT. See Fig. 10.

Case T-CLOSE, and T-WAIT. See Fig. 11.

Case T-SEND. See Fig. 12.

Case T-RECV. See Fig. 13.

$$\begin{array}{c}
\text{T-LINK} \\
\frac{x \leftrightarrow^A y \vdash x : A, y : A^\perp}{(\cdot)_M} \rightarrow \\
\\
\frac{\overline{\text{link} : (A) \times \overline{(A)} \multimap \mathbf{1}} \quad \frac{\overline{x : (A) \vdash^\perp x : (A)} \quad \overline{y : \overline{(A)} \vdash^\perp y : \overline{(A)}}}{x : (A), y : \overline{(A)} \vdash^\perp (x, y) : (A) \times \overline{(A)}}}{x : (A), y : \overline{(A)} \vdash^\perp \text{link}(x, y) : \mathbf{1}} \\
\\
\text{T-RES} \\
\frac{P \vdash \Gamma, x : A, y : A^\perp}{(\nu xy)P \vdash \Gamma} \quad (\cdot)_M \rightarrow \\
\\
\frac{\overline{\text{new} : \mathbf{1} \multimap (A) \times \overline{(A)}} \quad \overline{\emptyset \vdash^\perp () : \mathbf{1}}}{\emptyset \vdash^\perp \text{new}() : (A) \times \overline{(A)}} \quad \frac{(\Gamma), x : (A), y : \overline{(A)} \vdash^p (P)_M : \mathbf{1}}{(\Gamma) \vdash^p \text{let}(x, y) = \text{new}() \text{ in } (P)_M : \mathbf{1}} \\
\\
\text{T-PAR} \\
\frac{P \vdash \Gamma \quad Q \vdash \Delta}{P \parallel Q \vdash \Gamma, \Delta} \quad (\cdot)_M \rightarrow \\
\\
\frac{\overline{\text{spawn} : (\mathbf{1} \multimap^{\text{pr}(\Gamma), p} \mathbf{1}) \multimap \mathbf{1}} \quad \frac{(\Gamma) \vdash^p (P)_M : \mathbf{1}}{(\Gamma) \vdash^\perp \lambda().(P)_M : \mathbf{1} \multimap^{\text{pr}(\Gamma), p} \mathbf{1}}}{(\Gamma) \vdash^\perp \text{spawn}(\lambda().(P)_M) : \mathbf{1}} \quad \frac{(\Delta) \vdash^q (Q)_M : \mathbf{1}}{(\Gamma), (\Delta) \vdash^q \text{spawn}(\lambda().(P)_M); (Q)_M : \mathbf{1}} \\
\\
\text{T-HALT} \\
\frac{\overline{\mathbf{0} \vdash \emptyset}}{(\cdot)_M} \rightarrow \overline{\emptyset \vdash^\perp () : \mathbf{1}}
\end{array}$$

Figure 10: Translation $(\cdot)_M$ preserves typing (T-LINK, T-RES, T-PAR, and T-HALT).

Case T-SELECT-INL, T-SELECT-INR, and T-OFFER. *See Fig. 14.*

□

Theorem 4.6. *If $P \vdash \Gamma$, then $(\Gamma) \vdash^\circ (P)_c$.*

Proof. By induction on the derivation of $P \vdash \Gamma$.

Case T-RES. *Immediately, from the induction hypothesis.*

$$\text{T-RES} \\
\frac{\Gamma, x : A, y : A^\perp \vdash P}{\Gamma \vdash (\nu xy)P} \quad (\cdot)_c \rightarrow \frac{(\Gamma), x : (A), y : (B) \vdash^\circ (P)_c}{(\Gamma) \vdash^\circ (\nu xy)(P)_c}$$

$$\begin{array}{c}
\text{T-CLOSE} \\
\frac{P \vdash \Gamma \quad o < \text{pr}(\Gamma)}{x[] . P \vdash \Gamma, x : \mathbf{1}^o} \xrightarrow{(\cdot)_M} \\
\\
\frac{\overline{\text{close} : \text{end}_!^o \multimap^{\top, o} \mathbf{1}} \quad \overline{x : \text{end}_!^o \vdash^\perp x : \text{end}_!^o}}{x : \text{end}_!^o \vdash^o \text{close } x : \mathbf{1} \quad (\Gamma) \vdash^p (P)_M : \mathbf{1} \quad o < \text{pr}((\Gamma))} \\
\hline
(\Gamma), x : \text{end}_!^o \vdash^{o \sqcup p} \text{close } x; (P)_M : \mathbf{1} \\
\\
\text{T-WAIT} \\
\frac{P \vdash \Gamma \quad o < \text{pr}(\Gamma)}{x().P \vdash \Gamma, x : \mathbf{1}^o} \xrightarrow{(\cdot)_M} \\
\\
\frac{\overline{\text{wait} : \text{end}_?^o \multimap^{\top, o} \mathbf{1}} \quad \overline{x : \text{end}_?^o \vdash^\perp x : \text{end}_?^o}}{x : \text{end}_?^o \vdash^o \text{wait } x : \mathbf{1} \quad (\Gamma) \vdash^p (P)_M : \mathbf{1} \quad o < \text{pr}((\Gamma))} \\
\hline
(\Gamma), x : \text{end}_?^o \vdash^{o \sqcup p} \text{wait } x; (P)_M : \mathbf{1}
\end{array}$$

Figure 11: Translation $(\cdot)_M$ preserves typing (T-CLOSE and T-WAIT).

$$\begin{array}{c}
\text{T-SEND} \\
\frac{P \vdash \Gamma, y : A, x : B \quad o < \text{pr}(\Gamma, A, B)}{x[y].P \vdash \Gamma, x : A \otimes^o B} \xrightarrow{(\cdot)_M} \\
\\
\text{(A)} \\
\frac{\overline{\text{new} : \mathbf{1} \multimap (\overline{A}) \times (\overline{A})} \quad \overline{\emptyset \vdash^\perp () : \mathbf{1}}}{\emptyset \vdash^\perp \text{new } () : (\overline{A}) \times (\overline{A})} \\
\\
\text{(B)} \\
\frac{\overline{\text{send} : (\overline{A}) \times !^o(\overline{A}).(\overline{B}) \multimap^{\top, o} (\overline{B})}}{\frac{z : (\overline{A}) \vdash^\perp x : (\overline{A}) \quad x : !^o(\overline{A}).(\overline{B}) \vdash^\perp x : !^o(\overline{A}).(\overline{B})}{x : !^o(\overline{A}).(\overline{B}), z : (\overline{A}) \vdash^\perp (z, x) : (\overline{A}) \times !^o(\overline{A}).(\overline{B})}}{x : !^o(\overline{A}).(\overline{B}), z : (\overline{A}) \vdash^o \text{send } (z, x) : (\overline{B})} \\
\\
\text{(B)} \quad (\Gamma), y : (\overline{A}), x : (\overline{B}) \vdash^p (P)_M : \mathbf{1} \quad o < \text{pr}((\Gamma), (\overline{A}), (\overline{B})) \\
\text{(A)} \quad \frac{(\Gamma), x : !^o(\overline{A}).(\overline{B}), y : (\overline{A}), z : (\overline{A}) \vdash^{o \sqcup p} \text{let } x = \text{send } (z, x) \text{ in } (P)_M : \mathbf{1}}{(\Gamma), x : !^o(\overline{A}).(\overline{B}) \vdash^{o \sqcup p} \text{let } (y, z) = \text{new } () \text{ in let } x = \text{send } (z, x) \text{ in } (P)_M : \mathbf{1}}
\end{array}$$

Figure 12: Translation $(\cdot)_M$ preserves typing (T-SEND).

$$\begin{array}{c}
\text{T-RECV} \\
\frac{P \vdash \Gamma, y : A, x : B \quad o < \text{pr}(\Gamma, A, B)}{x(y).P \vdash \Gamma, x : A \wp^o B} \xrightarrow{(\cdot)_M} \\
\\
\text{(A)} \\
\frac{\frac{\text{recv} : ?^o(A). \langle B \rangle \multimap^{\top, o} \langle A \rangle \times \langle B \rangle \quad x : ?^o(A). \langle B \rangle \vdash^{\perp} x : ?^o(A). \langle B \rangle}{x : ?^o(A). \langle B \rangle \vdash^o \text{recv } x : \langle A \rangle \times \langle B \rangle}}{\\} \\
\\
\frac{\text{(A)} \quad \frac{\langle \Gamma \rangle, y : \langle A \rangle, x : \langle B \rangle \vdash^p \langle P \rangle_M : \mathbf{1} \quad o < \text{pr}(\langle \Gamma \rangle, \langle A \rangle, \langle B \rangle)}{\langle \Gamma \rangle, x : ?^o(A). \langle B \rangle, y : \langle A \rangle, z : \langle A \rangle \vdash^{o \sqcup p} \text{let } x = \text{recv } x \text{ in } \langle P \rangle_M : \mathbf{1}}}{\\}
\end{array}$$

Figure 13: Translation $(\cdot)_M$ preserves typing (T-RECV).

$$\begin{array}{c}
\text{T-SELECT-INL} \\
\frac{P \vdash \Gamma, x : A \quad o < \text{pr}(\Gamma)}{x \triangleleft \text{inl}.P \vdash \Gamma, x : A \oplus^o B} \xrightarrow{(\cdot)_M} \\
\\
\frac{\frac{\text{select inl} : \langle A \rangle \oplus^o \langle B \rangle \multimap^{\top, o} \langle A \rangle \quad x : \langle A \rangle \oplus^o \langle B \rangle \vdash^{\perp} x : \langle A \rangle \oplus^o \langle B \rangle}{x : \langle A \rangle \oplus^o \langle B \rangle \vdash^o \text{select inl } x : \langle A \rangle}}{\frac{\Gamma, x : \langle A \rangle \vdash^p \langle P \rangle_M : \mathbf{1} \quad o < \text{pr}(\Gamma)}{\Gamma, x : \langle A \rangle \oplus^o \langle B \rangle \vdash^{o \sqcup p} \text{let } x = \text{select inl } x \text{ in } \langle P \rangle_M : \mathbf{1}}} \\
\\
\text{T-SELECT-INR} \\
\frac{P \vdash \Gamma, x : A \quad o < \text{pr}(\Gamma)}{x \triangleleft \text{inr}.P \vdash \Gamma, x : A \oplus^o B} \xrightarrow{(\cdot)_M} \\
\\
\frac{\frac{\text{select inr} : \langle A \rangle \oplus^o \langle B \rangle \multimap^{\top, o} \langle B \rangle \quad x : \langle A \rangle \oplus^o \langle B \rangle \vdash^{\perp} x : \langle A \rangle \oplus^o \langle B \rangle}{x : \langle A \rangle \oplus^o \langle B \rangle \vdash^o \text{select inr } x : \langle B \rangle}}{\frac{\Gamma, x : \langle B \rangle \vdash^p \langle P \rangle_M : \mathbf{1} \quad o < \text{pr}(\Gamma)}{\Gamma, x : \langle A \rangle \oplus^o \langle B \rangle \vdash^{o \sqcup p} \text{let } x = \text{select inr } x \text{ in } \langle P \rangle_M : \mathbf{1}}} \\
\\
\text{T-OFFER} \\
\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B \quad o < \text{pr}(\Gamma, A, B)}{x \triangleright \{\text{inl} : P; \text{inr} : Q\} \vdash \Gamma, x : A \&^o B} \xrightarrow{(\cdot)_M} \\
\\
\frac{\frac{x : \langle A \rangle \&^o \langle B \rangle \vdash^{\perp} x : \langle A \rangle \&^o \langle B \rangle}{\langle \Gamma \rangle, x : \langle A \rangle \vdash^p \langle P \rangle_M : \mathbf{1} \quad \langle \Gamma \rangle, x : \langle B \rangle \vdash^p \langle Q \rangle_M : \mathbf{1} \quad o < \text{pr}(\langle \Gamma \rangle, \langle A \rangle, \langle B \rangle)}}{\langle \Gamma \rangle, x : \langle A \rangle \&^o \langle B \rangle \vdash^{o \sqcup p} \text{offer } x \{ \text{inl } x \mapsto \langle P \rangle_M; \text{inr } x \mapsto \langle Q \rangle_M \} : \mathbf{1}}
\end{array}$$

Figure 14: Translation $(\cdot)_M$ preserves typing (T-SELECT-INL, T-SELECT-INR, and T-OFFER).

Case T-PAR. *Immediately, from the induction hypotheses.*

$$\frac{\text{T-PAR} \quad \frac{\Gamma \vdash P \quad \Delta \vdash Q}{\Gamma, \Delta \vdash P \parallel Q}}{\frac{(\Gamma) \vdash^\circ \langle P \rangle_c \quad (\Delta) \vdash^\circ \langle Q \rangle_c}{(\Gamma), (\Delta) \vdash^\circ \langle P \rangle_c \parallel \langle Q \rangle_c}} \Downarrow_c$$

Case *. *By Lemma 4.5*

$$\Gamma \vdash P \xrightarrow{\Downarrow_c} \frac{(\Gamma) \vdash^P \langle P \rangle_M : \mathbf{1}}{(\Gamma) \vdash^\circ \langle P \rangle_M} \quad \square$$

Theorem 4.7. *If $P \vdash \Gamma$ and $\langle P \rangle_c \rightarrow_c \mathcal{C}$, there exists a Q such that $P \Rightarrow^+ Q$ and $\mathcal{C} \rightarrow_c^* \langle Q \rangle_c$*

Proof. By induction on the derivation of $\langle P \rangle_c \rightarrow_c \mathcal{C}$. We omit the cases which cannot occur as their left-hand side term forms are not in the image of the translation function, i.e., E-NEW, E-SPAWN, and E-LIFTM.

Case E-LINK.

$$(\nu x x')(\mathcal{F}[\mathbf{link}(w, x)] \parallel \mathcal{C}) \rightarrow_c \mathcal{F}[] \parallel \mathcal{C}\{w/x'\}$$

The source for $\mathbf{link}(w, x)$ must be $w \leftrightarrow x$. None of the translation rules introduce an evaluation context around the recursive call, hence $\mathcal{F}$ must be the empty context. Let P be the source term for $\mathcal{C}$, i.e., $\langle P \rangle_c = \mathcal{C}$. Hence, we have:

$$\begin{array}{ccc} (\nu x x')(w \leftrightarrow x \parallel P) & \xRightarrow{\quad} & P\{w/x'\} \\ \downarrow \Downarrow_c & & \downarrow \Downarrow_c \\ (\nu x x')(\circ \mathbf{link}(w, x) \parallel \langle P \rangle_c) & & \\ \downarrow \rightarrow_c^+ & & \downarrow \Downarrow_c \\ \langle P \rangle_c\{w/x'\} & \xRightarrow{\quad} & \langle P\{w/x'\} \rangle_c \end{array}$$

Case E-SEND.

$$(\nu x x')(\mathcal{F}[\mathbf{send}(V, x)] \parallel \mathcal{F}'[\mathbf{recv}(x')]) \rightarrow_c (\nu x x')(\mathcal{F}[x] \parallel \mathcal{F}'[(V, x')])$$

There are three possible sources for $\mathbf{send}$ and $\mathbf{recv}$: $x[y].P$ and $x'(y').Q$; $x \triangleleft \mathbf{inl}.P$ and $x' \triangleright \{\mathbf{inl} : Q; \mathbf{inr} : R\}$; or $x \triangleleft \mathbf{inr}.P$ and $x' \triangleright \{\mathbf{inl} : Q; \mathbf{inr} : R\}$.

Subcase $x[y].P$ and $x'(y').Q$. None of the translation rules introduce an evaluation context around the recursive call, hence $\mathcal{F}$ must be $\circ \mathbf{let } x = \square \mathbf{in } \langle P \rangle_M$. Similarly, $\mathcal{F}'$ must be $\circ \mathbf{let } (y', x') = \square \mathbf{in } \langle Q \rangle_M$. The value V must be an endpoint y , bound by the name restriction $(\nu y y')$ introduced by the translation. Hence, we have:

$$\begin{array}{ccc} (\nu x x')(x[y].P \parallel x'(y').Q) & \xRightarrow{\quad} & (\nu x x')(\nu y y')(P \parallel Q) \\ \downarrow \Downarrow_c & & \downarrow \Downarrow_c \\ (\nu x x')(\nu y y') \left(\begin{array}{l} \circ \mathbf{let } x = \mathbf{send}(y, x) \mathbf{in } \langle P \rangle_M \parallel \\ \circ \mathbf{let } (y', x') = \mathbf{recv } x' \mathbf{in } \langle Q \rangle_M \end{array} \right) & & \\ \downarrow \equiv \rightarrow_c^+ & & \downarrow \Downarrow_c \\ (\nu x x')(\nu y y')(\circ \langle P \rangle_M \parallel \circ \langle Q \rangle_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_c^*} & (\nu x x')(\nu y y')(\langle P \rangle_c \parallel \langle Q \rangle_c) \end{array}$$

Subcase $x \triangleleft \text{inl}.P$ and $x' \triangleright \{\text{inl} : Q; \text{inr} : R\}$. None of the translation rules introduce an evaluation context around the recursive call, hence $\mathcal{F}$ must be

$$\circ \text{let } x = \text{close } \square; y \text{ in } \langle P \rangle_M.$$

Similarly, $\mathcal{F}'$ must be

$$\circ \text{let } (y', x') = \square \text{ in wait } x'; \text{case } y' \{ \text{inl } y' \mapsto \langle Q \rangle_M; \text{inr } y' \mapsto \langle R \rangle_M \}.$$

Hence, we have:

$$\begin{array}{ccc} (\nu x x')(x \triangleleft \text{inl}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\}) & \xRightarrow{\quad} & (\nu x x')(P \parallel Q) \\ \downarrow \langle \cdot \rangle_M & & \downarrow \langle \cdot \rangle_C \\ (\nu x x') \left(\begin{array}{l} \circ \text{let } x = \text{select inl } x \text{ in } \langle P \rangle_M \parallel \\ \circ \text{offer } x' \{ \text{inl } x' \mapsto \langle Q \rangle_M; \text{inr } x' \mapsto \langle R \rangle_M \} \end{array} \right) & & \\ \downarrow \xrightarrow{+}_C & \xrightarrow[\text{(by Lemma 4.8)}]{\star}_C & \downarrow \langle \cdot \rangle_C \\ (\nu x x')(\circ \langle P \rangle_M \parallel \circ \langle Q \rangle_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\star}_C & (\nu x x')(\langle P \rangle_C \parallel \langle Q \rangle_C) \end{array}$$

Subcase $x \triangleleft \text{inr}.P$ and $x' \triangleright \{\text{inl} : Q; \text{inr} : R\}$. None of the translation rules introduce an evaluation context around the recursive call, hence $\mathcal{F}$ must be

$$\circ \text{let } x = \text{close } \square; y \text{ in } \langle P \rangle_M.$$

Similarly, $\mathcal{F}'$ must be

$$\circ \text{let } (y', x') = \square \text{ in wait } x'; \text{case } y' \{ \text{inl } y' \mapsto \langle Q \rangle_M; \text{inr } y' \mapsto \langle R \rangle_M \}.$$

Hence, we have:

$$\begin{array}{ccc} (\nu x x')(x \triangleleft \text{inr}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\}) & \xRightarrow{\quad} & (\nu x x')(P \parallel Q) \\ \downarrow \langle \cdot \rangle_M & & \downarrow \langle \cdot \rangle_C \\ (\nu x x') \left(\begin{array}{l} \circ \text{let } x = \text{select inr } x \text{ in } \langle P \rangle_M \parallel \\ \circ \text{offer } x' \{ \text{inl } x' \mapsto \langle Q \rangle_M; \text{inr } x' \mapsto \langle R \rangle_M \} \end{array} \right) & & \\ \downarrow \xrightarrow{+}_C & \xrightarrow[\text{(by Lemma 4.8)}]{\star}_C & \downarrow \langle \cdot \rangle_C \\ (\nu x x')(\circ \langle P \rangle_M \parallel \circ \langle R \rangle_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\star}_C & (\nu x x')(\langle P \rangle_C \parallel \langle R \rangle_C) \end{array}$$

Case E-CLOSE.

$$(\nu x x')(\mathcal{F}[\text{wait } x] \parallel \mathcal{F}'[\text{close } x']) \xrightarrow{C} \mathcal{F}[\langle \cdot \rangle] \parallel \mathcal{F}'[\langle \cdot \rangle]$$

The source for **wait** and **close** must be $x().P$ and $x'().Q$.

(The translation for $x \triangleright \{\text{inl} : P; \text{inr} : Q\}$ also introduces a **wait**, but it is blocked on another communication, and hence cannot be the first communication on a translated term. The translations for $x \triangleleft \text{inl}.P$ and $x \triangleleft \text{inr}.P$ also introduce a **close**, but these are similarly blocked.)

None of the translation rules introduce an evaluation context around the recursive call, hence $\mathcal{F}$ must be $\square; \langle P \rangle_M$. Similarly, $\mathcal{F}'$ must be $\square; \langle Q \rangle_M$. Hence, we have:

$$\begin{array}{ccc}
 (\nu x x')(x \square . P \parallel x'().Q) & \xRightarrow{\quad} & P \parallel Q \\
 \downarrow \langle \cdot \rangle_M & & \downarrow \langle \cdot \rangle_c \\
 (\nu x x')(\circ \text{close } x; \langle P \rangle_M \parallel \circ \text{wait } x'; \langle Q \rangle_M) & & \\
 \downarrow \rightarrow_C^+ & & \downarrow \langle \cdot \rangle_c \\
 \circ \langle P \rangle_M \parallel \circ \langle Q \rangle_M & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_C^*} & \langle P \rangle_c \parallel \langle Q \rangle_c
 \end{array}$$

Case E-LIFTC. By the induction hypothesis and E-LIFTC.

Case E-LIFTSC. By the induction hypothesis, E-LIFTSC, and Lemma 4.9. □

Lemma 4.8. For any P , either:

- $\circ \langle P \rangle_M = \langle P \rangle_c$; or
- $\circ \langle P \rangle_M \rightarrow_C^+ \langle P \rangle_c$, and for any $\mathcal{C}$, if $\circ \langle P \rangle_M \rightarrow_C \mathcal{C}$, then $\mathcal{C} \rightarrow_C^* \langle P \rangle_c$.

Proof. By induction on the structure of P .

Case $(\nu xy)P$. We have:

$$\begin{aligned}
 \langle (\nu xy)P \rangle_M &= \circ \text{let } (x, y) = \text{new } () \text{ in } \langle P \rangle_M \\
 &\rightarrow_C^+ (\nu xy)(\circ \langle P \rangle_M) \\
 &\rightarrow_C^* (\nu xy)\langle P \rangle_c \\
 &= \langle (\nu xy)P \rangle_c
 \end{aligned}$$

Case $P \parallel Q$.

$$\begin{aligned}
 \langle P \parallel Q \rangle_M &= \circ \text{spawn } (\lambda(). \langle P \rangle_M); \langle Q \rangle_M \\
 &\rightarrow_C^+ \circ \langle P \rangle_M \parallel \circ \langle Q \rangle_M \\
 &\rightarrow_C^* \langle P \rangle_c \parallel \langle Q \rangle_c \\
 &= \langle P \parallel Q \rangle_c
 \end{aligned}$$

Case $x[y].P$.

$$\begin{aligned}
 \langle x[y].P \rangle_M &= \text{let } (y, z) = \text{new } () \text{ in let } x = \text{send } (z, x) \text{ in } \langle P \rangle_M \\
 &\rightarrow_C^+ (\nu yz)(\circ \text{let } x = \text{send } (z, x) \text{ in } \langle P \rangle_M) \\
 &= \langle x[y].P \rangle_c
 \end{aligned}$$

Case $x \triangleleft \text{inl}.P$.

$$\begin{aligned}
 \langle x[y].P \rangle_M &= \text{let } x = \text{select inl } x \text{ in } \langle P \rangle_M \\
 &\triangleq \text{let } x = \text{let } (y, z) = \text{new } () \text{ in close } (\text{send } (\text{inl } y, x)); z \text{ in } \langle P \rangle_M \\
 &\rightarrow_C^+ (\nu yz)(\circ \text{let } x = \text{close } (\text{send } (\text{inl } y, x)); z \text{ in } \langle P \rangle_M) \\
 &= \langle x[y].P \rangle_c
 \end{aligned}$$

Case $x \triangleleft \text{inr}.P$.

$$\begin{aligned}
\llbracket x[y].P \rrbracket_M &= \text{let } x = \text{select inr } x \text{ in } \llbracket P \rrbracket_M \\
&\triangleq \text{let } x = \text{let } (y, z) = \text{new } () \text{ in close (send (inr } y, x)); z \text{ in } \llbracket P \rrbracket_M \\
&\rightarrow_C^+ (\nu yz)(\circ \text{let } x = \text{close (send (inr } y, x)); z \text{ in } \llbracket P \rrbracket_M) \\
&= \llbracket x[y].P \rrbracket_c
\end{aligned}$$

Case $*$. In all other cases, $\circ \llbracket P \rrbracket_M = \llbracket P \rrbracket_c$. □

Lemma 4.9. If $P \vdash \Gamma$ and $P \equiv Q$, then $\llbracket P \rrbracket_c \equiv \llbracket Q \rrbracket_c$.

Proof. Every axiom of the structural congruence in PCP maps directly to the axiom of the same name in PGV. □

Theorem 4.10.

If $P \vdash \Gamma$ and $P \Rightarrow Q$, then $\llbracket P \rrbracket_c \rightarrow_C^+ \llbracket Q \rrbracket_c$.

Proof. By induction on the derivation of $P \Rightarrow Q$.

Case E-LINK.

$$\begin{array}{ccc}
(\nu xx')(w \leftrightarrow x \parallel P) & \xRightarrow{\quad} & P\{w/x'\} \\
\downarrow \llbracket \cdot \rrbracket_c & & \downarrow \llbracket \cdot \rrbracket_c \\
(\nu xx')(\circ \text{link } (w, x) \parallel \llbracket P \rrbracket_c) & & \\
\downarrow \rightarrow_C^+ & & \downarrow \\
\llbracket P \rrbracket_c\{w/x'\} & \xrightarrow{\quad} & \llbracket P\{w/x'\} \rrbracket_c
\end{array}$$

Case E-SEND.

$$\begin{array}{ccc}
(\nu xx')(x[y].P \parallel x'(y').Q) & \xRightarrow{\quad} & (\nu xx')(\nu yy')(P \parallel Q) \\
\downarrow \llbracket \cdot \rrbracket_M & & \downarrow \llbracket \cdot \rrbracket_c \\
(\nu xx') \left(\circ \left(\text{let } (y, y') = \text{new } () \text{ in} \right. \right. & & \\
\quad \left. \left. \circ \text{let } x = \text{send } (y, x) \text{ in } \llbracket P \rrbracket_M \right) \parallel \right. & & \\
\quad \left. \circ \text{let } (y', x') = \text{recv } x' \text{ in } \llbracket Q \rrbracket_M \right) & & \\
\downarrow \rightarrow_C^+ & & \downarrow \\
(\nu xx')(\nu yy')(\circ \llbracket P \rrbracket_M \parallel \circ \llbracket Q \rrbracket_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_C^*} & (\nu xx')(\nu yy')(\llbracket P \rrbracket_c \parallel \llbracket Q \rrbracket_c)
\end{array}$$

Case E-CLOSE.

$$\begin{array}{ccc}
(\nu xx')(x \square .P \parallel x'().Q) & \xRightarrow{\quad} & P \parallel Q \\
\downarrow \llbracket \cdot \rrbracket_M & & \downarrow \llbracket \cdot \rrbracket_c \\
(\nu xx')(\circ \text{close } x; \llbracket P \rrbracket_M \parallel \circ \text{wait } x'; \llbracket Q \rrbracket_M) & & \\
\downarrow \rightarrow_C^+ & & \downarrow \\
\circ \llbracket P \rrbracket_M \parallel \circ \llbracket Q \rrbracket_M & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_C^*} & \llbracket P \rrbracket_c \parallel \llbracket Q \rrbracket_c
\end{array}$$

Case E-SELECT-INL.

$$\begin{array}{ccc}
(\nu xx')(x \triangleleft \text{inl}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\}) & \xRightarrow{\quad} & (\nu xx')(P \parallel Q) \\
\downarrow \Downarrow_M & & \downarrow \Downarrow_c \\
(\nu xx') \left(\begin{array}{l} \circ \text{let } x = \text{select inl } x \text{ in } \langle P \rangle_M \parallel \\ \circ \text{offer } x' \{ \text{inl } x' \mapsto \langle Q \rangle_M; \text{inr } x' \mapsto \langle R \rangle_M \} \end{array} \right) & & \\
\downarrow \rightarrow_c^+ & & \downarrow \\
(\nu xx')(\circ \langle P \rangle_M \parallel \circ \langle Q \rangle_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_c^*} & (\nu xx')(\langle P \rangle_c \parallel \langle Q \rangle_c)
\end{array}$$

Case E-SELECT-INR.

$$\begin{array}{ccc}
(\nu xx')(x \triangleleft \text{inr}.P \parallel x \triangleright \{\text{inl} : Q; \text{inr} : R\}) & \xRightarrow{\quad} & (\nu xx')(P \parallel R) \\
\downarrow \Downarrow_M & & \downarrow \Downarrow_c \\
(\nu xx') \left(\begin{array}{l} \circ \text{let } x = \text{select inr } x \text{ in } \langle P \rangle_M \parallel \\ \circ \text{offer } x' \{ \text{inl } x' \mapsto \langle Q \rangle_M; \text{inr } x' \mapsto \langle R \rangle_M \} \end{array} \right) & & \\
\downarrow \rightarrow_c^+ & & \downarrow \\
(\nu xx')(\circ \langle P \rangle_M \parallel \circ \langle R \rangle_M) & \xrightarrow[\text{(by Lemma 4.8)}]{\rightarrow_c^*} & (\nu xx')(\langle P \rangle_c \parallel \langle R \rangle_c)
\end{array}$$

Case E-LIFTRES. *By the induction hypothesis and E-LIFTC.*

Case E-LIFTPAR. *By the induction hypotheses and E-LIFTC.*

Case E-LIFTSC. *By the induction hypothesis, E-LIFTSC, and Lemma 4.9.* □

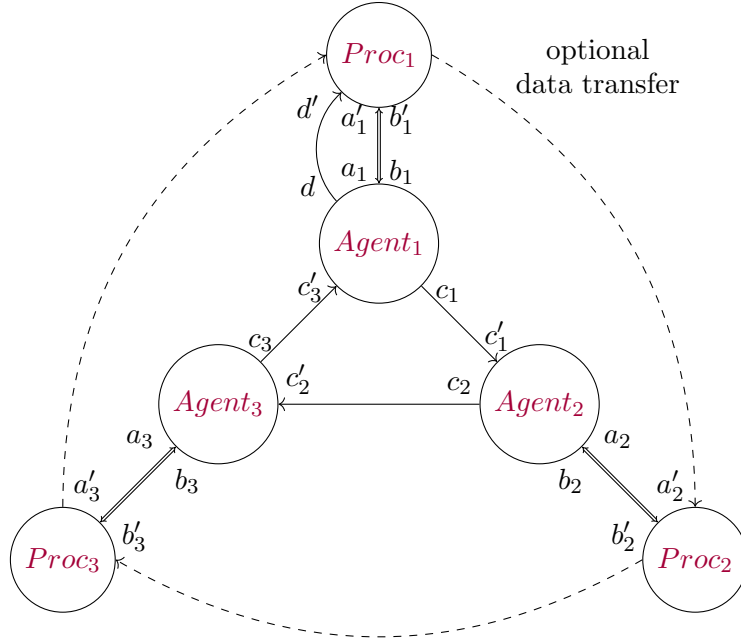
5. MILNER'S CYCLIC SCHEDULER

As an example of a deadlock-free cyclic process, Dardha and Gay [DG18b] introduce an implementation of Milner's cyclic scheduler [Mil89] in Priority CP. We reproduce that scheduler here, and show its translation to Priority GV.

Example 5.1. A set of processes Proc_i , $1 \leq i \leq n$, is scheduled to perform some tasks in cyclic order, starting with Proc_1 , ending with Proc_n , and notifying Proc_1 when all processes have finished.

Our scheduler Sched consists of set of agents Agent_i , $1 \leq i \leq n$, each representing their respective process. Each process Proc_i waits for the signal to start their task on a'_i , and signals completion on b'_i . Each agent signals their process to start on a_i , waits for their process to finish on b_i , and then signals for the next agent to continue on c_i . The agent Agent_1 initiates, then waits for every other process to finish, and signals Proc_1 on d . Every other agent Agent_i , $2 \leq i \leq n$ waits on c'_{i-1} for the signal to start. Each of the channels in the scheduler is of a terminated type, and is merely used to synchronise.

Below is a diagram of our scheduler instantiated with three processes:



We implement the scheduler as follows, using $\prod_I P_i$ to denote the parallel composition of the processes P_i , $i \in I$, and $P[Q]$ to denote the plugging of Q in the one-hole process-context P . The process-contexts P_i represent the computations performed by each process $Proc_i$. The process-contexts Q_i represent any post-processing, and any possible data transfer from $Proc_i$ to $Proc_{i+1}$. Finally, Q_1 should contain d' .

$$\begin{aligned}
 Sched &\triangleq (\nu a_1 a'_1) \dots (\nu a_n a'_n) (\nu b_1 b'_1) \dots (\nu b_n b'_n) (\nu c_1 c'_1) \dots (\nu c_n c'_n) (\nu d d') \\
 &\quad (Proc_1 \parallel Agent_1 \parallel \prod_{2 \leq i \leq n} (Proc_i \parallel c'_{i-1}().Agent_i)) \\
 Agent_1 &\triangleq a_i().b_i().c_i().c'_n().d().\mathbf{0} \\
 Agent_i &\triangleq a_i().b_i().c_i().\mathbf{0} \\
 Proc_i &\triangleq a'_i().P_i[b'_i].Q_i
 \end{aligned}$$

Example 5.2. The PGV scheduler has exactly the same behaviour as the PCP version in Example 5.1. It is implemented as follows, using $\prod_I C_i$ to denote the parallel composition of the processes C_i , $i \in I$, and $M[N]$ to denote the plugging of N in the one-hole term-context M . For simplicity, we let **sched** be a configuration. The terms M_i represent the computations performed by each process **proc** _{i} . The terms N_i represent any post-processing, and any possible data transfer from **proc** _{i} to **proc** _{$i+1$} . Finally, N_1 should contain **wait** d' .

$$\begin{aligned}
 \mathbf{sched} &\triangleq (\nu a_1 a'_1) \dots (\nu a_n a'_n) (\nu b_1 b'_1) \dots (\nu b_n b'_n) (\nu c_1 c'_1) \dots (\nu c_n c'_n) (\nu d d') \\
 &\quad (\phi \mathbf{proc}_1 \parallel \circ \mathbf{agent}_1; \mathbf{wait} \ c'_n; \mathbf{close} \ d \\
 &\quad \parallel \prod_{2 \leq i \leq n} (\circ \mathbf{proc}_i \parallel \circ \mathbf{wait} \ c'_{i-1}; \mathbf{agent}_i)) \\
 \mathbf{agent}_i &\triangleq \mathbf{close} \ a_i; \mathbf{wait} \ b_i; \mathbf{close} \ c_i \\
 \mathbf{proc}_i &\triangleq \mathbf{wait} \ a'_i; M_i[\mathbf{close} \ b'_i; N_i]
 \end{aligned}$$

If $\langle P_i \rangle_M = M_i$ and $\langle Q_i \rangle_M = N_i$, then the translation of $Sched$ (Example 5.1), $\langle Sched \rangle_c$, is exactly **sched** (Example 5.2).

6. RELATED WORK AND CONCLUSION

Deadlock freedom and progress. Deadlock freedom and progress are well studied properties in the π -calculus. For the “standard” typed π -calculus—types for channels used in input and output, an important line of work stems from Kobayashi’s approach to deadlock freedom [Kob98], where priorities are values from an abstract poset. Kobayashi [Kob06] simplifies the abstract poset to pairs of naturals, called *obligations* and *capabilities*. Padovani simplifies these further to a single natural, called a *priority* [Pad14], and adapts obligations/capabilities to session types [Pad13]. Later work by Kobayashi and co-authors [GKL14, KL17] address deadlock detection for a value-passing CCS (a predecessor of the π -calculus) where the number of nodes in a network is arbitrary, namely modelling unbounded networks. The authors define a sound inference algorithm for their type system, which guarantees deadlock freedom for these more complex kinds of communication networks. This type system is more expressive than previous work by Kobayashi.

For the session-typed π -calculus, an important line of work stems from Dezani and co-authors. In their work, Dezani *et al.* [DCMYD06] guarantee progress by allowing only one active session at a time. As sessions do not interleave, consequently they do not deadlock. Later on, Dezani *et al.* [DCdY07] introduce a partial order on channels, similar to Kobayashi [Kob98] and produce a type system for progress.

Carbone and Debois [CD10] define progress for session-typed π -calculus in terms of a *catalyser*, which provides the missing counterpart to a process. Intuitively, either the process is deadlock-free in which case the catalyser is structurally congruent to the inaction process, or if the process is “stuck” for *e.g.*, an input process $x(y).\mathbf{0}$ (using CP syntax) then the catalyser simply provides the missing output required for communication. Carbone *et al.* [CDM14] use catalysers to show that progress is a compositional form of livelock freedom and can be lifted to session types via the encoding of session types to linear types [Kob07, DGS12, Dar14, DGS17]. In this work, Carbone *et al.* [CDM14] sistematise and compare different notions of liveness properties: progress, deadlock freedom and livelock freedom. Their technique of using the encoding of session types combined with Kobayashi’s type systems for deadlock and livelock freedom allows for a more flexible deadlock/livelock detection. It is worth noting that livelock freedom is a stronger property than deadlock freedom in the presence of recursion, as the former also discards useless divergent processes; for finite processes the two properties coincide.

Vieira and Vasconcelos [VV13] use single priorities and an abstract partial order to guarantee deadlock freedom in a binary session-typed π -calculus and building on conversation types, which is akin to session types.

While our work focuses on binary session types, it is worth discussing related work on Multiparty Session Types (MPST), which describe communication among multiple agents in a distributed setting. The line of work on MPST starts with Honda *et al.* [HYC08], which guarantees deadlock freedom within a single session by design, but the property does not hold for session interleaving. Bettini *et al.* [BCD⁺08] follow a technique similar to Kobayashi’s for MPST. The main difference with our work is that we associate priorities with communication actions, where Bettini *et al.* [BCD⁺08] associate them with channels. Coppo *et al.* [CDPY13, CDYP16] present a deterministic, sound and complete, and compositional type inference algorithm for an *interaction type system*, which is used to guarantee global progress for processes in a calculus based on asynchronous as well as dynamically interleaved

and interfered multiparty sessions. The interaction type system allows to infer *causalities*—much in the line of Kobayashi’s priorities—of channels, thus guaranteeing that session-typed processes do not get stuck at intermediate stages of their sessions. Carbone and Montesi [CM13] combine MPST with choreographic programming and obtain a formalism that satisfies deadlock freedom. In the same vein as MPST, choreographic programming specifies communication among all participants in a distributed system. While MPST target mainly protocol descriptions, choreographies have mainly targetted implementations and programming languages as they are suitable for describing concrete system implementations. Deniérou and Yoshida [DY13] introduce *multiparty compatibility*, which generalises the notion of duality in binary session types. They synthesise safe and deadlock-free *global types*—specifying communication among all involved participants, from *local session types*—specifying communication from the viewpoint of one participant. To do so, they leverage Labelled Transition Systems (LTSs) and communicating automata. Castellani *et al.* [CDGH20] guarantee livelock freedom, a stronger property than deadlock freedom, for MPST with internal delegation, where participants in the same session are allowed to delegate tasks to each other, and internal delegation is captured by the global type. Scalas and Yoshida [SY19] provide a revision of the foundations for MPST, and offer a less complicated and more general theory, by removing duality/consistency. The type systems is parametric and type checking is decidable, but allows for a novel integration of model checking techniques. More protocols and processes can be typed and are guaranteed to be free of deadlocks.

Neubauer and Thiemann [NT04] and Vasconcelos *et al.* [VRG04, VGR06] introduce the first functional language with session types. Such works did not guarantee deadlock freedom until GV [LM15, Wad14]. Toninho *et al.* [TCP12] present a translation of simply-typed λ -calculus into session-typed π -calculus, but their focus is not on deadlock freedom. Fowler *et al.* [FKD⁺21] present Hypersequent GV (HGV), which is a variation of GV that uses hyper-environments, much in the same line as Hypersequent CP, and enjoys deadlock freedom, confluence, and strong normalisation.

Ties with logic. The correspondence between logic and types lays the foundation for functional programming [Wad15]. Since its inception by Girard [Gir87], linear logic has been a candidate for a foundational correspondence for concurrent programs. A correspondence with linear π -calculus was established early on by Abramsky [Abr94] and Bellin and Scott [BS94]. Many years later, several correspondences between linear logic and the π -calculus with binary session types were proposed. Caires and Pfenning [CP10] propose a correspondence with dual intuitionistic linear logic, while Wadler [Wad12] proposes a correspondence with classical linear logic. Both works guarantee deadlock freedom as a consequence of adopting cut elimination. Building on a previous work [CP10], Toninho *et al.* [TCP13] present a Curry-Howard correspondence between session types and linear logic for functional language via linear contextual monads, which are first-class values, thus giving rise to a higher-order session-typed language. In addition to the more standard results, the authors also prove a global progress theorem. Qian *et al.* [QKB21] extend Classical Linear Logic with co-exponentials, which allows to model servers receiving requests from an arbitrary set of clients, yielding an extension to the Curry-Howard correspondence between logic and session typed processes. Dardha and Gay [DG18b] define Priority CP by integrating Kobayashi and Padovani’s work on priorities [Kob06, Pad14] with CP, which as described in the introduction, weakens its ties to linear logic in exchange for expressivity. However, they show how PCP can be also viewed as an extension of linear logic, which they call Priority

Linear Logic (PLL), and uses mix and cycle rules as opposed to the cut rule. Dardha and Pérez [DP18, DP15, DP22] compare priorities à la Kobayashi with tree restrictions à la CP, and show that the latter is a subsystem of the former. In addition, they give a detailed account of comparing several type systems for deadlock freedom spanning across session types, linear logic, and linear types. Carbone *et al.* [CMSY15, CLM⁺16] give a logical view of MPST with a generalised duality. Caires and Pérez [CP16] give a presentation of MPST in terms of binary session types and the use of a *medium process* which guarantee protocol fidelity and deadlock freedom. Their binary session types are rooted in linear logic. Ciobanu and Horne [CH15] give the first Curry-Howard correspondence between MPST and BV [Gug07], a conservative extension of linear logic with a non-commutative operator for sequencing. Horne [Hor20] give a system for subtyping and multiparty compatibility where compatible processes are race free and deadlock free using a Curry-Howard correspondence, similar to the approach in [CH15]. Balzer *et al.* [BP17] introduce sharing at the cost of deadlock freedom, which they restore using *worlds*, an approach similar to priorities [BTP19]. Staying on sharing, Rocha and Caires [RC21] introduce an imperative feature, that of shared mutable states into a functional language based on Curry-Howard correspondence with linear logic. Their type system is thus able to capture programs which were not possible in previous works. The authors prove extensive technical results, including session fidelity, progress, confluence and normalisation. Lastly, Jacobs *et al.* [JBK22] present a novel technique to guarantee deadlock freedom based on the notion of connectivity graph, which is an abstract representation of the topology of concurrent systems, and separation logic used to substructurally treat connectivity graph edges and labels.

Conclusion and Future Work. We answered our research question by presenting Priority GV, a session-typed functional language which allows cyclic communication structures and uses priorities to ensure deadlock freedom. We showed its relation to Priority CP [DG18b] via an operational correspondence.

Our formalism so far only captures the core of GV. In future work, we plan to explore recursion, following Lindley and Morris [LM16] and Padovani and Novara [PN15], and sharing, following Balzer and Pfenning [BP17] or Voinea *et al.* [VDG19].

Acknowledgements. The authors would like to thank the anonymous reviewers for their detailed feedback, which helped produce a more complete and polished work. Also, the authors would like to thank Simon Fowler, April Gonçalves, and Philip Wadler for their comments on the manuscript.

REFERENCES

- [Abr94] Samson Abramsky. Proofs as processes. *Theor. Comput. Sci.*, 135(1):5–9, 1994. doi:10.1016/0304-3975(94)00103-0.
- [Bar96] Andrew Barber. Dual intuitionistic linear logic. 1996. URL: <https://www.lfcs.inf.ed.ac.uk/reports/96/ECS-LFCS-96-347/ECS-LFCS-96-347.pdf>.
- [BBN⁺18] Jean-Philippe Bernardy, Mathieu Boespflug, Ryan R. Newton, Simon Peyton Jones, and Arnaud Spiwack. Linear haskell: practical linearity in a higher-order polymorphic language. *Proc. of POPL*, 2:1–29, 2018. doi:10.1145/3158093.

- [BCD⁺08] Lorenzo Bettini, Mario Coppo, Loris D’Antoni, Marco De Luca, Mariangiola Dezani-Ciancaglini, and Nobuko Yoshida. Global progress in dynamically interleaved multiparty sessions. In *Proc. of CONCUR*, volume 5201 of *Lect. Notes Comput. Sci.*, pages 418–433. Springer, 2008. doi:10.1007/978-3-540-85361-9_33.
- [BP17] Stephanie Balzer and Frank Pfenning. Manifest sharing with session types. *Proc. ACM Program. Lang.*, 1(ICFP):37:1–37:29, 2017. doi:10.1145/3110281.
- [BS94] Gianluigi Bellin and Philip J. Scott. On the π -calculus and linear logic. *Theor. Comput. Sci.*, 135(1):11–65, 1994. doi:10.1016/0304-3975(94)00104-9.
- [BTP19] Stephanie Balzer, Bernardo Toninho, and Frank Pfenning. Manifest deadlock-freedom for shared session types. In *Proc. of ESOP*, volume 11423 of *Lect. Notes Comput. Sci.*, pages 611–639. Springer, 2019. doi:10.1007/978-3-030-17184-1_22.
- [CD10] Marco Carbone and Søren Debois. A graphical approach to progress for structured communication in web services. In *Proc. of ICE*, volume 38 of *Electron. Proc. in Theor. Comput. Sci.*, pages 13–27, 2010. doi:10.4204/EPTCS.38.4.
- [CDGH20] Iliaria Castellani, Mariangiola Dezani-Ciancaglini, Paola Giannini, and Ross Horne. Global types with internal delegation. *Theor. Comput. Sci.*, 807:128–153, 2020. doi:10.1016/j.tcs.2019.09.027.
- [CDM14] Marco Carbone, Ornella Dardha, and Fabrizio Montesi. Progress as compositional lock-freedom. In *Proc. of COORDINATION*, volume 8459 of *Lect. Notes Comput. Sci.*, pages 49–64. Springer, 2014. doi:10.1007/978-3-662-43376-8_4.
- [CDPY13] Mario Coppo, Mariangiola Dezani-Ciancaglini, Luca Padovani, and Nobuko Yoshida. Inference of global progress properties for dynamically interleaved multiparty sessions. In *Proc. of COORDINATION*, volume 7890 of *Lect. Notes Comput. Sci.*, pages 45–59. Springer, 2013. doi:10.1007/978-3-642-38493-6_4.
- [CDYP16] Mario Coppo, Mariangiola Dezani-Ciancaglini, Nobuko Yoshida, and Luca Padovani. Global progress for dynamically interleaved multiparty sessions. *Math. Struct. Comput. Sci.*, 26(2):238–302, 2016. doi:10.1017/S0960129514000188.
- [CH15] Gabriel Ciobanu and Ross Horne. Behavioural analysis of sessions using the calculus of structures. In *Proc. of PSI*, volume 9609 of *Lect. Notes Comput. Sci.*, pages 91–106. Springer, 2015. doi:10.1007/978-3-319-41579-6_8.
- [CLM⁺16] Marco Carbone, Sam Lindley, Fabrizio Montesi, Carsten Schürmann, and Philip Wadler. Coherence generalises duality: A logical explanation of multiparty session types. In *Proc. of CONCUR*, volume 59 of *LIPICs*, pages 33:1–33:15. Leibniz-Zentrum für Informatik, 2016. doi:10.4230/LIPICs.CONCUR.2016.33.
- [CM13] Marco Carbone and Fabrizio Montesi. Deadlock-freedom-by-design: Multiparty asynchronous global programming. In *Proc. of POPL*, pages 263–274, 2013. doi:10.1145/2480359.2429101.
- [CMS18] Marco Carbone, Fabrizio Montesi, and Carsten Schürmann. Choreographies, logically. *Distrib. Comput.*, 31(1):51–67, 2018. doi:10.1007/978-3-662-44584-6_5.
- [CMSY15] Marco Carbone, Fabrizio Montesi, Carsten Schürmann, and Nobuko Yoshida. Multiparty session types as coherence proofs. In *Proc. of CONCUR*, volume 42 of *LIPICs*, pages 412–426. Leibniz-Zentrum für Informatik, 2015. doi:10.1007/s00236-016-0285-y.
- [CP10] Luís Caires and Frank Pfenning. Session types as intuitionistic linear propositions. In *Proc. of CONCUR*, volume 6269 of *Lect. Notes Comput. Sci.*, pages 222–236. Springer, 2010. doi:10.1007/978-3-642-15375-4_16.
- [CP16] Luís Caires and Jorge A. Pérez. Multiparty session types within a canonical binary theory, and beyond. In *Proc. of FORTE*, volume 9688 of *Lect. Notes Comput. Sci.*, pages 74–95. Springer, 2016. doi:10.1007/978-3-319-39570-8_6.
- [Dar14] Ornella Dardha. Recursive session types revisited. In *Proc. of BEAT*, volume 162 of *Electron. Proc. in Theor. Comput. Sci.*, pages 27–34, 2014. doi:10.4204/EPTCS.162.4.
- [DCdY07] Mariangiola Dezani-Ciancaglini, Ugo de’Liguoro, and Nobuko Yoshida. On progress for structured communications. In *Proc. of TGC*, volume 4912 of *Lect. Notes Comput. Sci.*, pages 257–275. Springer, 2007. doi:10.1007/978-3-540-78663-4_18.
- [DCMYD06] Mariangiola Dezani-Ciancaglini, Dimitris Mostros, Nobuko Yoshida, and Sophia Drossopoulou. Session types for object-oriented languages. In *Proc. of ECOOP*, volume 4067 of *Lect. Notes Comput. Sci.*, pages 328–352. Springer, 2006. doi:10.1007/11785477_20.

- [DG18a] Ornella Dardha and Simon J. Gay. A new linear logic for deadlock-free session-typed processes. In *Proc. of FoSSaCS*, volume 10803 of *Lect. Notes Comput. Sci.*, pages 91–109. Springer, 2018. doi:10.1007/978-3-319-89366-2_5.
- [DG18b] Ornella Dardha and Simon J. Gay. A new linear logic for deadlock-free session typed processes. Extended version of [DG18a]. Available at <http://www.dcs.gla.ac.uk/~ornela/publications/DG18-Extended.pdf>, 2018.
- [DGS12] Ornella Dardha, Elena Giachino, and Davide Sangiorgi. Session types revisited. In *Proc. of PPDP*, pages 139–150. ACM, 2012. doi:10.1145/2370776.2370794.
- [DGS17] Ornella Dardha, Elena Giachino, and Davide Sangiorgi. Session types revisited. *Inf. Comput.*, 256:253–286, 2017. doi:10.1016/j.ic.2017.06.002.
- [DP15] Ornella Dardha and Jorge A. Pérez. Comparing deadlock-free session typed processes. In *Proc. of EXPRESS/SOS*, volume 190 of *Electron. Proc. in Theor. Comput. Sci.*, pages 1–15, 2015. doi:10.4204/EPTCS.190.1.
- [DP18] Ornella Dardha and Jorge A. Pérez. Comparing type systems for deadlock-freedom. *CoRR*, abs/1810.00635, 2018. URL: <http://arxiv.org/abs/1810.00635>, arXiv:1810.00635.
- [DP22] Ornella Dardha and Jorge A. Pérez. Comparing type systems for deadlock freedom. *J. Log. Algebraic Methods Program.*, 124:100717, 2022. doi:10.1016/j.jlamp.2021.100717.
- [DY13] Pierre-Malo Deniérou and Nobuko Yoshida. Multiparty compatibility in communicating automata: Characterisation and synthesis of global session types. In *Proc. of ICALP*, volume 7966 of *Lect. Notes Comput. Sci.*, pages 174–186. Springer, 2013. doi:10.1007/978-3-642-39212-2_18.
- [FKD⁺21] Simon Fowler, Wen Kokke, Ornella Dardha, Sam Lindley, and J. Garrett Morris. Separating sessions smoothly. In *Proc. of CONCUR*, volume 203 of *LIPICs*, pages 36:1–36:18. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2021. doi:10.4230/LIPICs.CONCUR.2021.36.
- [Gir87] Jean-Yves Girard. Linear logic. *Theor. Comput. Sci.*, 50:1–102, 1987. doi:10.1016/0304-3975(87)90045-4.
- [GKL14] Elena Giachino, Naoki Kobayashi, and Cosimo Laneve. Deadlock analysis of unbounded process networks. In *Proc. of CONCUR*, volume 8704 of *Lect. Notes Comput. Sci.*, pages 63–77. Springer, 2014. doi:10.1007/978-3-662-44584-6_6.
- [Gug07] Alessio Guglielmi. A system of interaction and structure. *ACM Trans. Comput. Log.*, 8(1):1, 2007. doi:10.1145/1182613.1182614.
- [GV10] Simon J. Gay and Vasco T. Vasconcelos. Linear type theory for asynchronous session types. *J. Funct. Program.*, 20(1):19–50, 2010. doi:10.1017/S0956796809990268.
- [Hon93] Kohei Honda. Types for dyadic interaction. In *Proc. of CONCUR*, volume 715 of *Lect. Notes Comput. Sci.*, pages 509–523. Springer, 1993. doi:10.1007/3-540-57208-2_35.
- [Hor20] Ross Horne. Session subtyping and multiparty compatibility using circular sequents. In *Proc. of CONCUR*, volume 171 of *LIPICs*, pages 12:1–12:22. Leibniz-Zentrum für Informatik, 2020. doi:10.4230/LIPICs.CONCUR.2020.12.
- [HVK98] Kohei Honda, Vasco Thudichum Vasconcelos, and Makoto Kubo. Language primitives and type discipline for structured communication-based programming. In *Proc. of ESOP*, volume 1381 of *Lect. Notes Comput. Sci.*, pages 122–138. Springer, 1998. doi:10.1007/BFb0053567.
- [HYC08] Kohei Honda, Nobuko Yoshida, and Marco Carbone. Multiparty asynchronous session types. In *Proc. of POPL*, volume 43(1), pages 273–284. ACM, 2008. doi:10.1145/2827695.
- [JBK22] Jules Jacobs, Stephanie Balzer, and Robbert Krebbers. Multiparty GV: functional multiparty session types with certified deadlock freedom. *Proc. ACM Program. Lang.*, 6(ICFP):466–495, 2022. doi:10.1145/3547638.
- [KD21a] Wen Kokke and Ornella Dardha. Deadlock-free session types in Linear Haskell. 2021. URL: <https://arxiv.org/abs/2103.14481>.
- [KD21b] Wen Kokke and Ornella Dardha. Prioritise the best variation. In *Proc. of FORTE*, volume 12719 of *Lect. Notes Comput. Sci.*, pages 100–119. Springer, 2021. doi:10.1007/978-3-030-78089-0_6.
- [KL17] Naoki Kobayashi and Cosimo Laneve. Deadlock analysis of unbounded process networks. *Inf. Comput.*, 252:48–70, 2017. doi:10.1016/j.ic.2016.03.004.

- [KMP19a] Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Better late than never: A fully-abstract semantics for classical processes. *Proc. ACM Program. Lang.*, 3(POPL), 2019. doi:10.1145/3290337.
- [KMP19b] Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Taking linear logic apart. In *Proc. of Linearity & TLLA*, volume 292 of *Electron. Proc. in Theor. Comput. Sci.*, pages 90–103. Open Publishing Association, 2019. doi:10.4204/EPTCS.292.5.
- [Kob98] Naoki Kobayashi. A partially deadlock-free typed process calculus. *ACM Trans. Program. Lang. Syst.*, 20(2):436–482, 1998. doi:10.1145/276393.278524.
- [Kob06] Naoki Kobayashi. A new type system for deadlock-free processes. In *Proc. of CONCUR*, volume 4137 of *Lect. Notes Comput. Sci.*, pages 233–247. Springer, 2006. doi:10.1007/11817949_16.
- [Kob07] Naoki Kobayashi. Type systems for concurrent programs. 2007.
- [LM15] Sam Lindley and J. Garrett Morris. A semantics for propositions as sessions. In *Proc. of ESOP*, pages 560–584, 2015. doi:10.1007/978-3-662-46669-8_23.
- [LM16] Sam Lindley and J. Garrett Morris. Talking bananas: Structural recursion for session types. In *Proc. of ICFP*. ACM, 2016. doi:10.1145/2951913.2951921.
- [Mil89] Robin Milner. *Communication and Concurrency*. Prentice Hall, 1989. doi:10.5555/63446.
- [MP18] Fabrizio Montesi and Marco Peressotti. Classical transitions. 2018. URL: <https://arxiv.org/abs/1803.01049>.
- [NT04] Matthias Neubauer and Peter Thiemann. An implementation of session types. In Bharat Jayaraman, editor, *Proc. of PADL*, volume 3057 of *Lect. Notes Comput. Sci.*, pages 56–70. Springer, 2004. doi:10.1007/978-3-540-24836-1_5.
- [Pad13] Luca Padovani. From lock freedom to progress using session types. In *Proc. of PLACES*, volume 137, pages 3–19. Electron. Proc. in Theor. Comput. Sci., 2013. doi:10.4204/EPTCS.137.2.
- [Pad14] Luca Padovani. Deadlock and lock freedom in the linear π -calculus. In *Proc. of CSL-LICS*, pages 72:1–72:10. ACM, 2014. doi:10.1145/2603088.2603116.
- [PG15] Frank Pfenning and Dennis Griffith. Polarized substructural session types. In *Proc. of FoS-SaCS*, volume 9034 of *Lect. Notes Comput. Sci.*, pages 3–22. Springer, 2015. doi:10.1007/978-3-662-46678-0_1.
- [PN15] Luca Padovani and Luca Novara. Types for deadlock-free higher-order programs. In *Proc. of FORTE*, volume 9039 of *Lect. Notes Comput. Sci.*, pages 3–18. Springer, 2015. doi:10.1007/978-3-319-19195-9_1.
- [QKB21] Zesen Qian, G. A. Kavvos, and Lars Birkedal. Client-server sessions in linear logic. *Proc. of ACM Program. Lang.*, 5(ICFP):1–31, 2021. doi:10.1145/3473567.
- [RC21] Pedro Rocha and Luís Caires. Propositions-as-types and shared state. *Proc. ACM Program. Lang.*, 5(ICFP):1–30, 2021. doi:10.1145/3473584.
- [SY19] Alceste Scalas and Nobuko Yoshida. Less is more: Multiparty session types revisited. *Proc. ACM Program. Lang.*, 3(POPL), 2019. doi:10.1145/3290343.
- [TCP12] Bernardo Toninho, Luís Caires, and Frank Pfenning. Functions as session-typed processes. In *Proc. of FoSSaCS*, volume 7213 of *Lect. Notes Comput. Sci.*, pages 346–360. Springer, 2012. doi:10.1007/978-3-642-28729-9_23.
- [TCP13] Bernardo Toninho, Luís Caires, and Frank Pfenning. Higher-order processes, functions, and sessions: A monadic integration. In *Proc. of ESOP*, volume 7792 of *Lect. Notes Comput. Sci.*, pages 350–369. Springer, 2013. doi:10.1007/978-3-642-37036-6_20.
- [THK94] Kaku Takeuchi, Kohei Honda, and Makoto Kubo. An interaction-based language and its typing system. In *Proc. of PARLE*, volume 817 of *Lect. Notes Comput. Sci.*, pages 398–413. Springer, 1994. doi:10.1007/3-540-58184-7_118.
- [VDG19] A. Laura Voinea, Ornela Dardha, and Simon J. Gay. Resource sharing via capability-based multiparty session types. In *Proc. of IFM*, volume 11918 of *Lect. Notes Comput. Sci.*, pages 437–455. Springer, 2019. doi:10.1007/978-3-030-34968-4_24.
- [VGR06] Vasco Thudichum Vasconcelos, Simon J. Gay, and António Ravara. Type checking a multi-threaded functional language with session types. *Theor. Comput. Sci.*, 368(1-2):64–87, 2006. doi:10.1016/j.tcs.2006.06.028.
- [VRG04] Vasco Vasconcelos, António Ravara, and Simon J. Gay. Session types for functional multithreading. In *Proc. of CONCUR*, volume 3170 of *Lect. Notes Comput. Sci.*, pages 497–511. Springer, 2004. doi:10.1007/978-3-540-28644-8_32.

- [VV13] Hugo Torres Vieira and Vasco Thudichum Vasconcelos. Typing progress in communication-centred systems. In *Proc. of COORDINATION*, volume 7890 of *Lect. Notes Comput. Sci.*, pages 236–250. Springer, 2013. doi:10.1007/978-3-662-43376-8_10.
- [Wad12] Philip Wadler. Propositions as sessions. In *Proc. of ICFP*, pages 273–286, 2012. doi:10.1145/2398856.2364568.
- [Wad14] Philip Wadler. Propositions as sessions. *J. Funct. Program.*, 24(2-3):384–418, 2014. doi:10.1017/S095679681400001X.
- [Wad15] Philip Wadler. Propositions as types. *Commun. ACM*, 58(12):75–84, 2015. doi:10.1145/2699407.

5.3 Discussion

This section proceeds as follows:

- In § 5.3.1, we discuss the relation between PCP and CP.
- In § 5.3.2, we define sound and complete priority inference for PCP.

In this section, PCP’s processes and PGV’s terms are printed in **red**, both of their types are printed in **blue**, and both of their priorities are printed in **yellow**, and all are rendered in a sans-serif font. To save on accessible colour combinations, the processes and terms, types, and priorities of *any other system* are printed in *pink*, *green*, and *periwinkle*, respectively, all are rendered in an italicised font with serif, and any relations, such as typing and reduction, are marked by a subscript.

5.3.1 Relation to Classical Processes

One notable omission in the literature is any proof that Priority CP is an extension of CP or that Priority CLL (PLL) is an extension of CLL. To discuss this matter formally, we must define what we mean by “extension”. What does it mean for one system to extend another? We consider two options:

- If PLL extends the *proofs* of CLL, any valid CLL proof is a valid PLL proof, and any well-typed CP process is a well-typed PCP process.
- If PLL extends the *theorems* of CLL, any proposition provable in CLL is provable in PLL, and any type inhabited in CP is inhabited in PCP.

An extension of the *proofs* is a much stronger notion than an extension of the *theorems*, and, from the perspective of a process calculus, it is much more useful. Unfortunately, the stronger notion does not hold, as we discuss shortly. At the time of writing, I do not know if the weaker notion—an extension of theorems—holds.

The question is complicated by the priority annotations. What does it mean for a CLL proposition to be provable in PLL? Does it suffice if the proposition is provable with *some* priority assignment? Or should it be provable with *every* priority assignment? Priorities leak a fair amount of the structure of the underlying proof. For instance, any proposition where the priorities decrease as we descend into the proposition is unprovable, other than by absurdity or the axiom, e.g. there is no proof of (a) even though there is a proof of (b):

$$(a) \not\vdash \mathbf{1}^1 \otimes \mathbf{3}^3 \mathbf{1}^2 \quad (b) \vdash \mathbf{1}^2 \otimes \mathbf{1}^1 \mathbf{3}^3$$

Hence, it would seem overly strict to require that a CLL proposition must be provable in PLL with *every* priority assignment.

PCP Conflates Tensor and Par

As Dardha and Gay [2018] discuss, PLL admits Mix and Mix_0 , under the names T-PAR and T-HALT . Consequently, it conflates one/bottom and partially conflates tensor/par

$$1^p \multimap^o \perp^q \quad \text{and} \quad A \otimes^p B \multimap^o A \wp^q B$$

For instance, the conversion of tensor to par holds by the following derivations, which hold as long as $o < p, q$ and either $p < q$ or $q < p$.

$$\begin{array}{c}
 \frac{\frac{\frac{\frac{\overline{\vdash \bar{A}, A}}{\vdash \bar{A}, \bar{B}, A, B}}{\vdash \bar{A}, \bar{B}, A \wp^q B}}{\vdash \bar{A} \wp^p \bar{B}, A \wp^q B}}{\vdash \bar{A} \wp^p \bar{B} \wp^o A \wp^q B}}{\vdash A \otimes^p B \multimap^o A \wp^q B} \quad \begin{array}{l} q < \text{pr}(A, B) \\ p < q \\ o < p, q \end{array}
 \end{array}
 \quad
 \begin{array}{c}
 \frac{\frac{\frac{\frac{\overline{\vdash \bar{A}, A}}{\vdash \bar{A}, \bar{B}, A, B}}{\vdash \bar{A} \wp^p \bar{B}, A, B}}{\vdash \bar{A} \wp^p \bar{B}, A \wp^q B}}{\vdash \bar{A} \wp^p \bar{B} \wp^o A \wp^q B}}{\vdash A \otimes^p B \multimap^o A \wp^q B} \quad \begin{array}{l} p < \text{pr}(A, B) \\ q < p \\ o < p, q \end{array}
 \end{array}$$

These conflations are the natural consequence of admitting Mix and Mix_0 , even in the absence of priorities. PLL also admits MULTICUT , which can be derived from T-RES and T-PAR , which fully conflates tensor/par

$$A \wp^p B \multimap^o A \otimes^q B$$

The conversion of par to tensor holds by the following derivations, which hold as long as $o < p, q$ and either $p < q$ or $q < p$.

$$\begin{array}{c}
 \frac{\frac{\frac{\frac{\overline{\vdash \bar{A}, A}}{\vdash \bar{A}, \bar{B}, A, B}}{\vdash \bar{A}, \bar{B}, A \otimes^q B}}{\vdash \bar{A} \otimes^p \bar{B}, A \otimes^q B}}{\vdash \bar{A} \otimes^p \bar{B} \wp^o A \otimes^q B}}{\vdash A \wp^p B \multimap^o A \otimes^q B} \quad \begin{array}{l} q < \text{pr}(A, B) \\ p < q \\ o < p, q \end{array}
 \end{array}
 \quad
 \begin{array}{c}
 \frac{\frac{\frac{\frac{\overline{\vdash \bar{A}, A}}{\vdash \bar{A}, \bar{B}, A, B}}{\vdash \bar{A} \otimes^p \bar{B}, A, B}}{\vdash \bar{A} \otimes^p \bar{B}, A \otimes^q B}}{\vdash \bar{A} \otimes^p \bar{B} \wp^o A \otimes^q B}}{\vdash A \wp^p B \multimap^o A \otimes^q B} \quad \begin{array}{l} p < \text{pr}(A, B) \\ q < p \\ o < p, q \end{array}
 \end{array}$$

In CLL, tensor and par capture independence and interdependence:

- If $\vdash_c \Gamma, A \otimes B$, the resources used by A and B are independent.
- If $\vdash_c \Delta, A \wp B$, the resources used by A and B are interdependent. How they depend on one another is unknown, but they are guaranteed to be free from cyclic dependencies.

These are naturally compositional. In $\vdash_c \Delta, A \wp B$, we do not know how the resources in A and B are used, so the only option that is guaranteed to avoid cyclic dependencies is to ensure that in $\vdash_c \Gamma, A \otimes B$ the resources used by A and B are independent.

PLL's “tensor” and “par” are not truly a tensor and par, and do not capture independence and interdependence. Morally, they are the

same self-dual connective, which may capture either independence or interdependence, but must specify *exactly* if and how its resources interdepend. The motivation to separate this self-dual connective into two dual connectives, “ $\otimes^\circ$ ” and “ $\wp^\circ$ ”, comes from the process calculus interpretation, rather than the logic itself, as the separate connectives are used to guarantee session fidelity under the interpretation where “ $\otimes^\circ$ ” is send and “ $\wp^\circ$ ” is receive.

PCP Does Not Extend CP

The priorities in PCP expose information about the order in which a process uses its resources. Consequently, the additives in PCP are less expressive than the additives in CP, and, while we have not discussed the exponentials in this thesis, the exponentials in PCP are also less expressive than in CP.

Let us consider the typing rules for the offer in CP and PCP:

$$\frac{P \vdash_c \Gamma, x : A \quad Q \vdash_c \Gamma, x : B}{x \triangleright \{\text{inl}: P; \text{inr}: Q\} \vdash_c \Gamma, x : A \& B} \text{T-OFFER}$$

$$\frac{P \vdash \Gamma, x : A \quad Q \vdash \Gamma, x : B \quad \circ < \text{pr}(\Gamma)}{x \triangleright \{\text{inl}: P; \text{inr}: Q\} \vdash \Gamma, x : A \&^\circ B} \text{T-OFFER}$$

The two typing rules ostensibly quite similar, with the only difference being the added priority constraint in PCP’s typing rule. However, recall that, while the processes P and Q may differ in their use of the sessions A and B , the priority annotations on the typing environment Γ mean that both processes must use all other resources in the exact same order. CP’s type system, on the other hand, imposes no such restriction.

Counterexample 5.1 (Extension). *It is not true that every well-typed CP process is a well-typed PCP process.*

For example, the process

$$a \triangleright \{\text{inl}: x().y().a[]; \text{inr}: y().x().a[]\}$$

is typeable in CP but not PCP.

In CP, its typing derivation is as follows:

$$\frac{\frac{\frac{a[] \vdash_c a : \mathbf{1}}{y().a[] \vdash_c y : \perp, a : \mathbf{1}}}{x().y().a[] \vdash_c x : \perp, y : \perp, a : \mathbf{1}} \quad \frac{\frac{\frac{a[] \vdash_c a : \mathbf{1}}{x().a[] \vdash_c x : \perp, a : \mathbf{1}}}{y().x().a[] \vdash_c x : \perp, y : \perp, a : \mathbf{1}}}{a \triangleright \{\text{inl}: x().y().a[]; \text{inr}: y().x().a[]\} \vdash_c x : \perp, y : \perp, a : \mathbf{1} \& \mathbf{1}}$$

In PCP, the process is not typeable, because the typing derivation requires that $p < q$ and $q < p$, where p and q are the priorities associated with the

actions on x and y .

$$\frac{\frac{\frac{a[] \vdash a : \mathbf{1}^{o_1} \quad q < o_1}{y().a[] \vdash y : \perp^q, a : \mathbf{1}^{o_1} \quad p < q}}{x().y().a[] \vdash x : \perp^p, y : \perp^q, a : \mathbf{1}^{o_2}} \quad \frac{\frac{\frac{a[] \vdash a : \mathbf{1}^{o_2} \quad p < o_2}{x().a[] \vdash x : \perp^p, a : \mathbf{1}^{o_1} \quad q < p}}{y().x().a[] \vdash x : \perp^p, y : \perp^q, a : \mathbf{1}^{o_2}}}{a \triangleright \{\text{inl: } x().y().a[]; \text{inr: } y().x().a[]\} \vdash x : \perp^p, y : \perp^q, a : \mathbf{1}^{o_1} \&^o \mathbf{1}^{o_2}}$$

Therefore, not every well-typed CP process is a well-typed PCP process.

For the additives, we can partially work around this restriction by combining the entire typing environment into a single type, offering a session of type (omitting the priority annotations on the $\wp$'s, to reduce eyestrain)

$$(\perp^{p_1} \wp \perp^{p_1} \wp \mathbf{1}^{o_1}) \&^o (\perp^{p_2} \wp \perp^{q_2} \wp \mathbf{1}^{o_2})$$

However, this requires a global process transformation and changes the type of the session.

A similar example exists for PCP's exponentials. We can construct two sessions of type $\perp^p \wp^o \perp^q$ that use the sessions corresponding to p and q in opposite orders, and apply dereliction and contracting the resulting sessions. The process

$$?x[x_1, x_2]. (\nu z \bar{z})(?x_1[x]. x(y). x(). y(). z[] \parallel ?x_2[x]. x(y). y(). x(). \bar{z}(). a[])$$

is well-typed in CP with exponentials, using the typing rules and alternative notation from Wadler [2014, 3.4], reproduced below:

$$\frac{\text{T-DERELICT} \quad \frac{P \vdash_c \Gamma, y : A}{?x[y]. P \vdash_c \Gamma, x : ?A}}{\quad} \quad \frac{\text{T-CONTRACT} \quad \frac{P \vdash_c \Gamma, x_1 : ?A, x_2 : ?A}{?x[x_1, x_2]. P \vdash_c \Gamma, x : ?A}}{\quad}$$

However, it is not typeable in PCP, using the typing rules from Dardha and Gay [2018, Figure 2], as typing derivation requires that $p < q$ and $q < p$, reproduced below:

$$\frac{\text{T-DERELICT} \quad \frac{P \vdash \Gamma, y : A \quad o < \text{pr}(\Gamma)}{?x[y]. P \vdash \Gamma, x : ?^o A}}{\quad} \quad \frac{\text{T-CONTRACT} \quad \frac{P \vdash \Gamma, x_1 : ?^{o_1} A, x_2 : ?^{o_2} A \quad o \leq o_1 \quad o \leq o_2 \quad o < \text{pr}(\Gamma)}{?x[x_1, x_2]. P \vdash \Gamma, x : ?^o A}}{\quad}$$

This restriction might be worked around to some extent by using the additives to encode the different usages of the typing environment and the duplicated session. However, this requires a global process transformation, and changes the type of the session.

PGV Does Not Extend GV

The counterexample to extension given for CP and PCP in the previous section is easily adapted to provide a counterexample to extension for GV and PGV.

Counterexample 5.2 (Extension). *It is not true that every well-typed GV term is a well-typed PGV term.*

For example, the term

$$\begin{aligned} & \text{case } x \{ \text{inl } x \mapsto x; M; \text{inr } x \mapsto x; N \} \\ & \text{where} \\ & \quad M \triangleq \text{wait } y; \text{wait } z \\ & \quad N \triangleq \text{wait } z; \text{wait } y \end{aligned}$$

is typeable in GV but not in PGV.

In GV, using the typing rules in Figure I.2, its typing derivation is as follows. (In the following derivation, the premises for TM-APP are stacked and the first premise of TM-LETUNIT is omitted.)

$$\frac{\begin{array}{c} \vdots \\ y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} M : \mathbf{1} \\ x : \mathbf{1}, y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} x; M : \mathbf{1} \end{array} \quad \frac{\begin{array}{c} \vdots \\ y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} N : \mathbf{1} \\ x : \mathbf{1}, y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} x; N : \mathbf{1} \end{array}}{x : \mathbf{1} + \mathbf{1}, y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} \text{case } x \{ \text{inl } x \mapsto x; M; \text{inr } x \mapsto x; N \} : \mathbf{1}}$$

where the typing derivations for M and N are

$$\frac{\frac{\frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{y : \text{end}_? \vdash_{\text{GV}} y : \text{end}_?} \quad \frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{z : \text{end}_? \vdash_{\text{GV}} z : \text{end}_?}}{y : \text{end}_? \vdash_{\text{GV}} \text{wait } y : \mathbf{1}} \quad \frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{z : \text{end}_? \vdash_{\text{GV}} \text{wait } z : \mathbf{1}}}{y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} M : \mathbf{1}}$$

and

$$\frac{\frac{\frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{z : \text{end}_? \vdash_{\text{GV}} z : \text{end}_?} \quad \frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{y : \text{end}_? \vdash_{\text{GV}} y : \text{end}_?}}{z : \text{end}_? \vdash_{\text{GV}} \text{wait } z : \mathbf{1}} \quad \frac{\vdash_{\text{GV}} \text{wait} : \text{end}_? \multimap \mathbf{1}}{y : \text{end}_? \vdash_{\text{GV}} \text{wait } y : \mathbf{1}}}{y : \text{end}_?, z : \text{end}_? \vdash_{\text{GV}} N : \mathbf{1}}$$

respectively.

In PGV, using the typing rules in Figure II.2, its typing derivation is as follows. (In the following derivation, the premises for T-APP are stacked and the first premise of T-LETUNIT is omitted.)

$$\frac{\begin{array}{c} \vdots \\ y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} M : \mathbf{1} \\ x : \mathbf{1}, y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} x; M : \mathbf{1} \end{array} \quad \frac{\begin{array}{c} \vdots \\ y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} N : \mathbf{1} \\ x : \mathbf{1}, y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} x; N : \mathbf{1} \end{array}}{x : \mathbf{1} + \mathbf{1}, y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} \text{case } x \{ \text{inl } x \mapsto x; M; \text{inr } x \mapsto x; N \} : \mathbf{1}}$$

where the typing derivations for M and N are

$$\frac{\frac{\frac{\vdash^\perp \text{wait} : \text{end}_?^p \multimap^{\tau, p} \mathbf{1}}{y : \text{end}_?^p \vdash^\perp y : \text{end}_?^p} \quad \frac{\vdash^\perp \text{wait} : \text{end}_?^q \multimap^{\tau, q} \mathbf{1}}{z : \text{end}_?^q \vdash^\perp z : \text{end}_?^q}}{y : \text{end}_?^p \vdash^p \text{wait } y : \mathbf{1}} \quad \frac{\vdash^\perp \text{wait} : \text{end}_?^q \multimap^{\tau, q} \mathbf{1}}{z : \text{end}_?^q \vdash^q \text{wait } z : \mathbf{1}} \quad p < q}{y : \text{end}_?^p, z : \text{end}_?^q \vdash^{\text{p} \sqcup q} M : \mathbf{1}}$$

and

$$\begin{array}{c}
 \frac{\frac{\vdash^\perp \text{wait} : \text{end}_?^q \multimap^{\top, q} \mathbf{1}}{z : \text{end}_?^q \vdash^\perp z : \text{end}_?^q}}{z : \text{end}_?^q \vdash^q \text{wait } z : \mathbf{1}} \quad \frac{\frac{\vdash^\perp \text{wait} : \text{end}_?^p \multimap^{\top, p} \mathbf{1}}{y : \text{end}_?^p \vdash^\perp y : \text{end}_?^p}}{y : \text{end}_?^p \vdash^p \text{wait } y : \mathbf{1}} \quad q < p \\
 \hline
 y : \text{end}_?^p, z : \text{end}_?^q \vdash^{p \sqcup q} N : \mathbf{1}
 \end{array}$$

respectively, which require that $p < q$ and $q < p$, respectively, where p and q are the priorities associated with the actions on y and z , respectively.

Therefore, not every well-typed GV term is a well-typed PGV term.

There is no corresponding counterexample for the exponentials, since the presentation of PGV in Paper II omits the exponentials. However, any variant of PGV that introduces replication, following the work by Lindley and Morris [2015] and Dardha and Gay [2018], without any further changes, would certainly admit such an example.

My Priorities Leave Me No Choice

What is the issue with priorities that causes PCP and PGV to be non-extensions of CP and GV, respectively?

The priorities impose a linear order in which a collection of resources is to be used. For instance, under the typing environment $x : A^1, y : B^2$, the resource x must be used first and the resource y must be used second. When the same typing environment is shared by different branches in a program, as is the case for the additives and the exponentials, each branch is required to use those resources in the same order, which fundamentally limits the expressiveness of branching. The structure offered by the ordering on priorities—which are the natural numbers extended with a lower and upper bound—is insufficient to capture different uses across branches.

To address this issue, we would have to extend the structure of priorities to be able to capture branching, e.g. by adding branching structure to the priorities, or by marking constraints by the choices under which they must hold. Any solution would likely be closely related to the introduction of proof-boxes in the proof nets for linear logic [Girard, 1987, p. 43].

On the other side, since the issue is that it is not generally possible to impose a linear order on processes with branching, this is a good indication that priorities *do* extend the multiplicative fragment of linear logic, as well as the fragment with fixed points [e.g. Lindley and Morris, 2016b], since neither introduces branching.

Identity Expansion Fails for Priority CP

The type system for PCP does not satisfy identity expansion.

Counterexample 5.3 (Identity Expansion). *It is not true that, if there exists a proof of $\vdash \Gamma$ that uses T-LINK, then there exists a proof of $\vdash \Gamma$ that does not use T-LINK.*

For example, given the following proof

$$\frac{}{\vdash \mathbf{1}^1 \otimes^3 \mathbf{1}^2, \perp^1 \wp^3 \perp^2} T\text{-LINK}$$

there exists no proof with the same conclusion that does not use T-LINK.

The reason is that PCP's type system omits certain constraints that are, generally speaking, syntactically impossible to violate. For instance, the typing rule T-RECV introduces the type $A \wp^o B$, but does not require that o is smaller than the priorities on A and B . This constraint is not required, since a process that uses an endpoint of type $A^1 \wp^3 B^2$ would have to act on the endpoint of type A^1 before receiving it. I refer to these constraints as the *priority tree* of a type, since they capture the tree structure of the type. For instance, for the type $A^p \wp^o B^q$, the priority tree contains the constraints $o < p$ and $o < q$ and all constraints imposed by the structure of A and B . The priority tree of a type, written T_A , is a set of constraints, defined as follows:

$$\begin{aligned} T_{A \otimes^o B}, T_{A \wp^o B}, T_{A \oplus^o B}, T_{A \&^o B} &\triangleq \{o < \text{pr}(A), o < \text{pr}(B)\} \cup T_A \cup T_B \\ T_{\mathbf{1}^o}, T_{\perp^o}, T_{0^o}, T_{\top^o} &\triangleq \emptyset \end{aligned}$$

PCP's type system omits the constraints in the priority tree for all types. This includes the types introduced by links and absurd offers. Unfortunately, the syntactic argument does not apply in these cases. Consequently, PCP's type system allows links, such as the process in Counterexample 5.3, whose priorities violate the constraints imposed by the priority trees. Since no process, other than the link and the absurd offer, can introduce such types, these links are, in essence, dead, as they can only ever be connected to other links or absurd offers, and, as mentioned, they do not satisfy identity expansion.

To address this issue, I conjecture that it suffices to require that the constraints in priority tree for all types introduced by links and absurd offers must hold.

5.3.2 Priority Inference

In this section, I introduce priority inference for PCP. Usually, a type inference system is presented as an algorithmic variant of the type system, whose derivations take terms as input and produce types as output. Such type inference derivations are compositional, and produce locally correct information.

Deadlock freedom is a global property. While CP and HCP guarantee deadlock freedom locally and compositionally, they do so by enforcing a much stronger and more restrictive invariant: the tree and forest structure of the connection graph. In essence, PCP guarantees global deadlock freedom simply by checking it, *globally*. Consequently, PCP's typing derivations are not meaningfully compositional. Priorities witness global deadlock freedom, and, as such, capture global information. Morally, if you connect two PCP processes, you must re-check if they are deadlock free. (Practically, you can connect two PCP processes if and only if you already chose their priorities to witness the global deadlock freedom of the resulting process.) As typing derivations are not compositional, there is little hope for a local, compositional priority inference system for PCP. Instead, I structure priority inference as a two-stage process, which factors out the local, compositional portion into the first stage, and defers the global check to the second stage.

Priority inference is structured as follows:

1. I define pre-processes and pre-types, which replace priorities with priority metavariables, and pre-typing, which ensures communication safety, but not deadlock freedom. Crucially, pre-typing is local and compositional.
2. I prove that a pre-process is a process only if it is deadlock-free.

In this section, processes, types, and priorities are printed in red, blue, and yellow, respectively, and all three are rendered in a sans-serif font, whereas pre-processes, pre-types, and priority metavariables are printed in pink, green, and periwinkle, respectively, and all three are rendered in an italicised font with serif. The pre-typing sequent is marked by a subscript “PI”.

Priority metavariables are names, whereas priorities are natural numbers. Let o , p , and q range over priority metavariables. Priority metavariables should be unique, i.e. no priority metavariable should occur more than once in a pre-process, pre-type, pre-typing environment, or pre-typing derivation. I use Barendregt's convention for priority metavariables, and assume this uniqueness, rather than explicitly renaming duplicate priority metavariables.

Pre-processes are the same as processes, but annotated by pre-types rather than types. Let x , y , z , and w range over endpoint names, and let P , Q , and R range over pre-processes. Binding for pre-processes is the same as binding for processes.

Pre-types are the same as types, but annotated by priority metavariables rather than priorities. Pre-types are well-formed if and only if each connective is annotated with a distinct priority metavariable. Let A , B , C , and D as well as $\bar{A}$, $\bar{B}$, $\bar{C}$, and $\bar{D}$ range over pre-types. Following our convention for endpoints, I write A and $\bar{A}$ to imply that the types

associated with these pre-types are dual, which I define shortly. The set of priority metavariables in a pre-type A , written $\text{fp}(A)$, is the set of all priority metavariables that occur in the pre-type A . Two pre-types are *equivalent*, written $A \approx B$, if they are equal up to priority metavariables.

$$\begin{array}{ll}
 A_1 \otimes^p A_2 \approx B_1 \otimes^q B_2 & \iff A_1 \approx B_1 \wedge A_2 \approx B_2 & \mathbf{1}^p \approx \mathbf{1}^q \\
 A_1 \wp^p A_2 \approx B_1 \wp^q B_2 & \iff A_1 \approx B_1 \wedge A_2 \approx B_2 & \perp^p \approx \perp^q \\
 A_1 \oplus^p A_2 \approx B_1 \oplus^q B_2 & \iff A_1 \approx B_1 \wedge A_2 \approx B_2 & \mathbf{0}^p \approx \mathbf{0}^q \\
 A_1 \&^p A_2 \approx B_1 \&^q B_2 & \iff A_1 \approx B_1 \wedge A_2 \approx B_2 & \top^p \approx \top^q
 \end{array}$$

Two pre-types are dual, written $A \sim \bar{A}$, if they are dual up to priority metavariables.

$$\begin{array}{ll}
 A_1 \otimes^p A_2 \sim \bar{A}_1 \wp^q \bar{A}_2 & \iff A_1 \sim \bar{A}_1 \wedge A_2 \sim \bar{A}_2 & \mathbf{1}^p \sim \perp^q \\
 A_1 \wp^p A_2 \sim \bar{A}_1 \otimes^q \bar{A}_2 & \iff A_1 \sim \bar{A}_1 \wedge A_2 \sim \bar{A}_2 & \perp^p \sim \mathbf{1}^q \\
 A_1 \oplus^p A_2 \sim \bar{A}_1 \&^q \bar{A}_2 & \iff A_1 \sim \bar{A}_1 \wedge A_2 \sim \bar{A}_2 & \mathbf{0}^p \sim \top^q \\
 A_1 \&^p A_2 \sim \bar{A}_1 \oplus^q \bar{A}_2 & \iff A_1 \sim \bar{A}_1 \wedge A_2 \sim \bar{A}_2 & \top^p \sim \mathbf{0}^q
 \end{array}$$

Pre-typing environments are the same as typing environments, but contain pre-type assignments, rather than type assignments. Let Γ and Δ range over pre-typing environments. Pre-typing environments are well-formed if and only if the priority metavariables in all pre-types are distinct. The set of priority metavariables in a pre-typing environment Γ , written $\text{fp}(\Gamma)$, is the set of priority metavariables that occur in all the pre-types in the pre-typing environment Γ .

Priority graph (ranged over by G, H) are mixed graphs whose vertices are priority metavariables. I informally revisit the relevant definitions. For a detailed discussion, see § A.1.

- A *mixed graph* G has a set of vertices (denoted V_G , ranged over by u, v), a set of edges (denoted E_G), a set of arcs (denoted A_G). Edges are unordered pairs denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$. The set of edges may not contain loops uu . Arcs are ordered pairs denoted by juxtaposition overset with an arrow to indicate the direction, i.e. $\overrightarrow{uv} \triangleq (u, v)$. The set of arcs may not contain loops $\overrightarrow{uu}$.
- For any graph G with vertices $u, v \in V_G$, u is *adjacent* to v when there exists some edge $uv \in E_G$ or some arc $\overrightarrow{uv} \in A_G$.
- A *walk* w is a sequence of pairwise adjacent vertices.
- A *path* p is a walk with no repeated vertices, except possibly the first and last.
- A *cycle* c is a path that begins and ends at the same vertex.
- A walk is *essentially directed* when it contains at least one arc.
- A graph is *essentially acyclic* if and only if it contains no essentially directed cycles.
- The *undirected reachability* relation (denoted by $\sim_G$) is the equivalence closure over E_G .
- The *essentially directed reachability* relation (denoted by $<_G$) is the transitive closure over A_G quotiented by $\sim_G$.

- The *empty graph* with vertices V , written $\bar{K}_V$, is the graph consisting of vertices V , with no edges or arcs.
- The *graph union* of G_1 and G_2 , denoted by $G_1 \cup G_2$, is defined by, for each projection, taking the union of the projection of G_1 and G_2 , e.g. $V_{G_1 \cup G_2} \triangleq V_{G_1} \cup V_{G_2}$, $E_{G_1 \cup G_2} \triangleq E_{G_1} \cup E_{G_2}$, etc..
- The *directed rooting* of G in v , written $v < G$, is the graph formed by adding the vertex v and adding an arc from v to every other vertex of G , i.e. $V_{v < G} \triangleq \{v\} \cup V_G$ and $A_{v < G} \triangleq \{\vec{v}u \mid u \in V_G\} \cup A_G$. Directed rooting preserves the remaining projections, e.g. $E_{v < G} \triangleq E_G$.

The *priority tree* of a pre-type, written T_A , is the priority graph whose vertices are the priority metavariables in A , with arcs flowing along the structure of the type.

$$\begin{array}{l} T_{A \otimes^o B}, T_{A \wp^o B}, T_{A \oplus^o B}, T_{A \&^o B} \triangleq O < (T_A \cup T_B) \\ T_{1^o}, T_{\perp^o}, T_{0^o}, T_{\top^o} \triangleq O \end{array}$$

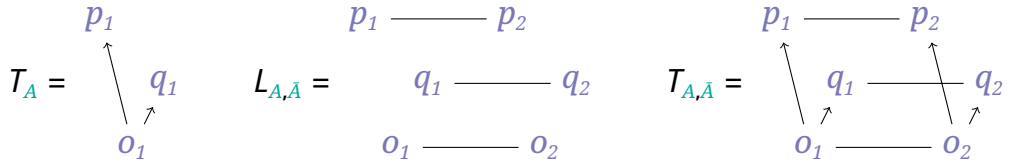
The *priority link* of two pre-types, written $L_{A, \bar{A}}$, is the bipartite priority graph with edges between the corresponding priority metavariables in A and $\bar{A}$. The priority link is defined if and only if $A \sim \bar{A}$.

$$\begin{array}{l} L_{A \otimes^p B, \bar{A} \wp^q \bar{B}}, L_{A \wp^p B, \bar{A} \otimes^q \bar{B}}, L_{A \oplus^p B, \bar{A} \&^q \bar{B}}, L_{A \&^p B, \bar{A} \oplus^q \bar{B}} \triangleq pq \cup L_{A, \bar{A}} \cup L_{B, \bar{B}} \\ L_{1^p, \perp^q}, L_{\perp^p, 1^q}, L_{0^p, \top^q}, L_{\top^p, 0^q} \triangleq pq \end{array}$$

The *priority link-tree* of two pre-types, written $T_{A, \bar{A}}$, is the union of the two priority trees for A and $\bar{A}$ and the priority link for A and $\bar{A}$.

$$T_{A, \bar{A}} \triangleq T_A \cup T_{\bar{A}} \cup L_{A, \bar{A}}$$

To illustrate these definitions, let us look at an example. Let A and $\bar{A}$ be the pre-types $1^{p_1} \otimes^{o_1} 1^{q_1}$ and $\perp^{p_1} \wp^{o_2} \perp^{q_2}$, respectively. The priority tree T_A , priority link $L_{A, \bar{A}}$, and priority link-tree $T_{A, \bar{A}}$ are as follows:



The pre-typing judgment $P \vdash_{\pi} \Gamma \mid G$ means that P is well-typed if, for each pre-type assignment $x : A$ in Γ , exactly one pre-process in P uses the endpoint x according to the session pre-type A . The priority graph G is an output of the pre-typing derivation.

Definition 5.4 (Pre-Typing). A pre-process P is well-typed under some pre-typing environment Γ if there exists a pre-typing derivation with conclusion $P \vdash_{\pi} \Gamma \mid G$ for some G that uses the pre-typing rules in Figure 5.1.

A *priority substitution* assigns a priority to each priority metavariable. Let σ range over priority substitutions. A priority substitution translates pre-processes to processes, pre-types to types, and pre-typing environments

$$\begin{array}{c}
\frac{}{x \leftrightarrow y \vdash_{PI} x : A, y : \bar{A} \mid L_{A, \bar{A}}} \text{PI-LINK} \qquad \frac{P \vdash_{PI} \Gamma, x : A, \bar{x} : \bar{A} \mid G}{(\nu x \bar{x})P \vdash_{PI} \Gamma \mid L_{A, \bar{A}} \cup G} \text{PI-RES} \\
\\
\frac{P \vdash_{PI} \Gamma \mid G \quad Q \vdash_{PI} \Delta \mid H}{P \parallel Q \vdash_{PI} \Gamma, \Delta \mid G \cup H} \text{PI-PAR} \qquad \frac{}{0 \vdash_{PI} \emptyset \mid \bar{K}_{\emptyset}} \text{PI-HALT} \\
\\
\frac{P \vdash_{PI} \Gamma, y : A, x : B \mid G}{x[y].P \vdash_{PI} \Gamma, x : A \oplus^o B \mid o < G} \text{PI-SEND} \qquad \frac{P \vdash_{PI} \Gamma, y : A, x : B \mid G}{x(y).P \vdash_{PI} \Gamma, x : A \wp^o B \mid o < G} \text{PI-RECV} \\
\\
\frac{P \vdash_{PI} \Gamma \mid G}{x[].P \vdash_{PI} \Gamma, x : \mathbf{1}^o \mid o < G} \text{PI-CLOSE} \qquad \frac{P \vdash_{PI} \Gamma \mid G}{x().P \vdash_{PI} \Gamma, x : \perp^o \mid o < G} \text{PI-WAIT} \\
\\
\frac{P \vdash_{PI} \Gamma, x : A \mid G}{x \triangleleft \text{inl}.P \vdash_{PI} \Gamma, x : A \oplus^o B \mid o < G} \text{PI-SELECT}_1 \\
\\
\frac{P \vdash_{PI} \Gamma, x : B \mid G}{x \triangleleft \text{inr}.P \vdash_{PI} \Gamma, x : A \oplus^o B \mid o < G} \text{PI-SELECT}_2 \\
\\
\frac{P \vdash_{PI} \Gamma, x : A \mid G_1 \quad Q \vdash_{PI} \Gamma, x : B \mid G_2}{x \triangleright \{\text{inl}; P; \text{inr}; Q\} \vdash_{PI} \Gamma, x : A \&^o B \mid o < (G_1 \cup G_2)} \text{PI-OFFER} \\
\\
\frac{N = \text{fn}(\Gamma)}{x \not\downarrow N \vdash_{PI} \Gamma, x : \top^o \mid \bar{K}_{\text{fn}(\Gamma, x : \top^o)}} \text{PI-ABSURD}
\end{array}$$

Figure 5.1: Pre-Typing Rules for Priority CP

to typing environments, by pointwise applying the priority substitution to each priority metavariable.

Priority inference is *sound*. If the priority graph produced by the pre-typing derivation is essentially acyclic, there exists some priority substitution such that the resulting process is typeable in PCP.

Proposition 5.5 (Soundness). *If $P \vdash_{\text{PT}} \Gamma \mid G$ and G is essentially acyclic, then there exists some priority substitution σ such that $\sigma(P) \vdash \sigma(\Gamma)$.*

Proof. Let H be the quotient graph of G by its edges E_G , i.e. G/E_G . Consequently, H has no edges and is a directed graph. Since G is essentially acyclic, H is acyclic. By topological sort, there exists a linear ordering S_n of the vertices of H such that, if $\vec{u}\vec{v} \in A_H$, then u comes before v in S_n . Let the priority substitution $\sigma : V_G \rightarrow \mathbb{N}$ be the function that maps each priority metavariable to its position in the linear ordering S_n , i.e. $\sigma = \{\sigma \mapsto i \mid \sigma \in S_i\}$.

The typing derivation for $\sigma(P) \vdash \sigma(\Gamma)$ is constructed by induction on the typing derivation for $P \vdash_{\text{PT}} \Gamma \mid G$.

- In the cases for PI-SEND, PI-RECV, PI-CLOSE, PI-WAIT, PI-SELECT₁, PI-SELECT₂, PI-OFFER, and PI-ABSURD, the result follows from the induction hypothesis and the rules T-SEND, T-RECV, T-CLOSE, T-WAIT, T-SELECT₁, T-SELECT₂, T-OFFER, and T-ABSURDZAP, respectively. The constraint $\sigma < \text{pr}(\Gamma)$ is satisfied by the definition of σ and the rooting of the priority graph $\sigma < G$.
- In the case for PI-LINK, the result follows from the rule PI-LINK. In the case for PI-RES, the result follows from the induction hypothesis and the rule PI-RES. In both cases, duality is satisfied by the definition of σ and the priority link $L_{A,\bar{A}}$.
- In the case for PI-PAR, the result follows from the induction hypotheses and the rule T-PAR.
- In the case for PI-HALT, the result follows from the rule T-HALT.

□

Priority erasure replaces priority annotations with fresh priority metavariables, and translates processes to pre-processes, types to pre-types, and typing environments to pre-typing environments, by pointwise applying the priority erasure to each priority metavariable.

Definition 5.6 (Priority Erasure). Priority erasure, written $[\cdot]_{\text{PE}}$, maps processes to pre-processes, types to pre-types, typing environments to pre-typing environments, and typing derivations to pre-typing derivations, by replacing all priorities with fresh priority metavariables.

Priority inference satisfies *completeness*. If the process is well-typed in PCP, then the priority graph produced by the pre-typing derivation

is essentially acyclic, and the order of the original priority assignment respects reachability in the priority graph.

Proposition 5.7 (Completeness). *If $P \vdash \Gamma$, then $\lfloor P \rfloor_{\pi} \vdash_{\pi} \lfloor \Gamma \rfloor_{\pi} \mid G$ such that the produced priority graph G is essentially acyclic, and, if σ is the priority substitution such that $\sigma(P) = P$ and $\sigma(\Gamma) = \Gamma$, then $p <_G q \implies \sigma(p) < \sigma(q)$.*

Proof. By induction on the derivation of $P \vdash \Gamma$. In each case, we construct the priority substitution σ by matching the priority metavariables in the erased types to the priorities on the original type, and prove that σ is a graph homomorphism from the produced priority graph G to the graph corresponding to the order on the natural numbers. Consequently, σ witnesses the fact that G is linearisable, and therefore is essentially acyclic. □

5.4 Conclusion

In this chapter, we introduced Priority GV with its typing rules, reduction semantics, and metatheory, and we revisited Priority CP with its typing rules, reduction semantics, and metatheory, and introduced an operational correspondence between PCP and PGV. Priority CP and Priority GV are variants of CP and GV that use priority typing to permit processes with benign cyclic connections. For PGV, we proved *preservation* (Theorem II.3.5) and *global progress* (Theorem II.3.14). For PCP, we dropped the commuting conversions, which cause Dardha and Gay’s PCP to be non-confluent, from the reduction semantics, we added the additive units, and we proved *preservation* (Theorem II.3.5), and proved that *global progress* (Theorem II.3.14) continues to hold. We related PCP to PGV by means of a translation, proved that it preserves types (Theorem II.4.6), and proved that it gives rise to a complete (Theorem II.4.7) and sound (Theorem II.4.10) operational correspondence. I demonstrated that PCP and PGV are not extensions of CP and GV, respectively, and that PCP does not satisfy identity expansion. Finally, I introduced priority inference for PCP.

I conjecture that multiplicative PCP—the fragment $(\otimes, 1, \wp, \perp)$ —does extend multiplicative CP, and proving this would be an interesting topic for future work. It is easily verified that most typing rules for multiplicative CP preserve essential acyclicity of the priority graph. The exception is cut. I conjecture that a successful proof for cut proceeds by induction on the cut formula, and relies on *splitting*—the property that the subgraphs of the priority graph which are reachable from the arguments of a tensor are disjoint.

It would be interesting to investigate the exact relation between priority

graphs and proof nets. Certainly, they are not equivalent, since priority graphs encode the order in which the sequent calculus rules are used. However, I conjecture that if priority graphs were extended with type information, they would contain sufficient information to reconstruct the entire typing derivation, and hence, they encode the entire process.

It would be interesting to extend PCP and PGV with fixed points, following Lindley and Morris [2016b] and Padovani [2014], and with first- and second-order quantifiers, interpreted as priority and type polymorphism.

Finally, it would be interesting to investigate solutions for the expressivity of additives in priority-based calculi by extending the structure of priority graphs with choice, e.g. by boxing, as done in the proof nets of linear logic [Girard, 1987].

Part IV

Deadlock-Free Session Types in Linear Haskell

Chapter 6

Implementations

In previous chapters, I have made the claim that GV’s session-typed communication is easy to implement as a library, especially when the host language already supports linear types. To evidence this claim, this chapter presents an implementation of both HGV and PGV’s session-typed concurrency primitives as a library, named `priority-sesh`, implemented in Linear Haskell [Marlow et al., 2010, Bernardy et al., 2018].

The library only implements HGV and PGV’s concurrency primitives. This is a double-edged sword. Only implementing the concurrency primitives makes GV easy to implement, but also easy to implement incorrectly, since this makes a number of assumptions about the host language that are difficult to prove.

- It relies on the host language to provide the basic ingredients of a functional language—e.g. functions, pairs, sums, etc.—and basic concurrency primitives—i.e. threads and channels.
- It assumes that the host language and GV agree on their evaluation strategy, or that GV’s correctness guarantees are invariant under the choice of evaluation strategy. (For Haskell, we can say with confidence that the first option does not hold, since GV’s semantics are call-by-value, whereas Haskell’s semantics are call-by-need.)
- It assumes that the host language and GV agree on the semantics for threads and channels.

The bulk of the chapter consists of the paper *Deadlock-Free Session Types in Linear Haskell* by Kokke and Dardha [2021b], hereafter referred to as Paper III. References made from the main body of this thesis into Paper III will be prefixed by an “III”, e.g. “Theorem III.3.1”. This chapter proceeds as follows:

- In § 6.1, I provide a legend and an errata for Paper III.
- In § 6.2, we present an implementation of both HGV and PGV’s session-typed communication as a library in Linear Haskell. This

section consists entirely of Paper III, and proceeds as follows:

- In § III.2.1, we implement one-shot channels.
- In § III.2.2, we implement session-typed channels.
- In § III.2.3, we show that a restriction of the interface defined in § III.2.2 corresponds to HGV’s communication primitives.
- In § III.2.4, we implement priority-based channels.
- In § III.3, we show that the interface defined in § III.2.4 is the monadic reflection of PGV’s communication primitives.
- In § III.4, we discuss the related work. Notably, we compare a variety of session types in Haskell, extending the analysis provided by Orchard and Yoshida [2017].
- In § III.5, we discuss our experience using Linear Haskell as well as potential directions for future work.

6.1 Legend and Errata

The conventions and terminology in Paper I are different from those used in the rest of this thesis.

- The concurrency primitives, session types, and priorities in the Linear Haskell implementation as well as the terms, types, and priorities of Priority GV are printed in **red**, **blue**, and **green**, respectively, and are rendered in an italicised or bolded font with serif.

There are several other notable differences between HGV and PGV and their implementations in the library.

First, the library implements several extensions to GV.

- The library implements session cancellation, following Exceptional GV [EGV, Fowler et al., 2019].

Strictly speaking, this means the implementation is *affine* rather than *linear*, in the sense that values can be left unused. This is important in any practical implementation of session types, regardless of what typing facilities are offered by the host language. We can design a session-typed language in the simplified setting where communication always succeeds and no process ever crashes, but this assumption is immediately shattered upon contact with reality.

This extension *also* allows us to implement GV as a library in host languages whose type systems are affine rather than linear, such as the predecessor to the library, *Rusty Variation* by Kokke [2019], which is implemented in Rust [Matsakis and Klock, 2014].

- The library has polymorphic and recursive session types, simply

by virtue of being implemented in Haskell. The recursive sessions in the library do not provide any termination guarantees. This is unsurprising, since Haskell does not offer any termination guarantees. Hence, the recursive sessions in the library do not correspond to the recursive sessions in Lindley and Morris’ μ GV, but they may correspond to the variant that identifies the greatest and least fixed points [Lindley and Morris, 2016b, § 2.3, under “Nontermination”].

Second, the library makes several simplifications that are common in practical variants of GV, which relax the correspondence with CLL.

- The `link` primitive is omitted from both the implementation of session-typed channels (in § III.2.2) and the implementation of priority-based session-typed channels (in § III.2.4). This is common in implementations of GV, or variants of GV which do not aim to prove a correspondence with a variant of CP.
- The type system conflates `endl` and `endr` into `End`. This is equivalent to postulating `Mix` and `Mix0` in CLL, i.e. it preserves deadlock-freedom, but relaxes the structure of the connection graph from a tree to a forest. Consequently, processes may become disconnected. Consequently, exceptions raised in disconnected child processes are not propagated to the main process. This is not a serious problem, as it cannot affect the outcome of the main computation, and does not justify the added complexity of separating `endl` and `endr`. (If the library’s concurrency primitives are used together with other forms of inter-process communication, exceptions raised in disconnected child processes can affect the outcome of the main computation, but such uses already lose the guarantee of deadlock-freedom.)

Finally, the library makes several changes that are required to correctly embed HGV and PGV into Linear Haskell.

- The library relies on Linear Haskell to provide linearity checking. However, Linear Haskell and GV have different notions of linearity. In GV, linearity is a property of *values*. If a value is linear, it must be used exactly once. In Linear Haskell, linearity is a property of *functions*. If a function is linear, it must use its argument exactly once. To accommodate this difference, it is important that session-typed endpoints, which must be used linearly, are only ever obtained as the arguments of linear functions.

In HGV’s implementation, the concurrency primitive `fork` is replaced with the equivalent concurrency primitive `connect` from Wadler’s GV (see § III.2.3). Whereas `fork` returns the dual endpoint, `connect` passes it to another continuation. As endpoints can only be

created by `connect`, this guarantees they are used linearly.

$$\text{fork} : (S \multimap \text{end}_l) \multimap \bar{S} \quad \text{connect} : (S \multimap \text{end}_l) \multimap (\bar{S} \multimap T) \multimap T$$

(The type of “`connect`” in the library differs from the type of `connect` presented above. We shall discuss this shortly.)

In PGV’s implementation, the concurrency primitives are lifted into a linear graded monad (see “`Sesh`”, in § III.2.4, under “The Communication Monad”). As the bind operation for this monad is linear, this guarantees that endpoints are used linearly.

- In HGV’s implementation, the type of “`connect`” differs from the type of Wadler’s `connect`, as presented above.

$$\text{connect} :: \text{Session } s \Rightarrow (s \multimap \text{IO } ()) \multimap (\text{Dual } s \multimap \text{IO } a) \multimap \text{IO } a$$

The function corresponding to the child thread has type $s \multimap \text{IO } ()$ instead of $S \multimap \text{end}_l$, which means it returns the unit, rather than an endpoint of type `endl`.

The change of type is due to two relaxations of GV’s typing. The first is the conflation of `endl` and `endr` into `End`, discussed above. The second is the switch from the synchronous `End` to the asynchronous unit, which lets child processes terminate asynchronously, without synchronising with their parents.

- In PGV’s implementation, the concurrency connectives are lifted into the graded linear monad Sesh_p^q , which permits us to encode PGV’s constraints in Haskell’s type system as type-level constraints.

There are minor errors in Paper III:

- A phrase in the first paragraph on Page III.6 reads “For instance, for *totallyFine* we can assign the number 0 to `send (chs1)` and `recv (chr2)`, and 1 to `send (chs2)` and `recv (chr1)`”, but should read “For instance, for *totallyFine* we can assign the number 0 to `send (chs1)` and `recv (chr1)`, and 1 to `send (chs2)` and `recv (chr2)`”.

6.2 Paper III: Deadlock-Free Session Types in Linear Haskell

This section contains the paper with the same title, in collaboration with Ornela Dardha, which was originally published in the proceedings of the ACM SIGPLAN International Symposium on Haskell (Haskell 2021).

The work presented in the paper was conceived of by all the authors. I implemented the Haskell library, developed the refined version of Priority GV's type system and its monadic reflection, wrote the initial draft of the paper, and co-authored the comparison with existing Haskell libraries for session-typed channels.



Deadlock-Free Session Types in Linear Haskell

Wen Kokke
University of Edinburgh
Edinburgh, Scotland
wen.kokke@ed.ac.uk

Ornela Dardha
University of Glasgow
Glasgow, Scotland
ornela.dardha@glasgow.ac.uk

Abstract

Priority Sesh is a library for session-typed communication in Linear Haskell which offers strong compile-time correctness guarantees. Priority Sesh offers two deadlock-free APIs for session-typed communication. The first guarantees deadlock freedom by restricting the process structure to trees and forests. It is simple and composable, but rules out cyclic structures. The second guarantees deadlock freedom via priorities, which allows the programmer to safely use cyclic structures as well.

Our library relies on Linear Haskell to guarantee linearity, which leads to easy-to-write session types and more idiomatic code, and lets us avoid the complex encodings of linearity in the Haskell type system that made previous libraries difficult to use.

CCS Concepts: • Theory of computation → Linear logic; Type theory.

Keywords: session types, linear haskell, deadlock freedom

ACM Reference Format:

Wen Kokke and Ornela Dardha. 2021. Deadlock-Free Session Types in Linear Haskell. In *Proceedings of the 14th ACM SIGPLAN International Haskell Symposium (Haskell '21), August 26–27, 2021, Virtual, Republic of Korea*. ACM, New York, NY, USA, 13 pages. <https://doi.org/10.1145/3471874.3472979>

1 Introduction

Session types are a type formalism used to specify and verify communication protocols [26–28, 62]. They’ve been studied extensively in the context of the π -calculus [58], a process calculus for communication an concurrency, and in the context of concurrent λ -calculi, such as the GV family of languages [“Good Variation”, 20, 23, 40, 64].

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
Haskell '21, August 26–27, 2021, Virtual, Republic of Korea
© 2021 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-8615-9/21/08...\$15.00
<https://doi.org/10.1145/3471874.3472979>

Session types have been implemented in various programming languages. We give a detailed overview in section 4, and Orchard and Yoshida [49] provide a complete survey of session type implementations in Haskell.

The main difficulty when implementing session types in most programming languages is *linearity*, i.e., the guarantee that each channel endpoint is used *exactly once*. There are several different approaches to guaranteeing linearity, but the main distinction is between *dynamic* [52, 59, 60] and *static* [41, 42, 56] usage checks. With dynamic checks, using a channel endpoint more than once simply throws a runtime error. With static checks, usage is *somehow* encoded into the type system of the host language usually by encoding the entire linear typing environment into the type system using a parameterised or graded monad. Such encodings are only possible if the type system of the host language is expressive enough. However, such encodings are often quite complex, and result in a trade-off between easy-to-write session types and idiomatic programs.

Moreover, these implementations only focus on the most basic features of session types and often ignore more advanced ones, such as channel delegation or deadlock freedom: Neubauer and Thiemann [44] only provide single session channels; Pucella and Tov [56] provide multiple channels, but only the building blocks for channel delegation; Imai et al. [30] extend Pucella and Tov [56] and provide full delegation. None of these works address deadlock freedom. Lindley and Morris [41] provide an implementation of GV into Haskell building on the work of Polakow [55]. To the best of our knowledge, this is the only work that guarantees deadlock freedom of session types in Haskell, albeit in a simple form. In GV, all programs must have *tree-shaped* process structures. The process structure of a program is an undirected graph, where nodes represent processes, and edges represent the channels connecting them. (We explore this in more detail in section 2.3.) Therefore, deadlock freedom is guaranteed by design: session types rule out deadlocks over a single channel, and the tree-restriction rules out sharing multiple channels between two processes. While Lindley and Morris [41] manage to implement more advanced properties, the tree restriction rules out many interesting programs which have *cyclic* process structure, but are deadlock free.

Recent works by Padovani and Novara [53] and Kokke and Dardha [PGV, 35] integrate *priorities* [32, 51] into functional languages. Priorities are natural numbers that abstractly represent the time at which a communication action happens.

Priority-based type systems check that there are *no* cycles in the communication graph. The communication graph is a directed graph where nodes represent dual communication actions, and directed edges represent one action must happen before another. (We explore this in more detail in section 2.4.) Such type systems are *more* expressive, as they allow programs to have *cyclic* process structure, as long as they have an *acyclic* communication graph.

With the above in mind, our research goals are as follows:

- Q1 Can we have easy-to-write session types, easy linearity checks and idiomatic code at the same time?
- Q2 Can we address not only the main features of session types, but also advanced ones, such as full delegation, recursion, and deadlock freedom of programs with cyclic process structure?

Our priority-sesh library answers both questions *mostly* positively. We sidestep the problems with encoding linearity in Haskell by using Linear Haskell [4], which has native support for linear types. The resulting session type library presented in sections 2.2 and 2.3 has both easy-to-write session types, easy linearity checks, and idiomatic code. Moving to Q2, the library has full delegation, recursion, and the variant in section 2.3 guarantees deadlock freedom, albeit by restricting the process structure to trees and forests. In section 2.4, we implement another variant which uses priorities to ensure deadlock freedom of programs with cyclic processes structure. The ease-of-writing suffers a little, as the programmer has to manually write priorities, though this isn't a huge inconvenience. Unfortunately, GHC's ability to reason about type-level naturals currently is *not* as powerful as to allow the programmer to easily write priority-polymorphic code, which is required for recursion. Hence, while we address deadlock freedom for cyclic process structures, we do so only for the finite setting.

Contributions. In section 2, we present Priority Sesh, an implementation of deadlock free session types in Linear Haskell which is:

- the *first* implementation of session types to take advantage of Linear Haskell for linearity checking, and producing easy-to-write session types and more idiomatic code;
- the *first* implementation of session types in Haskell to guarantee deadlock freedom of programs with cyclic process structure via *priorities*; and
- the *first* embedding of priorities into an existing mainstream programming language.

In section 3, we:

- present a variant of Priority GV [35]—the calculus upon which Priority Sesh is based—with asynchronous communication and session cancellation following Fowler et al. [20] and *explicit* lower bounds on the

sequent, rather than lower bounds inferred from the typing environment; and

- show that Priority Sesh is related to Priority GV via monadic reflection.

2 What is Priority Sesh?

In this section we introduce Priority Sesh in three steps:

- in section 2.1, we build a small library of *linear* or *one-shot channels* based on MVars [54];
- in section 2.2, we use these one-shot channels to build a small library of *session-typed channels* [12]; and
- in section 2.4, we decorate these session types with *priorities* to guarantee deadlock-freedom [35].

It is important to notice that the meaning of linearity in *one-shot channels* differs from linearity in *session channels*. A linear or one-shot channel originates from the linear π -calculus [33, 58], where each endpoint of a channel must be used for *exactly one* send or receive operation, whereas linearity in the context of session-typed channels, it means that each step in the protocol is performed *exactly once*, but the channel itself is used multiple times.

Priority Sesh is written in Linear Haskell [4]. The type $\multimap$ is syntactic sugar for the linear arrow $\%1 \multimap$. Familiar definitions refer to linear variants packaged with `linear-base`¹ (e.g., `IO`, `Functor`, `Bifunctor`, `Monad`) or with Priority Sesh (e.g., `MVar`).

We colour the Haskell definitions which are a part of Sesh: **red** for functions and constructors; **blue** for types and type families; and **emerald** for priorities and type families acting on priorities.

2.1 One-shot Channels

We start by building a small library of *linear* or *one-shot channels*, i.e., channels that must be use *exactly once* to send or receive a value.

The one-shot channels are at the core of our library, and their efficiency is crucial to the overall efficiency of Priority Sesh. However, we do not aim to present an efficient implementation here, rather we aim to present a compact implementation with the correct behaviour.

Channels. A one-shot channel has two endpoints, `Send1` and `Recv1`, which are two copies of the same `MVar`.

`newtype Send1 a = Send1 (MVar a)`

`newtype Recv1 a = Recv1 (MVar a)`

`new1 :: IO (Send1 a, Recv1 a)`

`new1 = do`

`(mvars, mvarr) ← dup2 <$> newEmptyMVar`

`return (Send1 (unur mvars), Recv1 (unur mvarr))`

The `newEmptyMVar` function returns an *unrestricted* `MVar`, which may be used non-linearly, i.e., as many times as one

¹<https://hackage.haskell.org/package/linear-base>

wants. The `dup2` function creates two (unrestricted) copies of the `MVar`. The `unur` function casts each *unrestricted* copy to a *linear* copy. Thus, we end up with two copies of an `MVar`, each of which must be used *exactly once*.

We implement `send1` and `recv1` as aliases for the corresponding `MVar` operations.

```
send1 :: Send1 a → a → IO ()
send1 (Send1 mvars) x = putMVar mvars x
recv1 :: Recv1 a → IO a
recv1 (Recv1 mvarr) = takeMVar mvarr
```

The `MVar` operations implement the correct blocking behaviour for asynchronous one-shot channels: the `send1` operation is non-blocking, and the `recv1` operations blocks until a value becomes available.

Synchronisation. We use `Send1` and `Recv1` to implement a construct for one-shot synchronisation between two processes, `Sync1`, which consists of two one-shot channels. To synchronise, each process sends a unit on the one channel, then waits to receive a unit on the other channel.

```
data Sync1 = Sync1 (Send1 ()) (Recv1 ())
newSync1 :: IO (Sync1, Sync1)
newSync1 = do (chs1, chr1) ← new1
              (chs2, chr2) ← new1
              return (Sync1 chs1 chr2, Sync1 chs2 chr1)
sync1 :: Sync1 → IO ()
sync1 (Sync1 chs chr) = do send1 chs (); recv1 chr
```

Cancellation. We implement *cancellation* for one-shot channels. One-shot channels are created in the linear `IO` monad, so *forgetting* to use a channel results in a complaint from the type-checker. However, it is possible to *explicitly* drop values whose types implement the *Consumable* class, using `consume :: a → ()`. The ability to cancel communications is important, as it allows us to safely throw an exception *without violating linearity*, assuming that we cancel all open channels before doing so.

One-shot channels implement *Consumable* by simply dropping their `MVars`. The Haskell runtime throws an exception when a “thread is blocked on an `MVar`, but there are no other references to the `MVar` so it can’t ever continue.”² Practically, `consumeAndRecv` throws a `BlockedIndefinitelyOnMVar` exception, whereas `consumeAndSend` does not:

```
consumeAndRecv = do
  (chs, chr) ← new1
  fork $ return (consume chs)
  recv1 chr
consumeAndSend = do
  (chs, chr) ← new1
  fork $ return (consume chr)
  send1 chs ()
```

²<https://downloads.haskell.org/~ghc/9.0.1/docs/html/libraries/base-4.15.0.0/Control-Exception.html#t:BlockedIndefinitelyOnMVar>

Where `fork` forks off a new thread using a linear `forkIO`. (In GV, this operation is called `spawn`.)

As the `BlockedIndefinitelyOnMVar` check is performed by the runtime, it’ll even happen when a channel is dropped for reasons other than consume, such as a process crashing.

2.2 Session-typed Channels

We use the one-shot channels to build a small library of *session-typed channels* based on the *continuation-passing style* encoding of session types in linear types by Dardha [9], Dardha et al. [12] and in line with other libraries for Scala [59, 60], OCaml [52], and Rust [34].

An Example. Let’s look at a simple example of a session-typed channel—a multiplication service, which receives two integers, sends back their product, and then terminates:

```
type MulServer = Recv Int (Recv Int (Send Int End))
type MulClient = Send Int (Send Int (Recv Int End))
```

We define `mulServer`, which acts on a channel of type `MulServer`, and `mulClient`, which acts on a channel of the *dual* type:

```
mulServer (s :: MulServer)  mulClient (s :: MulClient)
= do (x, s) ← recv s        = do s ← send (32, s)
  (y, s) ← recv s           s ← send (41, s)
  s ← send (x * y, s)       (z, s) ← recv s
  close s                   close s
  return ()                 return z
```

In order to encode the *sequence* of a session type using one-shot types, each action on a session-typed channel returns a channel for the *continuation* of the session—save for `close`, which ends the session. Furthermore, `mulServer` and `mulClient` act on endpoints with *dual* types. Duality is crucial to session types as it ensures that when one process sends, the other is ready to receive, and vice versa. This is the basis for communication safety guaranteed by a session type system.

Channels. We start by defining the `Session` type class, which has an associated type `Dual`. You may think of `Dual` as a type-level function associated with the `Session` class with *one* case for each instance. We encode the various restrictions on duality as constraints on the type class. Each session type must have a dual, which must itself be a session type—`Session (Dual s)` means the dual of `s` must also implement `Session`. Duality must be *injective*—the annotation `result → s` means `result` must uniquely determine `s` and *involutive*—`Dual (Dual s) ~ s` means `Dual (Dual s)` must equal `s`. These constraints are all captured by the `Session` class, along with `new` for constructing channels:

```
class (Session (Dual s), Dual (Dual s) ~ s) ⇒ Session s
  where
```

```

type Dual s = result | result → s
new :: IO (s, Dual s)

```

There are three primitive session types: *Send*, *Recv*, and *End*.

```

newtype Send a s = Send (Send1 (a, Dual s))
newtype Recv a s = Recv (Recv1 (a, s))
newtype End      = End Sync1

```

By following Dardha et al. [12], a channel *Send* wraps a one-shot channel *Send₁* over which we send some value—which is the intended value sent by the session channel, and the channel over which the communicating partner process continues the session—it'll make more sense once you read the definition for *send*. A channel *Recv* wraps a one-shot channel *Recv₁* over which we receive some value and the channel over which we continue the session. Finally, an channel *End* wraps a synchronisation.

We define duality for each session type—*Send* is dual to *Recv*, *Recv* is dual to *Send*, and *End* is dual to itself:

```

instance Session s ⇒ Session (Send a s)
  where
    type Dual (Send a s) = Recv a (Dual s)
    new = do (chs, chr) ← new1
           return (Send chs, Recv chr)

```

```

instance Session s ⇒ Session (Recv a s)
  where
    type Dual (Recv a s) = Send a (Dual s)
    new = do (chs, chr) ← new1
           return (Recv chr, Send chs)

```

```

instance Session End
  where
    type Dual End = End
    new = do (chsync1, chsync2) ← newSync1
           return (End chsync1, End chsync2)

```

The *send* operation constructs a channel for the continuation of the session, then sends one endpoint of that channel, along with the value, over its one-shot channel, and returns the other endpoint:

```

send :: Session s ⇒ (a, Send a s) → IO s
send (x, Send chs) = do (here, there) ← new
                        send1 chs (x, there)
                        return here

```

The *recv* and *close* operations simply wrap their corresponding one-shot operations:

```

recv :: Recv a s → IO (a, s)
recv (Recv chr) = recv1 chr

close :: End → IO ()
close (End chsync) = sync1 chsync

```

Cancellation. We implement session *cancellation* via the *Consumable* class. For convenience, we provide the *cancel* function:

```

cancel :: Session s ⇒ s → IO ()
cancel s = return (consume s)

```

As with one-shot channels, *consume* simply drops the channel, and relies on the *BlockedIndefinitelyOnMVar* check, which means that *cancelAndRecv* throws an exception and *cancelAndSend* does not:

```

cancelAndRecv = do
  (chs, chr) ← new
  fork $ cancel chs
  ((), ()) ← recv chr
  return ()

cancelAndSend = do u
  (chs, chr) ← new
  fork $ cancel chr
  () ← send chs ()
  return ()

```

These semantics correspond to EGV [20].

Asynchronous Close. We don't always want session-end to involve synchronisation. Unfortunately, the *close* operation is synchronous.

An advantage of defining session types via a type class is that its an *open* class, and we can add new primitives whenever. Let's make the unit type, *()*, a session type:

```

instance Session s ⇒ Session ()
  where
    type Dual () = ()
    new = return ((), ())

```

Units are naturally affine—they contain zero information, so dropping them won't harm—and the linear *Monad* class allows you to silently drop unit results of monadic computations. They're ideal for *asynchronous* session end!

Using *()* allows us to recover the semantics of one-shot channels while keeping a session-typed language for idiomatic protocol specification.

Choice. So far, we've only presented sending, receiving, and synchronisation. It is, however, possible to send and receive *channels* as well as values, and we leverage that to implement most other session types by using these primitives only!

For instance, we can implement *binary* choice by sending/receiving *Either* of two session continuations:

```

type Select s1 s2 = Send (Either (Dual s1) (Dual s2)) ()
type Offer s1 s2 = Recv (Either s1 s2) ()

selectLeft :: (Session s1) ⇒ Select s1 s2 → IO s1
selectLeft s = do (here, there) ← new
                  send (Left there, s)
                  return here

offerEither :: Offer s1 s2 → (Either s1 s2 → IO a) → IO a
offerEither s match = do (e, ()) ← recv s; match e

```

Differently from `()`, we don't have to implement the `Session` class for `Select` and `Offer`. They're already session types!

Recursion. We can write recursive session types by writing them as recursive Haskell types. Unfortunately, we cannot write recursive type synonyms, so we have to use a newtype. For instance, we can write the type for a recursive summation service, which receives numbers until the client indicates they're done, and then sends back the sum. We specify *two* newtypes:

```
newtype SumSrv
  = SumSrv (Offer (Recv Int SumSrv) (Send Int End))
newtype SumCnt
  = SumCnt (Select (Send Int SumCnt) (Recv Int End))
```

We implement the summation server as a recursive function:

```
sumSrv :: Int -> SumSrv -> IO ()
sumSrv tot (SumSrv s) = offerEither s $ \e. case x of
  Left s -> do (x, s) <- recv s; sumSrv (tot + x) s
  Right s -> do s <- send (tot, s); close s
```

As `SumSrv` and `SumCnt` are new types, we must provide instances of the `Session` class for them.

```
instance Session SumSrv
  where
    type Dual SumSrv = SumCnt
    new = do (chsrv, chcnt) <- new
           return (SumSrv chsrv, SumCnt chcnt)
```

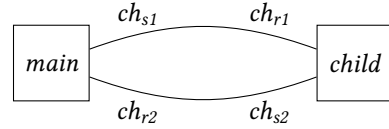
2.3 Deadlock Freedom via Process Structure

The session-typed channels presented in section 2.2 can be used to write deadlocking programs, e.g., by receiving before sending:

```
woops :: IO Void
woops = do (chs1, chr1) <- new
          (chs2, chr2) <- new
          fork $ do (void, ()) <- recv chr1
                  send (void, chs2)
          (void, ()) <- recv chr2
          let (void, voidcopy) = dup2 void
          send (void, chs1)
          return voidcopy
```

Counter to what the type says, this program doesn't actually produce an inhabitant of the *uninhabited* type `Void`. Instead, it deadlocks! We'd like to help the programmer avoid such programs.

As discussed in section 1, we can *structurally* guarantee deadlock freedom by ensuring that the *process structure* is always a tree or forest. The process structure of a program is an undirected graph, where nodes represent processes, and edges represent the channels connecting them. For instance, the process structure of `woops` is cyclic:



This restriction works by ensuring that between two processes there is *at most one* (series of) channels over which the two can communicate. As duality rules out deadlocks on any one channel, such configurations must be deadlock free.

We can rule out cyclic process structures by hiding *new*, and only exporting *connect*, which creates a new channel and, *crucially*, immediately passes one endpoint to a new thread:

```
connect :: Session s =>
  (s -> IO ()) -> (Dual s -> IO a) -> IO a
connect k1 k2 = do (s1, s2) <- new; fork (k1 s1); k2 s2
```

You can view *connect* as the node constructor for a binary process tree. If the programmer *only* uses *connect*, their process structure is guaranteed to be a *tree*. If they also use standalone *fork*, their process structure is a *forest*. Either way, their programs are guaranteed to be deadlock free.

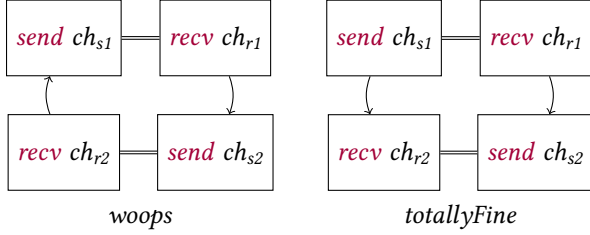
2.4 Deadlock Freedom via Priorities

The strategy for deadlock freedom presented in section 2.3 is simple, but *very* restrictive, since it rules out *all* cyclic communication structures, even the ones which don't deadlock:

```
totallyFine :: IO String
totallyFine = do (chs1, chr1) <- new
                (chs2, chr2) <- new
                fork $ do (x, ()) <- recv chr1
                        send (x, chs2)
                send ("Hiya!", chs1)
                (x, ()) <- recv chr2
                return x
```

This process has *exactly the same* process structure as `woops`, but it's totally fine, and returns "Hiya!" as you'd expect. We'd like to enable the programmer to write such programs while still guaranteeing their programs don't deadlock.

As discussed in section 1, there is another way to rule out deadlocks—by using *priorities*. Priorities are an approximation of the *communication graph* of a program. The communication graph of a program is a *directed graph* where nodes represent *actions on channels*, and directed edges represent that one action happens before the other. Dual actions are connected with double undirected edges. (You may consider the graph contracted along these edges.) If the communication graph is cyclic, the program deadlocks. The communication graphs for `woops` and `totallyFine` are as follows:



If the communication graph is acyclic, then we can assign each node a number such that directed edges only ever point to nodes with *bigger* numbers. For instance, for *totallyFine* we can assign the number 0 to *send* ch_{s1} and *recv* ch_{r2} , and 1 to *recv* ch_{r1} and *send* ch_{s2} . These numbers are *priorities*.

In this section, we present a type system in which *priorities* are used to ensure deadlock freedom, by tracking the time a process starts and finishes communicating using a graded monad [21, 48]. The bind operation registers the order of its actions in the type, requiring the sequentiality of their duals.

Priorities. The priorities assigned to *communication actions* are always natural numbers, which represent, *abstractly*, at which time the action happens. When tracking the start and finish times of a program, however, we also use $\perp$ and $\top$ for programs which don't communicate. These are used as the identities for $\sqcap$ and $\sqcup$ in lower and upper bounds, respectively. We let o range over natural numbers, p over *lower bounds*, and q over *upper bounds*.

data *Priority* = $\perp \mid \text{Nat} \mid \top$

We define strict inequality ($<$), minimum ($\sqcap$), and maximum ($\sqcup$) on priorities as usual.

Channels. We define *Send* ^{o} , *Recv* ^{o} , and *End* ^{o} , which decorate the *raw* sessions from section 2.2 with the priority o of the communication action, i.e., it denoted when the communication happens. Duality (*Dual*) preserves these priorities. These are implemented exactly as in section 2.2.

The Communication Monad. We define a graded monad *Sesh* _{p} ^{q} , which decorates *IO* with a lower bound p and an upper bound q on the priorities of its communication actions, i.e., if you run the monad, it denotes when communication begins and ends.

newtype *Sesh* _{p} ^{q} $a = \text{Sesh} \{ \text{runSeshIO} :: \text{IO } a \}$

The monad operations for *Sesh* _{p} ^{q} merely wrap those for *IO*, hence trivially obeys the monad laws.

The *ireturn* function returns a *pure* computation—the type *Sesh* _{$\perp$} ^{$\top$} guarantees that all communications happen between $\top$ and $\perp$, hence there can be no communication at all.

ireturn :: $a \rightarrow \text{Sesh}_{\perp}^{\top} a$
ireturn $x = \text{Sesh} \$ \text{return } x$

The $\gg$ operator sequences two actions with types *Sesh* _{p} ^{q} and *Sesh* _{p'} ^{q'} , and requires $q < p'$, i.e., the first action must

have finished before the second starts. The resulting action has lower bound $p \sqcap p'$ and upper bound $q \sqcup q'$.

$(\gg) :: (q < p') \Rightarrow \text{Sesh}_p^q a \rightarrow (a \rightarrow \text{Sesh}_{p'}^{q'} b) \rightarrow \text{Sesh}_{p \sqcap p'}^{q \sqcup q'} b$
 $mx \gg mf = \text{Sesh} \$ \text{runSeshIO } mx \gg \lambda x. \text{runSeshIO } (mf \ x)$

In what follows, we implicitly use $\gg$ with do-notation. This can be accomplished in Haskell using *RebindableSyntax*.

We define decorated variants of the concurrency and communication primitives: *send*, *recv*, and *close* each perform a communication action with some priority o , and return a computation of type *Sesh* _{o} ^{o} , i.e., with *exact* bounds; *new* and *cancel* don't perform any communication action, and so return a *pure* computation of type *Sesh* _{$\perp$} ^{$\top$} ; *fork* takes a computation which performs communication actions as an argument, forks it off into a separate thread, and masks the upper bound in its return type.

new :: *Session* $s \Rightarrow \text{Sesh}_{\perp}^{\top} (s, \text{Dual } s)$
fork :: *Sesh* _{p} ^{q} $() \rightarrow \text{Sesh}_{\perp}^{\top} ()$
cancel :: *Session* $s \Rightarrow s \rightarrow \text{Sesh}_{\perp}^{\top} ()$
send :: *Session* $s \Rightarrow (a, \text{Send}_o^o a \ s) \rightarrow \text{Sesh}_o^o s$
recv :: *Recv* ^{o} $a \ s \rightarrow \text{Sesh}_o^o (a, s)$
close :: *End* ^{o} $\rightarrow \text{Sesh}_o^o ()$

From these, we derive decorated choice, as before:

type *Select* ^{o} $s_1 \ s_2 = \text{Send}_o^o (\text{Either } (\text{Dual } s_1) (\text{Dual } s_2)) ()$
type *Offer* ^{o} $s_1 \ s_2 = \text{Recv}_o^o (\text{Either } s_1 \ s_2) ()$
selectLeft :: (*Session* s_1) $\Rightarrow \text{Select}_o^o s_1 \ s_2 \rightarrow \text{Sesh}_o^o s_1$
selectRight :: (*Session* s_2) $\Rightarrow \text{Select}_o^o s_1 \ s_2 \rightarrow \text{Sesh}_o^o s_2$
offerEither :: ($o < p$) $\Rightarrow \text{Offer}_o^o s_1 \ s_2 \rightarrow$
 $(\text{Either } s_1 \ s_2 \rightarrow \text{Sesh}_p^q a) \rightarrow \text{Sesh}_{o \sqcap p}^{o \sqcup q} a$

Safe IO. We can use a trick from the *ST* monad [38] to define a “pure” variant of *runSesh*, which encapsulates all use of *IO* within the *Sesh* _{p} ^{q} monad. The idea is to index the *Sesh* _{p} ^{q} and every session type constructor with an extra type parameter *tok*, which we'll call the *session token*:

send :: *Session* $s \Rightarrow (a, \text{Send}_o^o \text{ tok } a \ s) \rightarrow \text{Sesh}_o^o \text{ tok } s$
recv :: *Recv* ^{o} $\text{ tok } a \ s \rightarrow \text{Sesh}_o^o \text{ tok } (a, s)$
close :: *End* ^{o} $\text{ tok } \rightarrow \text{Sesh}_o^o \text{ tok } ()$

The session token should never be instantiated, except by *runSesh*, and every action under the same call to *runSesh* should use the same type variable *tok* as its session token:

runSesh :: $(\forall \text{ tok}. \text{Sesh}_p^q \text{ tok } a) \rightarrow a$
runSesh $x = \text{unsafePerformIO } (\text{runSeshIO } x)$

This ensures that none of the channels created in the session can escape out of the scope of *runSesh*.

We implement this encapsulation in *priority-sesh*, though the session token is the first argument, preceding the priority bounds.

Recursion. We could implement recursive session via priority-polymorphic types, or via priority-shifting [53]. For instance, we could give the *summation* service from section 2.2 the following type:

```
newtype SumSrvo
  = SumSrv (Offero (Recvo+1 Int (SumSrvo+2))
               (Sendo+1 Int (Endo+2)))
```

We'd then like to assign *sumSrv* the following type:

```
sumSrv : Int → SumSrvo → Seshoo⊤ ()
sumSrv tot (SumSrv s) = offerEither s $ λe. case x of
  Left s → do (x, s) ← recv s; sumSrv (tot + x) s
  Right s → do s ← send (tot, s); weaken (close s)
```

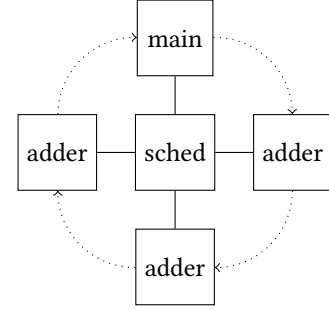
The upper bound for a recursive call should be $\top$, which ensures that recursive calls are only made in *tail* position [3, 22]. The recursive call naturally has upper bound $\top$. However, the *close* operation happens at some *concrete* priority $o + n$, which needs to be raised to $\top$, so we'd have to add a primitive *weaken*: $Sesh_p^q a \rightarrow Sesh_p^\top a$.

Unfortunately, writing such priority-polymorphic code relies heavily on GHC's ability to reason about type-level naturals, and GHC rejects *sumSrv* complaining that it cannot verify that $o < o + 1$, $o + 1 < o + 2$, etc. There's several possible solutions for this:

1. We could embrace the Hasochism [39], and provide GHC with explicit evidence, though this would make priority-sesh more difficult to use.
2. We could delegate some of these problems to a GHC plugin such as `type-nat-solver`³ or `ghc-typelits-presburger`⁴. Unfortunately, $\square$ and $\sqcup$ are beyond Presburger arithmetic, and `type-nat-solver` has not been maintained in recent years.
3. We could attempt to write type families which reduce in as many cases as possible. Unfortunately, a restriction in closed type families [16, §6.1] prevents us from checking *exactly these cases*.

Currently, the prioritised sessions don't support recursion, and implementing one of these solutions is future work.

Cyclic Scheduler. Dardha and Gay [10] and Kokke and Dardha [36] use a *finite* cyclic scheduler as an example. The cyclic scheduler has the following process structure, with the flow of information indicated by the dotted arrows:



We start by defining the types of the channels which connect each client process to the scheduler:

```
type SRo1o1 a = Sendo1 a (Recvo2 a ())
type RSo2o1 a = Dual (SRo2o1 a)
```

We then define the scheduler itself, which forwards messages from one process to the next in a cycle:

```
sched :: RS70 a → SR21 a → SR43 a → SR65 a → Sesh70⊤ ()
sched s1 s2 s3 s4 = do
  (x, s1) ← recv s1
  s2 ← send (x, s2); (x, ()) ← recv s2
  s3 ← send (x, s3); (x, ()) ← recv s3
  s4 ← send (x, s4); (x, ()) ← recv s4
  send (x, s1)
```

Finally, we define the *adder* and the *main* processes. The *adder* adds one to the value it receives, and the *main* process initiates the cycle and receives the result:

```
adder :: (o1 < o2) ⇒ RSo2o1 Int → Sesho2o1⊤ ()
adder s = do (x, s) ← recv s; send (x + 1, s)
main :: (o1 < o2) ⇒ Int → SRo2o1 Int → Sesho2o1⊤ Int
main x s = do ; s ← send (x, s); (x, ()) ← recv s; ireturn x
```

While the process structure of the cyclic scheduler *as presented* isn't cyclic, nothing prevents the user from adding communications between the various client processes, or from removing the scheduler and having the client processes communicate *directly* in a ring.

3 Relation to Priority GV

The *priority-sesh* library is based on a variant of Priority GV [35], which differs in three ways:

1. it marks lower bounds *explicitly* on the sequent, rather than implicitly inferring them from the typing environment;
2. it collapses the isomorphic types for session end, end_i^o and end_i^o , into end_i^o ;
3. it is extended with asynchronous communication and session cancellation following Fowler et al. [20].

These changes preserve subject reduction and progress properties, and give us *tighter* bounds on priorities. To see why, note that PCP [10] and PGV [35] use the *smallest* priority in the typing environment as an approximation for the lower

³<https://github.com/yav/type-nat-solver>

⁴<https://hackage.haskell.org/package/ghc-typelits-presburger>

bound. Unfortunately, this *underestimates* the lower bound in the rules T-VAR and T-LAM (check fig. 1). These rules type *values*, which are pure and could have lower bound $\top$, but the smallest priority in their typing environment is not necessarily $\top$.

Priority GV. We briefly revisit the syntax and type system of PGV, but a full discussion of PGV is out of scope for this paper. For a discussion of the *synchronous* semantics for PGV, and the proofs of subject reduction, progress, and deadlock freedom, please see Kokke and Dardha [35]. For a discussion of the *asynchronous* semantics and session cancellation, please see Fowler et al. [20].

As in section 2.4, we let o range over priorities, which are natural numbers, and p and q over priority bounds, which are either natural numbers, $\top$, or $\perp$.

PGV is based on the standard linear λ -calculus with product types $(\cdot \times \cdot)$, sum types $(\cdot + \cdot)$, and their units ($\mathbf{1}$ and $\mathbf{0}$). Linear functions $(\cdot \multimap_p^q \cdot)$ are annotated with priority bounds which tell us—when the function is applied—when communication begins and ends.

Types and session types are defined as follows:

$$\begin{aligned} S &:: !^o T.S \mid ?^o T.S \mid \mathbf{end}^o \\ T, U &:: T \times U \mid \mathbf{1} \mid T + U \mid \mathbf{0} \mid T \multimap_p^q U \mid S \end{aligned}$$

The types $!^o T.S$ and $?^o T.S$ mean “send” and “receive”, respectively, and $\mathbf{end}^o$ means, well, session end.

The term language is the standard linear λ -calculus extended with concurrency primitives K :

$$\begin{aligned} L, M, N &:: x \mid K \mid \lambda x. M \mid M N \\ &\mid () \mid M; N \\ &\mid (M, N) \mid \mathbf{let} (x, y) = M \mathbf{in} N \\ &\mid \mathbf{absurd} M \\ &\mid \mathbf{inl} M \mid \mathbf{inr} M \mid \mathbf{case} L \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} \\ K &:: \mathbf{new} \mid \mathbf{fork} \mid \mathbf{send} \mid \mathbf{recv} \mid \mathbf{close} \end{aligned}$$

The concurrency primitives are uninterpreted in the term language. Rather, they are interpreted in a configuration language based on the π -calculus, which we omit from this paper (see Kokke and Dardha [35]).

We present the typing rules for PGV in fig. 1. A sequent $\Gamma \vdash_p^q M : T$ should be read as “ M is well-typed PGV program with type T in typing environment Γ , and when run it starts communicating at time p and stops at time q .”

Monadic Reflection. The graded monad $\mathbf{Sesh}_p^q$ arises from the *monadic reflection* [17] of the typing rules in fig. 1. Monadic reflection is a technique for translating programs in an effectful language to *monadic* programs in a pure language. For instance, Filinski [17] demonstrates the reflection from programs of type T in a language with exceptions and handlers to programs of type $T + \mathbf{exn}$ in a pure language where $\mathbf{exn}$ is the type of exceptions.

We translate programs from PGV to Haskell programs in the $\mathbf{Sesh}_p^q$ monad. First, let’s look at the translation of types:

$$\begin{aligned} \llbracket T \multimap_p^q U \rrbracket &= \llbracket T \rrbracket \multimap \mathbf{Sesh}_p^q \llbracket U \rrbracket \quad \llbracket \mathbf{1} \rrbracket = () \\ \llbracket !^o T.S \rrbracket &= \mathbf{Send}^o \llbracket T \rrbracket \llbracket S \rrbracket \quad \llbracket T \times U \rrbracket = (\llbracket T \rrbracket, \llbracket U \rrbracket) \\ \llbracket ?^o T.S \rrbracket &= \mathbf{Recv}^o \llbracket T \rrbracket \llbracket S \rrbracket \quad \llbracket \mathbf{0} \rrbracket = \mathbf{Void} \\ \llbracket \mathbf{end}^o \rrbracket &= \mathbf{End}^o \quad \llbracket T + U \rrbracket = \mathbf{Either} \llbracket T \rrbracket \llbracket U \rrbracket \end{aligned}$$

Now, let’s look at the translation of terms. A term of type T with lower bound p and upper bound q is translated to a Haskell program of type $\mathbf{Sesh}_p^q \llbracket T \rrbracket$:

$$\begin{aligned} \llbracket x \rrbracket &= \mathbf{ireturn} x \\ \llbracket \lambda x. L \rrbracket &= \mathbf{ireturn} (\lambda x. \llbracket L \rrbracket) \\ \llbracket K \rrbracket &= \mathbf{ireturn} \llbracket K \rrbracket \\ \llbracket L M \rrbracket &= \llbracket L \rrbracket \gg \lambda f. \llbracket M \rrbracket \gg \lambda x. f x \\ \llbracket () \rrbracket &= \mathbf{ireturn} () \\ \llbracket \mathbf{let} () = L \mathbf{in} M \rrbracket &= \llbracket L \rrbracket \gg \lambda (). M \\ \llbracket (L, M) \rrbracket &= \llbracket L \rrbracket \gg \lambda x. \llbracket M \rrbracket \gg \lambda y. \mathbf{ireturn} (x, y) \\ \llbracket \mathbf{let} (x, y) = L \mathbf{in} M \rrbracket &= \llbracket L \rrbracket \gg \lambda (x, y). \llbracket M \rrbracket \\ \llbracket \mathbf{absurd} L \rrbracket &= \llbracket L \rrbracket \gg \lambda x. \mathbf{absurd} x \\ \llbracket \mathbf{inl} L \rrbracket &= \llbracket L \rrbracket \gg \lambda x. \mathbf{ireturn} (\mathbf{Left} x) \\ \llbracket \mathbf{inr} L \rrbracket &= \llbracket L \rrbracket \gg \lambda x. \mathbf{ireturn} (\mathbf{Right} x) \\ \llbracket \mathbf{case} L \{ \mathbf{inl} x \mapsto M; \mathbf{inr} y \mapsto N \} \rrbracket &= \\ &\llbracket L \rrbracket \gg \lambda x. \mathbf{case} x \mathbf{of} \{ \mathbf{Left} x \rightarrow \llbracket M \rrbracket; \mathbf{Right} y \rightarrow \llbracket N \rrbracket \} \end{aligned}$$

We translate the communication primitives from PGV to those with the same name in priority-sesh, with some minor changes in the translations of **new** and **fork**, where PGV needs some unit arguments to create thunks in PGV, as it’s call-by-value, which aren’t needed in Haskell:

$$\begin{aligned} \llbracket \mathbf{new} : \mathbf{1} \multimap_{\perp}^{\perp} S \times \bar{S} \rrbracket &= \lambda (). \mathbf{new} :: () \multimap (\llbracket S \rrbracket, \llbracket (\mathbf{Dual} S) \rrbracket) \\ \llbracket \mathbf{fork} : (\mathbf{1} \multimap_p^q \mathbf{1}) \multimap_{\perp}^{\perp} \mathbf{1} \rrbracket &= \lambda k. \mathbf{fork} (k ()) :: () \multimap \mathbf{Sesh}_p^q () \multimap \mathbf{Sesh}_{\perp}^{\perp} () \end{aligned}$$

The rest of PGV’s communication primitives line up exactly with those of priority-sesh:

$$\begin{aligned} \llbracket \mathbf{send} : T \times !^o T.S \multimap_o^o S \rrbracket &= \mathbf{send} :: \mathbf{Session} \llbracket S \rrbracket \Rightarrow (\llbracket T \rrbracket, \mathbf{Send}^o \llbracket T \rrbracket \llbracket S \rrbracket) \multimap \mathbf{Sesh}_o^o \llbracket S \rrbracket \\ \llbracket \mathbf{recv} : ?^o T.S \multimap_o^o T \times S \rrbracket &= \mathbf{recv} :: \mathbf{Recv}^o \llbracket T \rrbracket \llbracket S \rrbracket \multimap \mathbf{Sesh}_o^o (\llbracket T \rrbracket, \llbracket S \rrbracket) \\ \llbracket \mathbf{close} : \mathbf{end}^o \multimap_o^o \mathbf{1} \rrbracket &= \mathbf{close} :: \mathbf{End}^o \multimap \mathbf{Sesh}_o^o () \\ \llbracket \mathbf{cancel} : S \multimap_{\perp}^{\perp} \mathbf{1} \rrbracket &= \mathbf{cancel} :: \mathbf{Session} \llbracket S \rrbracket \Rightarrow \llbracket S \rrbracket \multimap \mathbf{Sesh}_{\perp}^{\perp} () \end{aligned}$$

These two translations, on types and terms, comprise a *monadic reflection* from PGV into priority-sesh, which preserves typing. We state this theorem formally, using $\Gamma \vdash x :: a$ to mean that the Haskell program x has type a in typing environment Γ :

Theorem 3.1. *If $\Gamma \vdash_p^q M : T$, then $\llbracket \Gamma \rrbracket \vdash \llbracket M \rrbracket :: \mathbf{Sesh}_p^q \llbracket T \rrbracket$.*

Proof. Figure 2 presents the translation from typing derivations in PGV to abbreviated typing derivations in Haskell with priority-sesh. $\square$

Static Typing Rules.

$$\boxed{\Gamma \vdash_p^q M : T}$$

$\text{T-VAR} \quad \frac{}{x : T \vdash_{\top}^{\perp} x : T}$	$\text{T-LAM} \quad \frac{\Gamma, x : T \vdash_p^q M : U}{\Gamma \vdash_{\top}^{\perp} \lambda x.M : T \multimap_p^q U}$	$\text{T-CONST} \quad \frac{}{\emptyset \vdash_{\top}^{\perp} K : T}$	$\text{T-APP} \quad \frac{\Gamma \vdash_p^q M : T \multimap_{p''}^{q''} U \quad \Delta \vdash_{p'}^{q'} N : T \quad q < p' \quad q' < p''}{\Gamma, \Delta \vdash_{p \sqcap p' \sqcap p''}^{q \sqcup q' \sqcup q''} M N : U}$
$\text{T-UNIT} \quad \frac{}{\emptyset \vdash_{\top}^{\perp} () : 1}$	$\text{T-LETUNIT} \quad \frac{\Gamma \vdash_p^q M : 1 \quad \Delta \vdash_{p'}^{q'} N : T \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} \text{let } () = M \text{ in } N : T}$	$\text{T-PAIR} \quad \frac{\Gamma \vdash_p^q M : T \quad \Delta \vdash_{p'}^{q'} N : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} (M, N) : T \times U}$	
$\text{T-LETPAIR} \quad \frac{\Gamma \vdash_p^q M : T \times T' \quad \Delta, x : T, y : T' \vdash_{p'}^{q'} N : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} \text{let } (x, y) = M \text{ in } N : U}$	$\text{T-INL} \quad \frac{\Gamma \vdash_p^q M : T}{\Gamma \vdash_p^q \text{inl } M : T + U}$	$\text{T-INR} \quad \frac{\Gamma \vdash_p^q M : T}{\Gamma \vdash_p^q \text{inr } M : T + U}$	
$\text{T-CASESUM} \quad \frac{\Gamma \vdash_p^q L : T + T' \quad \Delta, x : T \vdash_{p'}^{q'} M : U \quad \Delta, y : T' \vdash_{p'}^{q'} N : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcup p'}^{q \sqcup q'} \text{case } L \{ \text{inl } x \mapsto M; \text{inr } y \mapsto N \} : U}$	$\text{T-ABSURD} \quad \frac{\Gamma \vdash_p^q M : 0}{\Gamma \vdash_p^q \text{absurd } M : T}$		

Type Schemas for Constants.

$$\boxed{K : T}$$

$$\text{new} : 1 \multimap_{\top}^{\perp} S \times \bar{S}$$

$$\text{fork} : (1 \multimap_p^q 1) \multimap_{\top}^{\perp} 1$$

$$\text{cancel} : S \multimap_{\top}^{\perp} 1$$

$$\text{send} : T \times !^o T.S \multimap_o^o S$$

$$\text{recv} : ?^o T.S \multimap_o^o T \times S$$

$$\text{close} : \text{end}^o \multimap_o^o 1$$

Figure 1. Typing rules for Priority GV.

4 Related Work

Session Types in Haskell. Orchard and Yoshida [49] discuss various approaches to implementing session types in Haskell. Their overview is reproduced below:

- Neubauer and Thiemann [44] give an encoding of first-order single-channel session-types with recursion;
- Using *parameterised monads*, Pucella and Tov [56] provide multiple channels, recursion, and some building blocks for delegation, but require manual manipulation of a session typing environment;
- Sackman and Eisenbach [57] provide an alternate approach where session types are constructed via a value-level witnesses;
- Imai et al. [30] extend Pucella and Tov [56] with delegation and a more user-friendly approach to handling multiple channels;
- Orchard and Yoshida [50] use an embedding of effect systems into Haskell via graded monads based on a formal encoding of session-typed π -calculus into PCF with an effect system;
- Lindley and Morris [41] provide a *finally tagless* embedding of the GV session-typed functional calculus into Haskell, building on a linear λ -calculus embedding due to Polakow [55].

With respect to linearity, all works above—except Neubauer and Thiemann [44]—guarantee linearity by encoding a linear typing environment in the Haskell type system, which leads to a trade-off between having easy-to-write session types and having idiomatic programs. We side-step this trade-off by relying on Linear Haskell to check linearity. Furthermore, our implementation supports all relevant features, including multiple channels, full delegation, recursion, and more idiomatic code.

With respect to deadlock freedom, none of the works above—except Lindley and Morris [41]—guarantee deadlock freedom. However, Lindley and Morris [41] guarantee deadlock freedom *structurally*, by implementing GV. As discussed in section 1, structure-based deadlock freedom is more restrictive than priority-based deadlock freedom, as it restricts communication graphs to *trees*, whereas the priority-based approach allows programs to have *cyclic* process structures.

Orchard and Yoshida [49] summarise the capabilities of the various implementations of session types in Haskell in a table, which we adapted in table 1 by adding columns for the various versions of priority-sesh. In general, you may read $\checkmark$ as “Kinda” and $\checkmark$ as a resounding “Yes!” For instance, Pucella and Tov [56] only provide *partial* delegation, Neubauer and Thiemann [44], Pucella and Tov [56], and Lindley and Morris [41] still need to use combinators instead of standard Haskell application, abstraction, or variables in

$$\begin{array}{c}
\frac{}{x :: \llbracket T \rrbracket \vdash x :: \llbracket T \rrbracket} \\
\frac{}{x : T \vdash_{\frac{1}{2}} x : T} = \frac{x :: \llbracket T \rrbracket \vdash x :: \llbracket T \rrbracket}{\text{ireturn } x :: \text{Sesh}_{\frac{1}{2}}^{\perp} \llbracket T \rrbracket} \\
\\
\frac{\Gamma, x : T \vdash_p^q L : U}{\Gamma \vdash_{\frac{1}{2}} \lambda x. L : T \multimap_p^q U} = \frac{\llbracket \Gamma \rrbracket, x :: \llbracket T \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q \llbracket U \rrbracket}{\text{ireturn } (\lambda x. \llbracket L \rrbracket) :: \text{Sesh}_{\frac{1}{2}}^{\perp} (\llbracket T \rrbracket \multimap \text{Sesh}_p^q \llbracket U \rrbracket)} \\
\\
\frac{}{\emptyset \vdash_{\frac{1}{2}} K : T} = \frac{}{\text{ireturn } \llbracket K \rrbracket :: \text{Sesh}_{\frac{1}{2}}^{\perp} \llbracket T \rrbracket} \\
\\
\frac{\Gamma \vdash_p^q L : T \multimap_{p''}^{q''} U \quad \Delta \vdash_{p'}^{q'} M : T \quad q < p' \quad q' < p''}{\Gamma, \Delta \vdash_{p \sqcap p' \sqcap p''}^{q \sqcup q' \sqcup q''} L M : U} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q (\llbracket T \rrbracket \multimap \text{Sesh}_{p''}^{q''} \llbracket U \rrbracket) \quad \llbracket \Delta \rrbracket \vdash \llbracket M \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket T \rrbracket}{\llbracket L \rrbracket \gg \lambda f. \llbracket M \rrbracket \gg \lambda x. f \ x :: (q < p', q' < p'') \Rightarrow \text{Sesh}_{p \sqcap p' \sqcap p''}^{q \sqcup q' \sqcup q''} \llbracket U \rrbracket} \\
\\
\frac{}{\emptyset \vdash_{\frac{1}{2}} () : \mathbf{1}} = \frac{}{\text{ireturn } () :: \text{Sesh}_{\frac{1}{2}}^{\perp} ()} \\
\\
\frac{\Gamma \vdash_p^q L : \mathbf{1} \quad \Delta \vdash_{p'}^{q'} M : T \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} \text{let } () = L \text{ in } M : T} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q () \quad \llbracket \Delta \rrbracket \vdash \llbracket M \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket T \rrbracket}{\llbracket L \rrbracket \gg \lambda (). M :: (p < q') \Rightarrow \text{Sesh}_{p \sqcap p'}^{q \sqcup q'} \llbracket T \rrbracket} \\
\\
\frac{\Gamma \vdash_p^q L : T \quad \Delta \vdash_{p'}^{q'} M : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} (L, M) : T \times U} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q \llbracket T \rrbracket \quad \llbracket \Delta \rrbracket \vdash \llbracket M \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket U \rrbracket}{\llbracket L \rrbracket \gg \lambda x. \llbracket M \rrbracket \gg \lambda y. \text{ireturn } (x, y) :: (q < p') \Rightarrow \text{Sesh}_{p \sqcap p'}^{q \sqcup q'} (\llbracket T \rrbracket, \llbracket U \rrbracket)} \\
\\
\frac{\Gamma \vdash_p^q L : T \times T' \quad \Delta, x : T, y : T' \vdash_{p'}^{q'} M : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} \text{let } (x, y) = L \text{ in } M : U} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q (\llbracket T \rrbracket, \llbracket T' \rrbracket) \quad \llbracket \Delta \rrbracket, x :: \llbracket T \rrbracket, y :: \llbracket T' \rrbracket \vdash \llbracket M \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket U \rrbracket}{\llbracket L \rrbracket \gg \lambda (x, y). \llbracket M \rrbracket :: (q < p') \Rightarrow \text{Sesh}_{p \sqcap p'}^{q \sqcup q'} \llbracket U \rrbracket} \\
\\
\frac{\Gamma \vdash_p^q L : T}{\Gamma \vdash_p^q \text{inl } L : T + U} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q \llbracket T \rrbracket}{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket \gg \lambda x. \text{ireturn } (\text{Left } x) :: \text{Sesh}_p^q (\text{Either } \llbracket T \rrbracket \llbracket U \rrbracket)} \\
\\
\frac{\Gamma \vdash_p^q L : T}{\Gamma \vdash_p^q \text{inr } L : T + U} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q \llbracket T \rrbracket}{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket \gg \lambda x. \text{ireturn } (\text{Right } x) :: \text{Sesh}_p^q (\text{Either } \llbracket T \rrbracket \llbracket U \rrbracket)} \\
\\
\frac{\Gamma \vdash_p^q L : T + T' \quad \Delta, x : T \vdash_{p'}^{q'} M : U \quad \Delta, y : T' \vdash_{p'}^{q'} N : U \quad q < p'}{\Gamma, \Delta \vdash_{p \sqcap p'}^{q \sqcup q'} \text{case } L \text{ \{ inl } } x \mapsto M; \text{ inr } y \mapsto N \} : U} = \\
\frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q (\text{Either } \llbracket T \rrbracket \llbracket T' \rrbracket) \quad \llbracket \Delta \rrbracket \vdash x :: \llbracket T \rrbracket \vdash \llbracket M \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket U \rrbracket \quad \llbracket \Delta \rrbracket \vdash y :: \llbracket T' \rrbracket \vdash \llbracket N \rrbracket :: \text{Sesh}_{p'}^{q'} \llbracket U \rrbracket}{\llbracket \Gamma \rrbracket, \llbracket \Delta \rrbracket \vdash \llbracket L \rrbracket \gg \lambda x. \text{case } x \text{ \{ Left } } x \mapsto \llbracket M \rrbracket; \text{ Right } y \mapsto \llbracket N \rrbracket \} :: \text{Sesh}_{p \sqcap p'}^{q \sqcup q'} \llbracket U \rrbracket} \\
\\
\frac{\Gamma \vdash_p^q L : \mathbf{0}}{\Gamma \vdash_p^q \text{absurd } L : T} = \frac{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket :: \text{Sesh}_p^q \text{Void}}{\llbracket \Gamma \rrbracket \vdash \llbracket L \rrbracket \gg \lambda x. \text{absurd } x :: \text{Sesh}_p^q \llbracket T \rrbracket}
\end{array}$$

Figure 2. Translation from Priority GV to Sesh preserves types.

some places, and Neubauer and Thiemann [44] is only deadlock free on the technicality that they don't support multiple channels.

Session Types in other Programming Languages. Session types have been integrated in other programming language paradigms. Jespersen et al. [31], Padovani [52], Scalas and Yoshida [60] integrate *binary* session types in the *native* host language, without language extensions; this to

Table 1. Capabilities of various implementations of session types in Haskell [adapted from 49].

	NT04	PT08	SE08	IYA10	OY16	LM16	priority-sesh		
							section 2.2	section 2.3	section 2.4
Recursion	✓	✓	✓	✓	✓		✓	✓	
Delegation		✓	✓	✓	✓	✓	✓	✓	✓
Multiple channels		✓	✓	✓	✓	✓	✓	✓	✓
Idiomatic code	✓	✓		✓	✓	✓	✓	✓	✓
Easy-to-write session types	✓	✓	✓		✓	✓	✓	✓	✓
Deadlock freedom	✓					✓		✓	✓
via <i>process structure</i>	✓					✓		✓	
via <i>priorities</i>									✓

avoid hindering session types use in practice. To obtain this integration of session types without extensions Padovani [52], Scalas and Yoshida [60]) combine *static* typing of input and output actions with *runtime* checking of linearity of channel usage.

Implementations of *multiparty* session types (MPST) are less common than binary implementations. Scalas et al. [59] integrate MPST in Scala building upon Scalas and Yoshida [60] and a continuation-passing style encoding of session types into linear types Dardha et al. [11]. There are several works on implementations of MPST in Java: Sivaramakrishnan et al. [61] implement MPST leveraging an extension of Java with session primitives; Hu and Yoshida [29] develops a MPST-based API generation for Java leveraging CFSMs by Brand and Zafropulo [7]; and Kouzapas et al. [37] implement session types in the form of *typestates* in Java. Demangeon et al. [13] implement MPST in Python and Fowler [18], Neykova and Yoshida [46] in Erlang, focusing on purely dynamic MPST verification via runtime monitoring. Neykova et al. [45], Neykova and Yoshida [47] extend the work by Demangeon et al. [13] with actors and timed specifications. Lopez et al. [43] adopt a dependently-typed MPST theory to verify MPI programs.

Session Types, Linear Logic and Deadlock Freedom.

The main line of work regarding deadlock freedom in session-typed systems is that of Curry-Howard correspondences with linear logic [25]. Caires and Pfenning [8] defined a correspondence between session types and dual intuitionistic linear logic and Wadler [64] between session types and classical linear logic. These works guarantee deadlock freedom *by design* as the communication structures are restricted to trees and due to the *cut* rule, processes share *only* one channel between them. Dardha and Gay [10] extend Wadler [64] with *priorities* following Kobayashi [32], Padovani [51], thus allowing processes to share more than one channel in parallel, while guaranteeing deadlock freedom. Balzer et al. [2] introduce *sharing* and guarantee deadlock freedom via priorities. All the above works deal with deadlock freedom in a session-typed π -calculus. With regards to function languages, the original works on GV [23, 24] did not guarantee

deadlock freedom. This was later addressed by Lindley and Morris [40], Wadler [65] via syntactic restrictions where communication once again follows a tree structure. Kokke and Dardha [35] introduce PGV–Priority GV, by following Dardha and Gay [10] and allowing for more flexible programming in GV. Fowler et al. [19] present Hypersequent GV (HGV), a core calculus for functional programming with session types that enjoys deadlock freedom, confluence, and strong normalisation.

Other works on deadlock freedom in session-typed systems include the works by Dezani-Ciancaglini et al. [15], where deadlock freedom is guaranteed by allowing only one active session at a time and by Dezani-Ciancaglini et al. [14], where priorities are used for correct interleaving of channels. Honda et al. [28] guarantee deadlock freedom *within a single* session of MPST, but not for session interleaving. Kokke [34] guarantees deadlock freedom of session types in Rust by enforcing a tree structure of communication actions.

5 Discussion and Future Work

We presented *priority-sesh*, an implementation of deadlock-free session types in Linear Haskell. Using Linear Haskell allows us to check linearity—or more accurately, have linearity guaranteed for us—without relying on complex type-level machinery. Consequently, we have easy-to-write session types and idiomatic code—in fact, probably *the most* idiomatic code when compared with previous work, though in fairness, all previous work predates Linear Haskell. Unfortunately, there are some drawbacks to using Linear Haskell. Most importantly, Linear Haskell is not very mature at this stage. For instance:

- Anonymous functions are assumed to be unrestricted rather than linear, meaning anonymous functions must be factored out into a let-binding or where-clause with *at least* a minimal type signature such as $_ \multimap _$.
- There is no integration with base or popular Haskell packages, and given that `LinearTypes` is an extension, there likely won't be for quite a while. There's `linear-base`, which provides linear variants of many

of the constructs in base. However, linear-base relies heavily on `unsafeCoerce`, which, ironically, may affect Haskell's performance.

- Generally, there is little integration with the Haskell ecosystem, *e.g.*, one other contribution we made are the formatting directives for Linear Haskell in `lhs2TeX` [1].

However, we believe that many of these drawbacks will disappear as the Linear Haskell ecosystem matures.

Our work also provides a library which guarantees deadlock freedom via *priorities*, which allows for more flexible typing than previous work on deadlock freedom via a tree process structure.

In the future, we plan to address the issue of priority-polymorphic code and recursion session types in our implementation. (While the versions of our library in sections 2.2 and 2.3 support recursion, that is not yet the case for the priority-based version in section 2.4.) This is a challenging task, as it requires complex reasoning about type-level naturals. We outlined various approaches in section 2.4. However, an alternative we would like to investigate, would be to implement priority-sesh in Idris2 [5, 6], which supports both linear types and complex type-level reasoning.

Acknowledgments

We thank Simon Fowler and April Gonçalves for comments on the manuscript. This work is supported by the EU HORIZON 2020 MSCA RISE project 778233 “Behavioural Application Program Interfaces” (BehAPI).

References

- [1] Accessed: 2021-08-06. `lhs2tex`: Preprocessor for typesetting Haskell sources with LaTeX. <https://hackage.haskell.org/package/lhs2tex>.
- [2] Stephanie Balzer, Bernardo Toninho, and Frank Pfenning. 2019. Manifest Deadlock-Freedom for Shared Session Types. In *Proc. of ESOP (Lecture Notes in Computer Science, Vol. 11423)*. Springer, 611–639.
- [3] Giovanni Bernardi, Ornela Dardha, Simon J. Gay, and Dimitrios Kouzapas. 2014. On Duality Relations for Session Types. In *Trustworthy Global Computing*. Springer Berlin Heidelberg, 51–66. https://doi.org/10.1007/978-3-662-45917-1_4
- [4] Jean-Philippe Bernardy, Mathieu Boespflug, Ryan R. Newton, Simon Peyton Jones, and Arnaud Spiwack. 2018. Linear Haskell: practical linearity in a higher-order polymorphic language. *Proc. of POPL 2* (2018), 1–29. <https://doi.org/10.1145/3158093>
- [5] Edwin Brady. 2013. Idris, a general-purpose dependently typed programming language: Design and implementation. *Journal of Functional Programming* 23, 5 (2013), 552–593. <https://doi.org/10.1017/S095679681300018X>
- [6] Edwin Brady. 2017. Type-Driven Development of Concurrent Communicating Systems. *Computer Science* 18, 3 (2017), 219. <https://doi.org/10.7494/csci.2017.18.3.1413>
- [7] Daniel Brand and Pitro Zafiropolo. 1983. On Communicating Finite-State Machines. *J. ACM* 30, 2 (April 1983), 323–342. <https://doi.org/10.1145/322374.322380>
- [8] Luís Caires and Frank Pfenning. 2010. Session Types as Intuitionistic Linear Propositions. In *CONCUR (LNCS, Vol. 6269)*. Springer, 222–236. https://doi.org/10.1007/978-3-642-15375-4_16
- [9] Ornela Dardha. 2016. *Type Systems for Distributed Programs: Components and Sessions*. Atlantis Studies in Computing, Vol. 7. Springer / Atlantis Press. <https://doi.org/10.2991/978-94-6239-204-5>
- [10] Ornela Dardha and Simon J. Gay. 2018. A New Linear Logic for Deadlock-Free Session-Typed Processes. In *Proc. of FoSSaCS (LNCS, Vol. 10803)*. Springer, 91–109.
- [11] Ornela Dardha, Elena Giachino, and Davide Sangiorgi. 2012. Session types revisited. In *Proc. of PPDP*. ACM, 139–150.
- [12] Ornela Dardha, Elena Giachino, and Davide Sangiorgi. 2017. Session types revisited. *Inf. Comput.* 256 (2017), 253–286. <https://doi.org/10.1016/j.ic.2017.06.002> Extended version of [11].
- [13] Romain Demangeon, Kohei Honda, Raymond Hu, Rumyana Neykova, and Nobuko Yoshida. 2015. Practical Interruptible Conversations: Distributed Dynamic Verification with Multiparty Session Types and Python. *Formal Methods in System Design* (2015). <https://doi.org/10.1007/s10703-014-0218-8>
- [14] Mariangiola Dezani-Ciancaglini, Ugo de'Liguoro, and Nobuko Yoshida. 2009. On Progress for Structured Communications. In *Proc. of TGC (LNCS, Vol. 4912)*. Springer, 257–275.
- [15] Mariangiola Dezani-Ciancaglini, Dimitris Mostrous, Nobuko Yoshida, and Sophia Drossopoulou. 2006. Session Types for Object-Oriented Languages. In *Proc. of ECOOP (LNCS, Vol. 4067)*. Springer, 328–352.
- [16] Richard A. Eisenberg, Dimitrios Vytiniotis, Simon Peyton Jones, and Stephanie Weirich. 2014. Closed type families with overlapping equations. In *Proceedings of the 41st ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*. ACM. <https://doi.org/10.1145/2535838.2535856>
- [17] Andrzej Filinski. 1994. Representing monads. In *Proceedings of the 21st ACM SIGPLAN-SIGACT symposium on Principles of programming languages - POPL '94*. ACM Press. <https://doi.org/10.1145/174675.178047>
- [18] Simon Fowler. 2016. An Erlang Implementation of Multiparty Session Actors. In *ICE*. <https://doi.org/10.4204/EPTCS.223.3>
- [19] Simon Fowler, Wen Kokke, Ornela Dardha, Sam Lindley, and J. Garrett Morris. 2021. Separating Sessions Smoothly. *CoRR* abs/2105.08996 (2021). [arXiv:2105.08996](https://arxiv.org/abs/2105.08996) <https://arxiv.org/abs/2105.08996>
- [20] Simon Fowler, Sam Lindley, J. Garrett Morris, and Sára Decova. 2019. Exceptional Asynchronous Session Types: Session Types without Tiers. *Proc. of POPL 3*, Article 28 (2019), 29 pages. <https://doi.org/10.1145/3290341>
- [21] Marco Gaboardi, Shin ya Katsumata, Dominic Orchard, Flavien Breuvar, and Tarmo Uustalu. 2016. Combining effects and coefficients via grading. In *Proceedings of the 21st ACM SIGPLAN International Conference on Functional Programming*. ACM. <https://doi.org/10.1145/2951913.2951939>
- [22] Simon J. Gay, Peter Thiemann, and Vasco T. Vasconcelos. 2020. Duality of Session Types: The Final Cut. *Electronic Proceedings in Theoretical Computer Science* 314 (April 2020), 23–33. <https://doi.org/10.4204/eptcs.314.3>
- [23] Simon J. Gay and Vasco T. Vasconcelos. 2010. Linear type theory for asynchronous session types. *Journal of Functional Programming* 20, 1 (2010), 19–50.
- [24] Simon J. Gay and Vasco T. Vasconcelos. 2012. Linear type theory for asynchronous session types. *JFP* 20, 1 (2012), 19–50. Extended version of [23].
- [25] Jean-Yves Girard. 1987. Linear Logic. *Theoretical Computer Science* 50 (1987), 1–102.
- [26] Kohei Honda. 1993. Types for Dyadic Interaction. In *Proc. of CONCUR (LNCS, Vol. 715)*. Springer, 509–523.
- [27] Kohei Honda, Vasco Thudichum Vasconcelos, and Makoto Kubo. 1998. Language Primitives and Type Discipline for Structured Communication-Based Programming. In *Proc. of ESOP (LNCS, Vol. 1381)*. Springer, 122–138.
- [28] Kohei Honda, Nobuko Yoshida, and Marco Carbone. 2008. Multiparty asynchronous session types. In *Proc. of POPL*, Vol. 43(1). ACM, 273–284.

- [29] Raymond Hu and Nobuko Yoshida. 2016. Hybrid Session Verification Through Endpoint API Generation. In *Proc. of FASE*. https://doi.org/10.1007/978-3-662-49665-7_24
- [30] Keigo Imai, Shoji Yuen, and Kiyoshi Agusa. 2010. Session Type Inference in Haskell. In *Proc. of PLACES (EPTCS, Vol. 69)*. 74–91. <https://doi.org/10.4204/EPTCS.69.6>
- [31] Thomas Bracht Laumann Jespersen, Philip Munksgaard, and Ken Friis Larsen. 2015. Session types for Rust. In *Proc. of WGP@ICFP*. <https://doi.org/10.1145/2808098.2808100>
- [32] Naoki Kobayashi. 2006. A New Type System for Deadlock-Free Processes. In *Proc. of CONCUR (LNCS, Vol. 4137)*. Springer, 233–247.
- [33] Naoki Kobayashi, Benjamin C. Pierce, and David N. Turner. 1999. Linearity and the pi-calculus. *ACM Trans. Program. Lang. Syst.* 21, 5 (1999), 914–947. <https://doi.org/10.1145/330249.330251>
- [34] Wen Kokke. 2019. Rusty Variation: Deadlock-free Sessions with Failure in Rust. *EPTCS* 304 (Sept. 2019), 48–60. <https://doi.org/10.4204/eptcs.304.4> Renamed to Sesh.
- [35] Wen Kokke and Ornela Dardha. 2021. Prioritise the Best Variation. In *Proc. of FORTE (Lect. Not. in Comput. Sci., Vol. 12719)*. Springer, 100–119. https://doi.org/10.1007/978-3-030-78089-0_6
- [36] Wen Kokke and Ornela Dardha. 2021. Prioritise the Best Variation. *CoRR* abs/2103.14466 (2021). arXiv:2103.14466 <https://arxiv.org/abs/2103.14466> Extended version of [35].
- [37] Dimitrios Kouzapas, Ornela Dardha, Roly Perera, and Simon J. Gay. 2016. Typechecking protocols with Mungo and StMungo. In *PPDP*. 146–159.
- [38] John Launchbury and Simon L. Peyton Jones. 1994. Lazy Functional State Threads. In *Proc. of PLDI (Orlando, Florida, USA)*. ACM, New York, NY, USA, 24–35. <https://doi.org/10.1145/178243.178246>
- [39] Sam Lindley and Conor McBride. 2013. Hasochism. In *Proceedings of the 2013 ACM SIGPLAN symposium on Haskell - Haskell '13*. ACM Press. <https://doi.org/10.1145/2503778.2503786>
- [40] Sam Lindley and J. Garrett Morris. 2015. A Semantics for Propositions as Sessions. In *Proc. of ESOP*. 560–584.
- [41] Sam Lindley and J. Garrett Morris. 2016. Embedding session types in Haskell. In *Proc. of Haskell*. ACM, 133–145. <https://doi.org/10.1145/2976002.2976018>
- [42] Sam Lindley and J. Garrett Morris. 2017. Lightweight Functional Session Types. In *Behavioural Types: from Theory to Tools*. River Publishers, 265–286.
- [43] Hugo A. Lopez, Eduardo R. B. Marques, Francisco Martins, Nicholas Ng, Casar Santos, Vasco Thudichum Vasconcelos, and Nobuko Yoshida. 2015. Protocol-Based Verification of Message-Passing Parallel Programs. In *OOPSLA*. <https://doi.org/10.1145/2814270.2814302>
- [44] Matthias Neubauer and Peter Thiemann. 2004. An Implementation of Session Types. In *Proc. of PADL (Lecture Notes in Computer Science, Vol. 3057)*. Springer, 56–70. https://doi.org/10.1007/978-3-540-24836-1_5
- [45] Romyana Neykova, Laura Bocchi, and Nobuko Yoshida. 2017. Timed Runtime Monitoring for Multiparty Conversations. *Formal Aspects of Computing* (2017). <https://doi.org/10.1007/s00165-017-0420-8>
- [46] Romyana Neykova and Nobuko Yoshida. 2017. Let It Recover: Multiparty Protocol-Induced Recovery. In *CC*. <https://doi.org/10.1145/3033019.3033031>
- [47] Romyana Neykova and Nobuko Yoshida. 2017. Multiparty Session Actors. *Logical Methods in Computer Science* 13, 1 (March 2017). [https://doi.org/10.23638/LMCS-13\(1:17\)2017](https://doi.org/10.23638/LMCS-13(1:17)2017)
- [48] Dominic Orchard, Philip Wadler, and Harley Eades. 2020. Unifying graded and parameterised monads. *Electronic Proceedings in Theoretical Computer Science* 317 (May 2020), 18–38. <https://doi.org/10.4204/eptcs.317.2>
- [49] Dominic Orchard and Nobuko Yoshida. 2017. Session types with linearity in Haskell. *Behavioural Types: from Theory to Tools* (2017), 219.
- [50] Dominic A. Orchard and Nobuko Yoshida. 2016. Effects as sessions, sessions as effects. In *Proc. of POPL*. ACM, 568–581. <https://doi.org/10.1145/2837614.2837634>
- [51] Luca Padovani. 2014. Deadlock and Lock Freedom in the Linear π -Calculus. In *Proc. of CSL-LICS*. ACM, 72:1–72:10.
- [52] Luca Padovani. 2017. A simple library implementation of binary sessions. *Journal of Functional Programming* 27 (2017). <https://doi.org/10.1017/S0956796816000289> Website: <http://www.di.unito.it/~padovani/Software/FuSe/FuSe.html>.
- [53] Luca Padovani and Luca Novara. 2015. Types for Deadlock-Free Higher-Order Programs. In *Proc. of FORTE (LNCS, Vol. 9039)*. Springer, 3–18.
- [54] Simon L. Peyton Jones, Andrew D. Gordon, and Sigbjørn Finne. 1996. Concurrent Haskell. In *Proc. of POPL*. ACM, 295–308. <https://doi.org/10.1145/237721.237794>
- [55] Jeff Polakow. 2015. Embedding a full linear Lambda calculus in Haskell. In *Proc. of the Symposium on Haskell*. ACM. <https://doi.org/10.1145/2804302.2804309>
- [56] Riccardo Pucella and Jesse A. Tov. 2008. Haskell session types with (almost) no class. In *Proc. of Haskell*. ACM. <https://doi.org/10.1145/1411286.1411290>
- [57] Matthew Sackman and Susan Eisenbach. 2008. Session Types in Haskell Updating Message Passing for the 21st Century. (01 2008).
- [58] Davide Sangiorgi and David Walker. 2001. *The π -calculus: a Theory of Mobile Processes*. Cambridge University Press.
- [59] Alceste Scalas, Ornela Dardha, Raymond Hu, and Nobuko Yoshida. 2017. A Linear Decomposition of Multiparty Sessions for Safe Distributed Programming. In *Proc. of ECOOP (LIPIcs, Vol. 74)*. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 24:1–24:31. <https://doi.org/10.4230/LIPIcs.ECOOP.2017.24>
- [60] Alceste Scalas and Nobuko Yoshida. 2016. Lightweight Session Programming in Scala. In *ECOOP*. <https://doi.org/10.4230/LIPIcs.ECOOP.2016.21>
- [61] K. C. Sivaramakrishnan, Karthik Nagaraj, Lukasz Ziarek, and Patrick Eugster. 2010. Efficient Session Type Guided Distributed Interaction. In *Proc. of COORDINATION*, Vol. 6116. https://doi.org/10.1007/978-3-642-13414-2_11
- [62] Kaku Takeuchi, Kohei Honda, and Makoto Kubo. 1994. An Interaction-Based Language and its Typing System. In *Proc. of PARLE (LNCS, Vol. 817)*. Springer, 398–413.
- [63] Philip Wadler. 2012. Propositions as sessions. In *Proc. of ICFP*. 273–286.
- [64] Philip Wadler. 2014. Propositions as sessions. *Journal of Functional Programming* 24, 2-3 (Jan. 2014), 384–418. Extended version of [63].
- [65] Philip Wadler. 2015. Propositions as types. *Commun. ACM* 58, 12 (2015), 75–84.

6.3 Conclusion

In this chapter, we introduced a library in Linear Haskell which implements session-typed channels based on GV and Priority GV. We defined the monadic reflection of PGV's type system into a graded linear monad, to permit an easy embedding of PGV's priority typing into Linear Haskell's type system. Finally, we compared our Haskell library to existing Haskell libraries for session-typed channels.

In the future, it would be interesting to implement a variant of the library that uses type-level programming for priority inference, as the current design requires the user to manually annotate actions with their priorities. Furthermore, it would be interesting to describe the relation between the library and GV more formally, and formalise the correspondence using a proof-assistant. Finally, it would be useful to extend the library with the link primitive, both from a practical standpoint, and to more closely align the implementation with the formal calculus.

Appendix A

Glossary

A.1 Graph Theory

In this section, I introduce various graph theoretic notions that are used throughout the thesis. I need undirected labelled graphs when discussing connection graphs in § 2.2.5 and § 3.2.5, mixed graphs when discussing dependency graphs and priority graphs in § 2.2.4, § 3.2.4, and § 5.3.2, and undirected multigraphs when discussing abstract process structure in I.3.1.

Commonly, graphs are introduced as tuples of their components. However, since I require undirected, directed, and mixed graphs, I intend to save some ink by defining graphs in terms of their projections.

In the following, if some projection is undefined, it is the empty set, e.g. for an undirected graph, the set of directed arcs is empty, and, consequently, directed reachability is the empty relation.

- A *graph* (ranged over by G) has a set of vertices, denoted by V_G (ranged over by u, v).
- In an abuse of notation, I write u for the singleton graph consisting of the single vertex u , i.e. $V_u \triangleq \{u\}$.
- Two graphs G_1 and G_2 are *disjoint* if and only if their sets of vertices are disjoint, i.e. $V_{G_1} \cap V_{G_2} = \emptyset$.
- An *edge* (ranged over by e) is an unordered pair of vertices, denoted by juxtaposition, i.e. $uv \triangleq \{u, v\}$.
- An *edge-loop* or *loop* is an edge uu that connects a vertex to itself.
- In an abuse of notation, I write uv for the graph consisting of the single edge uv , i.e. $V_{uv} \triangleq \{u, v\}$ and $E_{uv} \triangleq \{uv\}$.
- An *undirected graph* G is a graph with a set of undirected edges with no edge-loops, denoted by E_G , i.e. $E_G \subseteq \{uv \mid u, v \in V_G \wedge u \neq v\}$.

- An *arc* (ranged over by a) is an ordered pair of vertices, denoted by juxtaposition overset with an arrow to indicate the direction, i.e. $\vec{uv} \triangleq (u, v)$.
- An *arc-loop* or *loop* is an arc $\vec{uu}$ that connects a vertex to itself.
- In an abuse of notation, I write $\vec{uv}$ for the graph consisting of the single arc $\vec{uv}$, i.e. $V_{\vec{uv}} \triangleq \{u, v\}$ and $A_{\vec{uv}} \triangleq \{\vec{uv}\}$.
- A *directed graph* G is a graph with a set of directed arcs with no arc-loops, denoted by A_G , i.e. $A_G \subseteq \{\vec{uv} \mid u, v \in V_G \wedge u \neq v\}$.
- A *mixed graph* G is a graph with a set of undirected edges with no edge-loops, as above, and a set of directed arcs with no arc-loops, as above.
- An *edge-labelled graph* G is a graph with a set of edges, as above, a set of edge labels, denoted by $\mathcal{L}_G$, and an edge-labelling function, denoted by ℓ_G , where $\ell_G : E_G \rightarrow \mathcal{L}_G$. The definitions for E_G and $\mathcal{L}_G$ may be omitted, since $E_G \triangleq \text{dom}(\ell_G)$ and $\mathcal{L}_G \triangleq \text{cod}(\ell_G)$.
- An *undirected multigraph* G is a graph with a set of edge names, denoted by $\mathcal{E}_G$, and an edge-connection function, denoted by r_G , where $r_G : \mathcal{E}_G \rightarrow \{uv \mid u, v \in V_G\}$. The set of edges E_G is defined as the union of all edges, i.e. $E_G \triangleq \bigcup_{e \in \mathcal{E}_G} r_G(e)$. An undirected multigraph is similar to an undirected edge-labelled graph but differs in that it permits edge-loops and permits multiple edges between any two vertices.
- The *empty graph* with vertices V , written $\vec{K}_V$, is the graph consisting of vertices V with no edges or arcs, i.e. $V_{\vec{K}_V} \triangleq V$, $E_{\vec{K}_V} \triangleq \emptyset$, and $A_{\vec{K}_V} \triangleq \emptyset$.
- The *graph union* of G_1 and G_2 , denoted by $G_1 \cup G_2$, is defined by, for each projection, taking the union of the projection of G_1 and G_2 , e.g. $V_{G_1 \cup G_2} \triangleq V_{G_1} \cup V_{G_2}$, and $E_{G_1 \cup G_2} \triangleq E_{G_1} \cup E_{G_2}$, etc..
- The *directed rooting* of G in v , written $v < G$, is the graph formed by adding the vertex v and adding an arc from v to every other vertex of G , i.e. $V_{v < G} \triangleq \{v\} \cup V_G$ and $A_{v < G} \triangleq \{\vec{vu} \mid u \in V_G\} \cup A_G$. Directed rooting preserves the remaining projections, e.g. $E_{v < G} \triangleq E_G$.
- For any graph G with vertices $u, v \in V_G$, u is *adjacent* to v if and only if there exists some edge $uv \in E_G$ or some arc $\vec{uv} \in A_G$.
- For any graph G , a *walk* (ranged over by w) is a sequence of pairwise adjacent vertices. Equivalently, a walk is a sequence of edges and arcs that join a sequence of vertices, where all arcs have the same direction.
- A walk w *visits* a vertex v if and only if v occurs in the walk, i.e. if and only if $w_i = v$ for some $i \in \mathbb{N}$.
- A walk is *closed* if and only if its first and last vertex are the same.

- A walk is *undirected* if and only if it contains only edges.
- A walk is *directed* if and only if it contains only arcs.
- A walk is *essentially directed* if and only if it contains at least one arc.
- A *path* (ranged over by p) is a walk without repeated vertices.
- A *cycle* (ranged over by c) is a closed path.
- A graph is *acyclic* if and only if the undirected graph formed by replacing all arcs in with edges contains no cycles.
- A graph is *essentially acyclic* if and only if it contains no essentially directed cycles.
- A graph is *strongly connected* if and only if there exists a path between any two distinct vertices.
- A graph is *connected* if and only if the undirected graph formed by replacing all arcs in with edges is strongly connected. (Commonly, connected is also referred to as *weakly connected*.)
- For any graph G , the *subgraph* of G induced by U , denoted by $G[U]$, where $U \subseteq V_G$, is the graph constructed by taking the subset of vertices U and restricting G 's projections to vertices in U .
- For any graph G , a *strongly connected component* of G (ranged over by C) is a maximal strongly connected subgraph of G .
- For any graph G , a *component* of G (ranged over by C) is a maximal connected subgraph of G .
- A *tree* (ranged over by T) is a graph that is connected and acyclic.
- A *forest* (ranged over by F) is a graph whose components are trees.
- For any graph G , *undirected reachability*, denoted by $\sim_G$, is the equivalence closure over E_G , i.e. $u \sim_G v$ holds if $u = v$ or there exists an undirected path from u to v . (If E_G is undefined, undirected reachability is the smallest reflexive relation over V_G .)
- For any graph G , *essentially directed reachability*, denoted by $<_G$, is the transitive closure over A_G quotiented by $\sim_G$, i.e. $u <_G v$ holds if and only if there exists an essentially directed path from u to v . (If A_G is undefined, essentially directed reachability is the empty relation.)
- For any graph G , *reachability*, denoted by $\leqslant_G$, is the union of undirected and essentially directed reachability, i.e. $u \leqslant_G v$ holds if and only if there exists an undirected or essentially directed path from u to v .
- For any graph G , its *sources* and *sinks*, written $\text{sources}(G)$ and $\text{sinks}(G)$, respectively, are the sets of minimals and maximals of

essentially directed reachability, respectively, i.e. $\text{sources}(G) \triangleq \{u \in V_G \mid \nexists v.v <_G u\}$ and $\text{sinks}(G) \triangleq \{u \in V_G \mid \nexists v.u <_G v\}$.

Lemma A.1. *If T_1 and T_2 are disjoint trees and u and v are vertices in T_1 and T_2 , respectively, then the graph G formed by connecting T_1 and T_2 with the edge uv is a tree.*

Lemma A.2. *If G_1 and G_2 are disjoint essentially acyclic mixed graphs, $(u_1, \dots, u_n)$ and $(v_1, \dots, v_n)$ are vertices in G_1 and G_2 , respectively, such that $u_1 < \dots < u_n$ and $v_1 < \dots < v_n$, respectively, then the mixed graph G formed by connecting G_1 and G_2 with the edges $u_1v_1, \dots, u_nv_n$ is essentially acyclic.*

Proof. By contradiction. Assume c is an essentially directed cycle in G .

- If c visits no vertex in G_2 , then c is an essentially directed cycle in G_1 .
- If c visits no vertex in G_1 , then c is an essentially directed cycle in G_2 .
- Otherwise, c must visit some vertex u in G_1 and some vertex v in G_2 .

As c is an essentially directed cycle, there must be distinct paths $p_{uv} = (u, \dots, v)$ and $p_{vu} = (v, \dots, u)$. At least one of p_{uv} and p_{vu} must be essentially directed.

As G_1 and G_2 are disjoint, p_{uv} and p_{vu} must each contain at least one of the edges $u_1v_1, \dots, u_nv_n$. Hence, there must be paths $p_{uu_i} = (u, \dots, u_i)$, $p_{v_iv} = (v_i, \dots, v)$, $p_{vv_j} = (v, \dots, v_j)$, and $p_{u_ju} = (u_j, \dots, u)$ for some $1 \leq i, j \leq n$. At least one of p_{uu_i} , p_{u_ju} , p_{vv_j} , and p_{v_iv} must be essentially directed.

- If $i = j$, then either c contains an arc in G_1 and $p_{uu_i}p_{u_ju}$ is an essentially directed cycle in G_1 , or c contains an arc in G_2 and $p_{v_iv}p_{vv_j}$ is an essentially directed cycle in G_2 .
- If $i < j$, then $u_i < u_j$ and $v_i < v_j$. There must be essentially directed paths $p_{u_iu_j} = (u_i, \dots, u_j)$ in G_1 and $p_{v_iv_j} = (v_i, \dots, v_j)$ in G_2 , and $p_{uu_i}p_{u_iu_j}p_{u_ju}$ is an essentially directed cycle in G_1 .
- If $i > j$, then $u_j < u_i$ and $v_j < v_i$. There must be essentially directed paths $p_{u_ju_i} = (u_j, \dots, u_i)$ in G_1 and $p_{v_jv_i} = (v_j, \dots, v_i)$ in G_2 , and $p_{v_iv}p_{vv_j}p_{v_jv_i}$ is an essentially directed cycle in G_2 .

□

A.2 Multisets

Definition A.3 (Multiset). *A multiset is a variant of a set that allows multiple occurrences of each element. Formally, a multiset $\mathcal{X}$ is a tuple (X, μ_X) where X is the support set of the multiset, and μ_X is a function, giving each element its multiplicity, from X to the class of nonzero cardinal*

numbers. I write multisets as lists of elements between “{” and “}”, e.g. $\{a, b, b, b\}$.

I define the following operations on multisets:

Membership	$a \in^k \mathcal{X} \triangleq a \in X \wedge \mu_X(a) = k$
Union	$\mathcal{X} \cup \mathcal{Y} \triangleq (X \cup Y, \max(\mu_X, \mu_Y))$
Intersection	$\mathcal{X} \cap \mathcal{Y} \triangleq (X \cap Y, \min(\mu_X, \mu_Y))$
Deletion	$\mathcal{X} \setminus \mathcal{Y} \triangleq (X \setminus Y, \mu_X \triangleleft Y)$
Subtraction	$\mathcal{X} - \mathcal{Y} \triangleq (\{a \in X \cup Y \mid \mu_X(a) > \mu_Y(a)\}, \mu_X - \mu_Y)$
Sum	$\mathcal{X} + \mathcal{Y} \triangleq (\mathcal{X} \cup \mathcal{Y}, \mu_X + \mu_Y)$
Product	$\mathcal{X}\mathcal{Y} \triangleq (\mathcal{X} \cap \mathcal{Y}, \mu_X \mu_Y)$
Symmetric Difference	$\mathcal{X} \Delta \mathcal{Y} \triangleq (\mathcal{X} - \mathcal{Y}) \cup (\mathcal{Y} - \mathcal{X})$

Definition A.4 (Domain Subtraction). Domain subtraction, $f \triangleleft X$, is the operation that removes all elements in X from the domain of the function f , i.e. $f \triangleleft X = a \mapsto f(a)$ if $a \notin X$.

Bibliography

- Samson Abramsky. Proofs as processes. *Theoretical Computer Science*, 135 (1):5–9, April 1994. ISSN 0304-3975. doi: 10.1016/0304-3975(94)00103-0.
- Jean-Marc Andreoli. Logic programming with focusing proofs in linear logic. *Journal of Logic and Computation*, 2(3):297–347, 6 1992. ISSN 0955-792X. doi: 10.1093/logcom/2.3.297.
- Robert Atkey. Parameterised notions of computation. *Journal of Functional Programming*, 19(3 & 4):335–376, 2009. doi: 10.1017/S095679680900728X.
- Robert Atkey. Observed communication semantics for classical processes. In Hongseok Yang, editor, *Programming Languages and Systems - 26th European Symposium on Programming, ESOP 2017*, volume 10201 of *Lecture Notes in Computer Science*, pages 56–82, 2017. doi: 10.1007/978-3-662-54434-1_3.
- Robert Atkey, Sam Lindley, and J. Garrett Morris. *Conflation Confers Concurrency*, pages 32–55. Springer International Publishing, Cham, 2016. doi: 10.1007/978-3-319-30936-1_2.
- Arnon Avron. Hypersequents, logical consequence and intermediate logics for concurrency. *Ann. Math. Artif. Intell.*, 4:225–248, 1991. doi: 10.1007/BF01531058. Arnon Avron credits Garrel Pottinger [Pottinger, 1983] with the first use of hypersequents (Footnote 3, p. 227).
- Arnon Avron. The Method of Hypersequents in the Proof Theory of Propositional Non-classical Logics. In *Logic: from Foundations to Applications: European logic colloquium*. Oxford University Press, June 1996. ISBN 9780198538622. doi: 10.1093/oso/9780198538622.003.0001.
- Hendrik Pieter Barendregt. The lambda calculus - its syntax and semantics. In *Studies in Logic and the Foundations of Mathematics*, 1985. URL <https://api.semanticscholar.org/CorpusID:263892017>.
- Gianluigi Bellin and Philip J. Scott. On the π -calculus and linear logic. *Theoretical Computer Science*, 135(1):11–65, April 1994. ISSN 0304-3975. doi: 10.1016/0304-3975(94)00104-9.

- Nuel Belnap. Linear logic displayed. *Notre Dame Journal of Formal Logic*, 31(1), December 1989. ISSN 0029-4527. doi: 10.1305/ndjfl/1093635329.
- Nuel D. Belnap. Display logic. *Journal of Philosophical Logic*, 11(4):375–417, 1982. ISSN 00223611, 15730433. URL <http://www.jstor.org/stable/30226258>.
- Jean-Philippe Bernardy, Mathieu Boespflug, Ryan R. Newton, Simon Peyton Jones, and Arnaud Spiwack. Linear haskell: practical linearity in a higher-order polymorphic language. *Proc. of POPL*, 2: 1–29, 2018. doi: 10.1145/3158093.
- Michele Boreale. *On the expressiveness of internal mobility in name-passing calculi*, page 163–178. Springer Berlin Heidelberg, 1996. ISBN 9783540706250. doi: 10.1007/3-540-61604-7_54. URL http://dx.doi.org/10.1007/3-540-61604-7_54.
- Luís Caires and Frank Pfenning. Session types as intuitionistic linear propositions. In *CONCUR 2010 - Concurrency Theory*, pages 222–236. Springer Berlin Heidelberg, 2010. doi: 10.1007/978-3-642-15375-4_16.
- Luís Caires, Frank Pfenning, and Bernardo Toninho. Linear logic propositions as session types. *Mathematical Structures in Computer Science*, 26(3):367–423, 2016. doi: 10.1017/S0960129514000218. Journal version of Caires and Pfenning [2010].
- Marco Carbone, Sam Lindley, Fabrizio Montesi, Carsten Schürmann, and Philip Wadler. Coherence Generalises Duality: A Logical Explanation of Multiparty Session Types. In Josée Desharnais and Radha Jagadeesan, editors, *27th International Conference on Concurrency Theory (CONCUR 2016)*, volume 59 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pages 33:1–33:15, Dagstuhl, Germany, 2016. Schloss Dagstuhl – Leibniz-Zentrum für Informatik. ISBN 978-3-95977-017-0. doi: 10.4230/LIPIcs.CONCUR.2016.33. URL <https://drops.dagstuhl.de/entities/document/10.4230/LIPIcs.CONCUR.2016.33>.
- B. Jack Copeland, editor. *The Essential Turing*. Clarendon Press, Oxford, England, July 2004. Via Wikipedia [2024] under History.
- Pierre-Évariste Dagand and Conor McBride. Transporting functions across ornaments. *Journal of Functional Programming*, 24(2–3):316–383, 2014. doi: 10.1017/S0956796814000069.
- Vincent Danos and Laurent Regnier. The structure of multiplicatives. *Archive for Mathematical Logic*, 28(3):181–203, October 1989. ISSN 1432-0665. doi: 10.1007/bf01622878.
- Ornela Dardha and Simon J. Gay. A new linear logic for deadlock-free session-typed processes. In Christel Baier and Ugo Dal Lago, editors, *Foundations of Software Science and Computation Structures - 21st International Conference, FOSSACS 2018, Held as Part of the*

- European Joint Conferences on Theory and Practice of Software, ETAPS 2018, Thessaloniki, Greece, April 14-20, 2018, Proceedings*, volume 10803 of *Lecture Notes in Computer Science*, pages 91–109. Springer, 2018. doi: 10.1007/978-3-319-89366-2_5.
- Ornela Dardha, Elena Giachino, and Davide Sangiorgi. Session types revisited. *Inf. Comput.*, 256:253–286, 2017. doi: 10.1016/J.IC.2017.06.002. URL <https://doi.org/10.1016/j.ic.2017.06.002>.
- Martin Davis. Lecture at the Control Systems Laboratory at the University of Illinois, 1952. Via Copeland [2004], Footnote 61, p. 40.
- Martin Davis. *Computability & Unsolvability*. McGraw-Hill, New York, USA, 1958.
- Thomas Ehrhard. Finiteness spaces. *Mathematical Structures in Computer Science*, 15(4):615–646, 2005. doi: 10.1017/S0960129504004645.
- Thomas Ehrhard. An introduction to differential linear logic: proof-nets, models and antiderivatives. *Mathematical Structures in Computer Science*, 28(7):995–1060, 2018. doi: 10.1017/S0960129516000372.
- Thomas Ehrhard and Laurent Regnier. Differential interaction nets. *Theoretical Computer Science*, 364(2):166–195, November 2006. ISSN 0304-3975. doi: 10.1016/j.tcs.2006.08.003.
- Andrzej Filinski. Representing monads. In *Proceedings of the 21st ACM SIGPLAN-SIGACT symposium on Principles of programming languages - POPL '94*. ACM Press, 1994. doi: 10.1145/174675.178047. URL <https://doi.org/10.1145/174675.178047>.
- Simon Fowler. *Typed concurrent functional programming with channels, actors and sessions*. PhD thesis, The University of Edinburgh, 2019.
- Simon Fowler, Sam Lindley, J. Garrett Morris, and Sára Decova. Exceptional asynchronous session types: session types without tiers. *Proc. ACM Program. Lang.*, 3(POPL):28:1–28:29, 2019. doi: 10.1145/3290341.
- Simon Fowler, Wen Kokke, Ornela Dardha, J. Garrett Morris, and Sam Lindley. Separating Sessions Smoothly. In Serge Haddad and Daniele Varacca, editors, *32nd International Conference on Concurrency Theory (CONCUR 2021)*, volume 203 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pages 36:1–36:18, Dagstuhl, Germany, 2021. Schloss Dagstuhl – Leibniz-Zentrum für Informatik. ISBN 978-3-95977-203-7. doi: 10.4230/LIPIcs.CONCUR.2021.36.
- Simon Fowler, Wen Kokke, Ornela Dardha, Sam Lindley, and J. Garrett Morris. Separating Sessions Smoothly. *Logical Methods in Computer Science*, 19(3), July 2023. doi: 10.46298/lmcs-19(3:3)2023. Journal version of Fowler et al. [2021].

- Simon J. Gay and Vasco T. Vasconcelos. Linear type theory for asynchronous session types. *Journal of Functional Programming*, 20(1): 19–50, January 2010. ISSN 0956-7968. doi: 10.1017/S0956796809990268.
- J. Y. Girard and Y. Lafont. Linear logic and lazy computation. In Hartmut Ehrig, Robert Kowalski, Giorgio Levi, and Ugo Montanari, editors, *TAPSOFT '87*, pages 52–66, Berlin, Heidelberg, 1987. Springer Berlin Heidelberg. ISBN 978-3-540-47717-4.
- Jean-Yves Girard. Linear logic. *Theoretical Computer Science*, 50:1–102, 1987.
- Jean-Yves Girard, Paul Taylor, and Yves Lafont. *Proofs and Types*. Cambridge University Press, USA, 1989. ISBN 0521371813.
- Alessio Guglielmi. A system of interaction and structure. *ACM Trans. Comput. Logic*, 8(1):1–es, jan 2007. ISSN 1529-3785. doi: 10.1145/1182613.1182614.
- Kohei Honda. Types for dyadic interaction. In Eike Best, editor, *CONCUR'93*, pages 509–523, Berlin, Heidelberg, 1993. Springer Berlin Heidelberg. ISBN 978-3-540-47968-0.
- Kohei Honda and Olivier Laurent. An exact correspondence between a typed pi-calculus and polarised proof-nets. *Theoretical Computer Science*, 411(22):2223–2238, 2010. ISSN 0304-3975. doi: <https://doi.org/10.1016/j.tcs.2010.01.028>.
- Kohei Honda, Vasco T. Vasconcelos, and Makoto Kubo. Language primitives and type discipline for structured communication-based programming. In Chris Hankin, editor, *Programming Languages and Systems*, pages 122–138, Berlin, Heidelberg, 1998. Springer Berlin Heidelberg. ISBN 978-3-540-69722-0.
- Stephen Cole Kleene. *Introduction to Metamathematics*. D. Van Nostrand Company, Inc., January 1952. Via Wikipedia Wikipedia [2024] under History.
- Naoki Kobayashi. *A New Type System for Deadlock-Free Processes*, page 233–247. Springer Berlin Heidelberg, 2006. ISBN 9783540373773. doi: 10.1007/11817949_16. URL http://dx.doi.org/10.1007/11817949_16.
- Naoki Kobayashi, Benjamin C. Pierce, and David N. Turner. Linearity and the pi-calculus. In *Proceedings of the 23rd ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*, POPL '96, page 358–371, New York, NY, USA, 1996. Association for Computing Machinery. ISBN 0897917693. doi: 10.1145/237721.237804.
- Wen Kokke. Races in Classical Linear Logic. Master's thesis, University of Edinburgh, August 2017.
- Wen Kokke. Rusty Variation: Deadlock-free Sessions with Failure in

- Rust. *Electronic Proceedings in Theoretical Computer Science*, 304:48–60, September 2019. ISSN 2075-2180. doi: 10.4204/eptcs.304.4. Renamed to Sesh.
- Wen Kokke and Ornela Dardha. Prioritise the best variation. In Kirstin Peters and Tim A. C. Willemse, editors, *Formal Techniques for Distributed Objects, Components, and Systems*, pages 100–119, Cham, February 2021a. Springer International Publishing. ISBN 978-3-030-78089-0.
- Wen Kokke and Ornela Dardha. Deadlock-free session types in linear haskell. In *Proceedings of the 14th ACM SIGPLAN International Symposium on Haskell*, Haskell 2021, page 1–13, New York, NY, USA, 2021b. Association for Computing Machinery. ISBN 9781450386159. doi: 10.1145/3471874.3472979.
- Wen Kokke and Ornela Dardha. Prioritise the Best Variation. *Logical Methods in Computer Science*, Volume 19, Issue 4, December 2023. doi: 10.46298/lmcs-19(4:28)2023. Journal version of Kokke and Dardha [2021a].
- Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Better Late Than Never: A fully-abstract semantics for Classical Processes. *Proceedings of the ACM on Programming Languages*, 3(POPL), January 2019a. doi: 10.1145/3290337.
- Wen Kokke, Fabrizio Montesi, and Marco Peressotti. Taking Linear Logic Apart. In Thomas Ehrhard, Maribel Fernández, Valeria de Paiva, and Lorenzo Tortora de Falco, editors, *Proceedings Joint International Workshop on Linearity & Trends in Linear Logic and Applications, Oxford, UK, 7-8 July 2018*, volume 292 of *Electronic Proceedings in Theoretical Computer Science*, pages 90–103. Open Publishing Association, April 2019b. doi: 10.4204/EPTCS.292.5. Version with errata: <https://marcoperessotti.com/files/papers/KMP18-revised.pdf>.
- Sam Lindley and J. Garrett Morris. Sessions as propositions. In Alastair F. Donaldson and Vasco T. Vasconcelos, editors, *Proceedings 7th Workshop on Programming Language Approaches to Concurrency and Communication-cEntric Software, PLACES 2014, Grenoble, France, 12 April 2014*, volume 155 of *EPTCS*, pages 9–16, 2014. doi: 10.4204/EPTCS.155.2. URL <https://doi.org/10.4204/EPTCS.155.2>.
- Sam Lindley and J. Garrett Morris. A semantics for propositions as sessions. In Jan Vitek, editor, *Programming Languages and Systems*, pages 560–584, Berlin, Heidelberg, 2015. Springer Berlin Heidelberg. ISBN 978-3-662-46669-8. doi: 10.1007/978-3-662-46669-8_23.
- Sam Lindley and J. Garrett Morris. Embedding session types in haskell. In *Proceedings of the 9th International Symposium on Haskell*,

- Haskell 2016, page 133–145, New York, NY, USA, 2016a. Association for Computing Machinery. ISBN 9781450344340. doi: 10.1145/2976002.2976018.
- Sam Lindley and J. Garrett Morris. Talking bananas: structural recursion for session types. In *Proceedings of the 21st ACM SIGPLAN International Conference on Functional Programming*. ACM, September 2016b. doi: 10.1145/2951913.2951921.
- Sam Lindley and J. Garrett Morris. *Lightweight Functional Session Types*, pages 265–286. River Publisher, 2017. doi: 10.13052/rp-9788793519817.
- Simon Marlow et al. Haskell 2010 language report. Available online <https://www.haskell.org/onlinereport/haskell2010/>, 2010. Available online <https://www.haskell.org/onlinereport/haskell2010/>.
- Nicholas D. Matsakis and Felix S. Klock. The rust language. *Ada Lett.*, 34(3):103–104, oct 2014. ISSN 1094-3641. doi: 10.1145/2692956.2663188. URL <https://doi.org/10.1145/2692956.2663188>.
- Damiano Mazza. The true concurrency of differential interaction nets. *Mathematical Structures in Computer Science*, 28(7):1097–1125, 2018. doi: 10.1017/S0960129516000402.
- Conor McBride. Ornamental algebras, algebraic ornaments. *Journal of functional programming*, 47, 2010.
- Paul-André Melliès. *A Topological Correctness Criterion for Multiplicative Non-Commutative Logic*, page 283–322. Cambridge University Press, Cambridge, 2004.
- Robin Milner. The polyadic pi-calculus: A tutorial. Technical report, The University of Edinburgh, 1991. URL <http://www.lfcs.inf.ed.ac.uk/reports/91/ECS-LFCS-91-180/>.
- Robin Milner, Joachim Parrow, and David Walker. A calculus of mobile processes, i. *Information and Computation*, 100(1):1–40, 1992a. ISSN 0890-5401. doi: 10.1016/0890-5401(92)90008-4.
- Robin Milner, Joachim Parrow, and David Walker. A calculus of mobile processes, ii. *Information and Computation*, 100(1):41–77, 1992b. ISSN 0890-5401. doi: 10.1016/0890-5401(92)90009-5.
- Fabrizio Montesi and Marco Peressotti. Classical transitions. Technical report, University of Southern Denmark, Odense, Denmark, 2018.
- Fabrizio Montesi and Marco Peressotti. Linear logic, the π -calculus, and their metatheory: A recipe for proofs as processes. *CoRR*, abs/2106.11818, 2021. URL <https://arxiv.org/abs/2106.11818>.
- Dominic Orchard and Nobuko Yoshida. Session types with linearity in Haskell. *Behavioural Types: from Theory to Tools*, page 219, 2017.

- Luca Padovani. Deadlock and lock freedom in the linear π -calculus. In *Proceedings of the Joint Meeting of the Twenty-Third EACSL Annual Conference on Computer Science Logic (CSL) and the Twenty-Ninth Annual ACM/IEEE Symposium on Logic in Computer Science (LICS)*, CSL-LICS '14. ACM, July 2014. doi: 10.1145/2603088.2603116. URL <http://dx.doi.org/10.1145/2603088.2603116>.
- Garrel Pottinger. Uniform, cut-free formulations of t , s_4 and s_5 . *Journal of Symbolic Logic*, 48(3):900, 1983. (In the notes for the Annual Meeting of the Association for Symbolic Logic, Philadelphia, 1981.).
- Zesen Qian. *Concurrency And Races In Classical Linear Logic*. PhD thesis, Århus Universitet, January 2023.
- Zesen Qian, G. A. Kavvos, and Lars Birkedal. Client-server sessions in linear logic. *Proceedings of the ACM on Programming Languages*, 5 (ICFP), August 2021. doi: 10.1145/3473567. Expanded on in Qian's Ph.D. thesis [Qian, 2023].
- Christian Retoré. Pomset logic: A non-commutative extension of classical linear logic. In *Lecture Notes in Computer Science*, pages 300–318. Springer Berlin Heidelberg, 1997. doi: 10.1007/3-540-62688-3_43.
- Davide Sangiorgi. π -calculus, internal mobility, and agent-passing calculi. *Theoretical Computer Science*, 167(1):235–274, 1996. ISSN 0304-3975. doi: 10.1016/0304-3975(96)00075-8.
- Davide Sangiorgi and David Walker. *The π -calculus*. Cambridge University Press, Cambridge, England, October 2003.
- Sergey Slavnov. On noncommutative extensions of linear logic. *Logical Methods in Computer Science*, Volume 15, Issue 3, September 2019. doi: 10.23638/LMCS-15(3:30)2019.
- Lutz Straßburger. A local system for linear logic. In Matthias Baaz and Andrei Voronkov, editors, *Logic for Programming, Artificial Intelligence, and Reasoning*, pages 388–402, Berlin, Heidelberg, 2002. Springer Berlin Heidelberg. ISBN 978-3-540-36078-0.
- Aaron Stump. The calculus of dependent lambda eliminations. *Journal of Functional Programming*, 27:e14, 2017. doi: 10.1017/S0956796817000053.
- Alwen Tiu. A System of Interaction and Structure II: The Need for Deep Inference. *Logical Methods in Computer Science*, Volume 2, Issue 2, April 2006. doi: 10.2168/LMCS-2(2:4)2006. URL <https://lmcs.episciences.org/2252>.
- Philip Wadler. Propositions as sessions. In *Proceedings of the 17th ACM SIGPLAN International Conference on Functional Programming*, ICFP '12, page 273–286, New York, NY, USA, 2012. Association for Computing Machinery. ISBN 9781450310543. doi: 10.1145/2364527.2364568.

- Philip Wadler. Propositions as sessions. *Journal of Functional Programming*, 24(2-3):384–418, 2014. doi: 10.1017/S095679681400001X. Journal version of Wadler [2012].
- Wikipedia. Halting problem — Wikipedia, the free encyclopedia, 2024. URL <https://en.wikipedia.org/w/index.php?title=Halting%20problem&oldid=1184812236>. [Online; accessed 12-January-2024].
- Nobuko Yoshida, Kohei Honda, and Martin Berger. Linearity and bisimulation. In Mogens Nielsen and Uffe Engberg, editors, *Foundations of Software Science and Computation Structures*, pages 417–433, Berlin, Heidelberg, 2002. Springer Berlin Heidelberg. ISBN 978-3-540-45931-6.