



Delft University of Technology

Actris 2.0

Asynchronous session-type based reasoning in separation logic

Hinrichsen, Jonas Kastberg; Bengtson, Jesper; Krebbers, Robbert

DOI

[10.46298/lmcs-18\(2:16\)2022](https://doi.org/10.46298/lmcs-18(2:16)2022)

Publication date

2022

Document Version

Final published version

Published in

Logical Methods in Computer Science

Citation (APA)

Hinrichsen, J. K., Bengtson, J., & Krebbers, R. (2022). Actris 2.0: Asynchronous session-type based reasoning in separation logic. *Logical Methods in Computer Science*, 18(2), 16:1-16:64.
[https://doi.org/10.46298/lmcs-18\(2:16\)2022](https://doi.org/10.46298/lmcs-18(2:16)2022)

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

ACTRIS 2.0: ASYNCHRONOUS SESSION-TYPE BASED REASONING IN SEPARATION LOGIC

JONAS KASTBERG HINRICHSSEN^a, JESPER BENGTON^b, AND ROBBERT KREBBERS^c

^a IT University of Copenhagen and Aarhus University, Denmark
e-mail address: hinrichsen@cs.au.dk

^b IT University of Copenhagen, Denmark
e-mail address: bengton@itu.dk

^c Radboud University and Delft University of Technology, The Netherlands
e-mail address: mail@robertkrebbers.nl

ABSTRACT. Message passing is a useful abstraction for implementing concurrent programs. For real-world systems, however, it is often combined with other programming and concurrency paradigms, such as higher-order functions, mutable state, shared-memory concurrency, and locks. We present **Actris**: a logic for proving functional correctness of programs that use a combination of the aforementioned features. Actris combines the power of modern concurrent separation logics with a first-class protocol mechanism—based on session types—for reasoning about message passing in the presence of other concurrency paradigms. We show that Actris provides a suitable level of abstraction by proving functional correctness of a variety of examples, including a channel-based merge sort, a channel-based load-balancing mapper, and a variant of the map-reduce model, using concise specifications.

While Actris was already presented in a conference paper (POPL’20), this paper expands the prior presentation significantly. Moreover, it extends Actris to **Actris 2.0** with a notion of *subprotocols*—based on session-type subtyping—that permits additional flexibility when composing channel endpoints, and that takes full advantage of the asynchronous semantics of message passing in Actris. Soundness of Actris 2.0 is proven using a model of its protocol mechanism in the Iris framework. We have mechanised the theory of Actris, together with custom tactics, as well as all examples in the paper, in the Coq proof assistant.

1. INTRODUCTION

Message-passing programs are ubiquitous in modern computer systems, emphasising the importance of their functional correctness. Programming languages, like Erlang, Elixir, and Go, have built-in primitives that handle spawning of processes and intra-process communication, while other mainstream languages, such as Java, Scala, F#, and C#, have introduced an Actor model [HBS73] to achieve similar functionality. In both cases the goal remains the same—help design reliable systems, often with close to constant up-time, using lightweight processes that can be spawned by the hundreds of thousands and that communicate via asynchronous message passing.

Key words and phrases: Message passing, actor model, concurrency, session types, Iris.

While message passing is a useful abstraction, it is not a silver bullet of concurrent programming. In a qualitative study of larger Scala projects Tasharofi *et al.* [TDJ13] write:

We studied 15 large, mature, and actively maintained actor programs written in Scala and found that 80% of them mix the actor model with another concurrency model.

In this study, 12 out of 15 projects did not entirely stick to the Actor model, hinting that even for projects that embrace message passing, low-level concurrency primitives like locks (*i.e.*, mutexes) still have their place. Tu *et al.* [TLSZ19] came to a similar conclusion when studying 6 large and popular Go programs. A suitable solution for reasoning about message-passing programs should thus integrate with other programming and concurrency paradigms.

In this paper we introduce **Actris**—a concurrent separation logic for proving functional correctness of programs that combine message passing with other programming and concurrency paradigms. Actris can be used to reason about programs written in a language that mimics the important features found in aforementioned languages such as higher-order functions, higher-order references, fork-based concurrency, locks, and primitives for asynchronous message passing over channels. The channels of our language are first-class and can be sent as arguments to functions, be sent over other channels (often referred to as delegation), and be stored in references.

Program specifications in Actris are written in an impredicative higher-order concurrent separation logic built on top of the Iris framework [JSS⁺15; KJB⁺17; JKBD16; JKJ⁺18]. In addition to the usual features of Iris, Actris provides a notion of *dependent separation protocols* to reason about message passing over channels, inspired by binary session types [HVK98]. We show that dependent separation protocols integrate seamlessly with other concurrency paradigms, allow delegation of resources, support channel sharing over multiple concurrent threads using locks, and more.

1.1. Message passing in concurrent separation logic. Over the last decade, there has been much work on extensions of concurrent separation logic with reasoning principles for message passing [FRS11; LV12; CKC15; OBH16]. These logics typically include some form of mechanism for writing protocol specifications in a high-level manner, to elegantly reason about message passing in some specific context.

In a different line of work, researchers have developed more expressive extensions of concurrent separation logic that support proving strong specifications of programs involving features such as higher-order functions, fine-grained shared-memory concurrency, and locks. Examples of such logics are TaDA [dRPDG14], iCAP [SB14], Iris [JSS⁺15], FCSL [NLS14], and VST [App14]. However, only a few variants and extensions of these logics provide a high-level reasoning mechanism specific to message-passing concurrency.

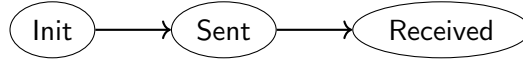
First off, there has been work on the use of Iris-like separation logic to reason about programs that communicate via message passing over a network. The reasoning principles in such logics are geared towards different programming patterns than the ones used in high-level languages like Erlang, Elixir, Go, and Scala. Namely, on networks all data must be serialised, and packets can be lost or delivered out of order. In high-level languages messages cannot get lost, are ensured to be delivered in order, and are allowed to contain many types of data, including functions, references, and even channel endpoints. Two examples of network logics are Disel by Sergey *et al.* [SWT18] and Aneris by Krogh-Jespersen *et al.* [KTO⁺20].

Additionally, there has been work on the use of separation logic to prove compiler correctness of high-level message-passing languages. Tassarotti *et al.* [TJH17] verified a small compiler of a session-typed language into a language where channel buffers are modelled on the heap.

The primary reasoning principle to model the interaction between processes in the aforementioned logics is the notion of a State Transition System (STS). As a simple example, consider the following program, which is borrowed from Tassarotti *et al.* [TJH17]:

$$prog_1 := \text{let } (c, c') := \text{new_chan } () \text{ in fork } \{ \text{send } c' \ 42 \} ; \text{recv } c$$

This program creates two channel endpoints c and c' , forks off a new thread, and sends the number 42 over the channel c' , which is then received by the initiating thread. Modelling the behaviour of this program in an STS typically requires three states:



The three states model that no message has been sent (**Init**), that a message has been sent but not received (**Sent**), and finally that the message has been sent and received (**Received**). Exactly what this STS represents is made precise by the underlying logic, which determines what constitutes a state and a transition, and how these are related to the channel buffers.

While STSs appear like a flexible and intuitive abstraction to reason about message-passing concurrency, they have their problems:

- Coming up with a good STS that makes the appropriate abstractions is difficult because the STS has to keep track of all possible states that the channel buffers can be in, including all possible interleavings of messages in transit.
- While STSs used for the verification of different modules can be composed at the level of the logic, there is no canonical way of composing them due to their unrestrained structure.
- Finally, STSs are first-order meaning that their states and transitions cannot be indexed by propositions of the underlying logic, which limits what they can express when sending messages containing functions or other channels.

1.2. Actris 1.0: Dependent separation protocols. Actris extends separation logic with a notion called *dependent separation protocols*. This notion is inspired by the session type community, pioneered by Honda *et al.* [HVK98], where channel endpoints are given types that describe the expected exchanges. Using session types, the channels c and c' in the program $prog_1$ in § 1.1 would have the types $c : ?\mathbb{Z}.\text{end}$ and $c' : !\mathbb{Z}.\text{end}$, where $!T$ and $?T$ denotes that a value of type T is sent or received, respectively. Moreover, the types of the channels c and c' are *duals*—when one does a send the other does a receive, and *vice versa*.

While session types provide a compact way of specifying the behaviour of channels, they can only be used to talk about the type of data that is being passed around—not their payloads. In this paper, we build on prior work by Bocchi *et al.* [BHTY10] and Craciun *et al.* [CKC15] to attach logical predicates to session types to say more about the payloads, thus vastly extending the expressivity. Concretely, we port session types into separation logic in the form of a construct $c \mapsto \text{prot}$, which denotes ownership of a channel c with dependent separation protocol prot . Dependent separation protocols prot are streams of $!\vec{x}:\vec{\tau}\langle v \rangle\{P\}.\text{prot}$ and $? \vec{x}:\vec{\tau}\langle v \rangle\{P\}.\text{prot}$ constructors that are either infinite or finite, where finite streams are ultimately terminated by an **end** constructor. Here, v is the value that is being sent or received, P is a separation logic proposition denoting the ownership of the

resources being transferred as part of the message, and the variables $\vec{x}:\vec{\tau}$ bind into v , P , and $prot$. The dependent separation protocols for the above example are:

$$c \mapsto ?\langle 42 \rangle \{ \text{True} \}. \text{end} \quad \text{and} \quad c' \mapsto !\langle 42 \rangle \{ \text{True} \}. \text{end}$$

These protocols state that the endpoint c expects the number 42 to be sent along it, and that the endpoint c' expects to send the number 42. Using this protocol, we can prove that $prog_1$ has the specification $\{ \text{True} \} prog_1 \{ v. v = 42 \}$, where v is its resulting value.

Dependent separation protocols $!\vec{x}:\vec{\tau}\langle v \rangle \{ P \}. prot$ and $?\vec{x}:\vec{\tau}\langle v \rangle \{ P \}. prot$ are *dependent*, meaning that the tail $prot$ can be defined in terms of the previously bound variables $\vec{x}:\vec{\tau}$. A sample program showing the use of such dependency is:

```
prog2 := let (c, c') := new_chan () in
         fork { let x := recv c' in send c' (x + 2) };
         send c 40; recv c
```

In this program, the main thread sends the number 40 to the forked-off thread, which then adds two to it, and sends it back. This program has the same specification as $prog_1$, while we change the dependent separation protocol as follows (we omit the dependent separation protocol for the dual endpoint c'):

$$c \mapsto ! (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ?\langle x + 2 \rangle \{ \text{True} \}. \text{end}$$

This protocol states that the second exchanged value is exactly the first with two added to it. To do so, it makes use of a dependency on the variable x , which is used to describe the contents of the first message, which the second message then depends on. This variable is bound in the protocol and it is instantiated only when a message is sent. This is different from the logic by Craciun *et al.* [CKC15], which does not support dependent protocols. Their logic is limited to protocols analogous to $!\langle x \rangle \{ \text{True} \}. ?\langle x + 2 \rangle \{ \text{True} \}. \text{end}$ where x is free, which means the value of x must be known when the protocol is created.

While the prior examples could have been type-checked and verified using the formalisms of Bocchi *et al.* [BHTY10] and Craciun *et al.* [CKC15], the following stateful example cannot:

```
prog3 := let (c, c') := new_chan () in
         fork { let l := recv c' in l ← (!l + 2); send c' () };
         let l := ref 40 in send c l; recv c; !l
```

Here, the main thread stores the value 40 on the heap, and sends a reference l over the channel c to the forked-off thread. The main thread then awaits a signal $()$, notifying that the reference has been updated to 42 by the forked-off thread. This program has the same specification as $prog_1$ and $prog_2$, but the dependent separation protocol is updated:

$$c \mapsto ! (\ell : \text{Loc}) (x : \mathbb{Z}) \langle \ell \rangle \{ \ell \mapsto x \}. ?\langle () \rangle \{ \ell \mapsto (x + 2) \}. \text{end}$$

This protocol denotes that the endpoints first exchange a reference ℓ , as well as a *points-to* connective $\ell \mapsto x$ that describes the ownership and value of the reference ℓ . To perform the exchange c has to give up ownership of the location, while c' acquires it—which is why it can then safely update the received location to 42 before sending the ownership back along with the notification $()$.

The type system by Bocchi *et al.* [BHTY10] cannot verify this program because it does not support mutable state, while Actris can verify the program because it is a separation logic. The logic by Craciun *et al.* [CKC15] cannot verify this program because it does not

support dependent protocols, which are crucial here as they make it possible to delay picking the location ℓ used in the protocol until the send operation is performed.

Dependent protocols are also useful to define recursive protocols to reason about programs that use a channel in a loop. Consider the following variant of $prog_1$:

```

prog4 := let (c, c') := new_chan () in
  fork { let go () := (send c' (recv c' + 2); go ()) in go () };
  send c 18; let x := recv c in
  send c 20; let y := recv c in x + y

```

The forked-off thread will repeatedly interleave receiving values with sending those values back incremented by two. The program $prog_4$ has the same specification as before, but now we use the following recursive dependent separation protocol:

$$c \mapsto \mu rec. ! (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ? \langle x + 2 \rangle \{ \text{True} \}. rec$$

This protocol expresses that it is possible to make repeated exchanges with the forked-off thread to increment a number by two. The fact that the variable x is bound in the protocol is once again crucial—it allows the use of different numbers for each exchange.

Furthermore, Actris inherently includes some features of conventional session types. One such example is the *delegation* of channels as seen in the following program:

```

prog5 := let (c1, c'1) := new_chan () in
  fork { let c := recv c'1 in let y := recv c'1 in send c y; send c'1 () };
  let (c2, c'2) := new_chan () in
  fork { let x := recv c'2 in send c'2 (x + 2) };
  send c1 c2; send c1 40; recv c1; recv c2

```

This program uses the channel pair c_2, c'_2 to exchange the number 40 with the second forked-off thread, which adds 2 to it, and sends it back. Contrary to the programs we have seen before, it uses the additional channel pair c_1, c'_1 to delegate the endpoint c_2 to the first forked-off thread, which then sends the number over c_2 . While this program is intricate, the following dependent separation protocols describe the communication concisely:

$$\begin{aligned}
c_1 &\mapsto ! (c : \text{Val}) \langle c \rangle \{ c \mapsto ! (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ? \langle x + 2 \rangle \{ \text{True} \}. \text{end} \}. \\
&\quad ! (y : \mathbb{Z}) \langle y \rangle \{ \text{True} \}. ? \langle () \rangle \{ c \mapsto ? \langle y + 2 \rangle \{ \text{True} \}. \text{end} \}. \text{end} \\
c_2 &\mapsto ! (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ? \langle x + 2 \rangle \{ \text{True} \}. \text{end}
\end{aligned}$$

The first protocol states that the initial value sent must be a channel endpoint c with the protocol used in $prog_1$. This means that the main thread must give up ownership of the channel endpoint c_2 , thereby delegating it. The first protocol then expects a value y to be sent, and finally to receive a notification $()$, along with ownership of the channel c_2 , which has since taken one step by sending y . Note that c is of the type Val of programming language values due to the programming language being untyped.

Lastly, the dependencies in dependent separation protocols are not limited to first-order data, but can also be used in combination with functions. For example:

```

prog6 := let (c, c') := new_chan () in
  fork { let f := recv c' in send c' (λ_. f() + 2) };
  let l := ref 40 in send c (λ_. !l); recv c ()

```

This program exchanges a value to which 2 is added, but postpones the evaluation by wrapping the computation in a closure. The following protocol is used to verify this program:

$$c \mapsto !(P \ Q : \text{iProp}) (f : \text{Val}) \langle f \rangle \{ \{P\} f () \{v. v \in \mathbb{Z} * Q(v)\} \}. \\ ?(g : \text{Val}) \langle g \rangle \{ \{P\} g () \{v. \exists w. (v = w + 2) * Q(w)\} \}. \text{end}$$

The send constructor (!) does not just bind the function value f , but also the precondition P and postcondition Q of its Hoare triple. In the second message, a Hoare triple is returned that maintains the original pre- and postconditions, but returns an integer of two higher. To send the function, the main thread would let $P \triangleq \ell \mapsto 40$ and $Q(v) \triangleq (v = 40)$, and prove $\{P\} (\lambda_. !\ell) () \{Q\}$. This example demonstrates that the state space of dependent separation protocols can be higher-order—it is indexed by the precondition P and postcondition Q of f —which means that they do not have to be agreed upon when creating the protocol, masking the internals of the function from the forked-off thread.

It is worth noting that using dependent recursive protocols it is possible to keep track of a history of what actions have been performed, which, as is shown in § 7, is especially useful when combining channels with locks.

1.3. Actris 2.0: Subprotocols. While Actris 1.0’s notion of dependent separation protocols is expressive enough to specify advanced exchanges, as indicated by the examples in the previous section, they can only reason about interactions that are strictly dual. In particular, the dual nature of Actris 1.0 requires that:

- Sends $(!\vec{x}:\vec{\tau}\langle v \rangle\{P\})$ are matched up with receives $(?\vec{x}:\vec{\tau}\langle v \rangle\{P\})$, and *vice versa*,
- The logical variables $\vec{x}:\vec{\tau}$ of matched sends and receives are the same, and,
- The propositions P of matched send and receives are the same.

Reasoning about programs with a more relaxed duality principle has been studied in the session type community, namely in the context of *asynchronous session subtyping* [MYH09; MY15]. A subtyping relation $S_1 <: S_2$ captures that the session type S_2 can be used in place of S_1 when type checking a program. Channel endpoints are then allocated with strictly dual session types, after which either side can be weakened based on the subtyping relation. For one, the subtyping relation captures that sends can be swapped *ahead of* receives $?T.!U.S <: !U.?T.S$. Swapping sends ahead of receives is safe to do, as the messages are simply enqueued into the corresponding channel buffer earlier than necessary. The following program illustrates such a non-dual yet safe interaction:

```
prog7 := let (c, c') := new_chan () in
  fork { send c' 20; send c' (recv c' + 2) };
  send c 20;
  let x := recv c in
  let y := recv c in x + y
```

Here, both threads first send the value 20, which is enqueued into both of the channel buffers, after which they receive the value of the other thread. After this, they follow a dual behaviour, where the forked-off thread sends a value, which the main thread receives.

In this paper, we show that dependent separation protocols are compatible with the idea of asynchronous session subtyping. This gives rise to **Actris 2.0**, which supports so-called *subprotocols*. Subprotocols are formalised by a preorder $prot_1 \sqsubseteq prot_2$, which captures (among others) a notion of swapping sends ahead of receives (provided that the send does

not depend on the logical variables of the receive). We can prove that $prog_7$ results in 42 by picking the following dependent separation protocols:

$$\begin{aligned} c &\mapsto ! (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ? \langle 20 \rangle \{ \text{True} \}. ? \langle x + 2 \rangle \{ \text{True} \}. \text{end} \quad \text{and} \\ c' &\mapsto ? (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ! \langle 20 \rangle \{ \text{True} \}. ! \langle x + 2 \rangle \{ \text{True} \}. \text{end} \end{aligned}$$

While the main thread satisfies the protocol of c immediately, the forked-off thread does not satisfy the protocol of c' , as it sends the first value before receiving. However, it is possible to weaken the protocol of c' using Actris 2.0's notion of subprotocols:

$$\begin{aligned} &? (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ! \langle 20 \rangle \{ \text{True} \}. ! \langle x + 2 \rangle \{ \text{True} \}. \text{end} \\ \sqsubseteq &! \langle 20 \rangle \{ \text{True} \}. ? (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ! \langle x + 2 \rangle \{ \text{True} \}. \text{end} \end{aligned}$$

This gives $c' \mapsto ! \langle 20 \rangle \{ \text{True} \}. ? (x : \mathbb{Z}) \langle x \rangle \{ \text{True} \}. ! \langle x + 2 \rangle \{ \text{True} \}. \text{end}$. Since the first send (with value 20) is independent of the variable x bound by the receive, the subprotocol relation follows immediately from the swapping property. Note that it is *not* possible to swap the second send (with value $x + 2$) ahead of the receive, as it does in fact depend on variable x bound by the receive.

In addition to allowing the verification of a larger class of programs, Actris 2.0's subprotocols also provide a more extensional approach to reasoning about dependent separation protocols. This is beneficial whenever we want to reuse existing specifications that might use a syntactically different protocol, but that nonetheless logically entail each another. For example, the ordering of logical variables can be changed using the subprotocol relation:

$$! (x : \mathbb{Z}) (y : \mathbb{Z}) \langle (x, y) \rangle \{ \text{True} \}. \text{prot} \sqsubseteq ! (y : \mathbb{Z}) (x : \mathbb{Z}) \langle (x, y) \rangle \{ \text{True} \}. \text{prot}$$

Since the subprotocol relation is a first-class logical proposition of Actris 2.0, it also allows the manipulation of separation logic resources, such as moving in ownership. For example, we can show the following *conditional* subprotocol relation:

$$\begin{aligned} \ell'_1 &\mapsto 20 \quad \multimap \\ ! (\ell_1, \ell_2 : \text{Loc}) \langle (\ell_1, \ell_2) \rangle \{ \ell_1 \mapsto 20 * \ell_2 \mapsto 22 \}. \text{prot} &\sqsubseteq ! (\ell_2 : \text{Loc}) \langle (\ell'_1, \ell_2) \rangle \{ \ell_2 \mapsto 22 \}. \text{prot} \end{aligned}$$

Here, we move the ownership of $\ell'_1 \mapsto 20$ into the protocol, to resolve the eventual obligation of sending it, while instantiating the logical variable ℓ_1 with ℓ'_1 .

In addition to the demonstrated features, in the rest of this paper we show that Actris 2.0's subprotocol relation is capable of moving resources from one message to another. This gives rise to a principle similar to *framing*, known from conventional separation logic, but applied to dependent separation protocols. Lastly, inspired by the work of Brandt and Henglein [BH98], the subprotocol relation is defined coinductively, allowing us to use the principle of Löb induction to prove subprotocol relations for recursive protocols.

1.4. Formal correspondence to session types. Even though Actris's notion of dependent separation protocols is influenced by binary session types, this paper does not provide a formal correspondence between the two systems. However, since Actris is built on top of Iris, it forms a suitable foundation for building logical relation models of type systems. In related work by Hinrichsen *et al.* [HLKB21], Actris has been used to define a logical relations model of binary session types, with support for various forms of polymorphism and recursion, asynchronous subtyping, references, and locks/mutexes. Similar to RustBelt [JJKD18; JJKD21], the work by Hinrichsen *et al.* [HLKB21] gives rise to an extensible approach for proving type safety, which can be used to manually prove the typing judgements of racy, but safe, programs that cannot be type checked using only the rules of the type system.

1.5. Contributions and outline. This paper introduces **Actris 2.0**: a higher-order impredicative concurrent separation logic built on top of the Iris framework for reasoning about functional correctness of programs with asynchronous message-passing that combine higher-order functions, higher-order references, fork-based concurrency, and locks. Concretely, this paper makes the following contributions:

- We introduce *dependent separation protocols* inspired by affine binary session types to model the transfer of resources (including higher-order functions) between channel endpoints. We show that they can be used to handle choice, recursion, and delegation (§ 2 to 5).
- We introduce *subprotocols* inspired by asynchronous session subtyping. This notion relaxes duality, allowing channels to send messages before receiving others, and gives rise to a more extensional approach to reasoning about dependent separation protocols, providing more flexibility in the design and reuse of protocols. We moreover show how Löb induction is used to reason about recursive subprotocols (§ 6).
- We demonstrate the benefits obtained from building Actris on top of Iris by showing how Iris’s support for ghost state and locks can be used to prove functional correctness of programs using manifest sharing, *i.e.*, channel endpoints shared by multiple parties (§ 7).
- We provide a case study on Actris and its mechanisation in Coq by proving functional correctness of a variant of the map-reduce model by Dean and Ghemawat [DG04] (§ 8).
- We give a model of dependent separation protocols in the Iris framework to prove safety and postcondition validity of our Hoare triples (§ 9).
- We provide a full mechanisation of Actris [HBK21] using the interactive theorem prover Coq. On top of our Coq mechanisation, we provide custom tactics, which we use to mechanise all examples in the paper (§ 10).

1.6. Differences from the conference version. This paper is an extension of the paper “Actris: Session-type based reasoning in separation logic” presented at the POPL’20 conference [HBK20]. In this paper we present Actris 2.0, which extends Actris 1.0 with the notion of subprotocols. This extension introduces new logical connectives and proof rules, but also involves a significant overhaul of the original model and its Coq mechanisation. We extend the presentation of the programming language semantics, model and mechanisation substantially, with additional details, considerations, and examples, to give a better understanding of how Actris works and how it can be used. Concretely, this paper includes the following extensions compared to the conference version:

- An overview of subprotocols in the introduction (§ 1.3).
- A new background section on the programming language semantics (§ 2) and Iris (§ 3).
- A section with an expanded overview of Actris (§ 4, moved from § 5).
- A new section on Actris 2.0’s notion of subprotocols (§ 6).
- An updated and expanded description of the model of Actris in Iris (§ 9).
- An extension of the section on the Coq mechanisation with sample proofs (§ 10).

2. PROGRAMMING LANGUAGE SEMANTICS

The Iris program logic is parametric in the programming language that is used. As a result there are multiple approaches to extend Iris with support for channels:

- Instantiate Iris with a language that has native support for channels. This approach was carried out in the original Iris paper [JSS⁺15] and by Tassarotti *et al.* [TJH17].

- Instantiate Iris with a language that has low-level concurrency primitives, but no native support for channels, and implement channels as a library in that language. This approach was carried out by Bizjak *et al.* [BGKB19] for a lock-free implementation of channels.

In this paper we take the second approach. We implement channels in HeapLang—the default programming language that is shipped with Iris’s Coq development [Iri21]. HeapLang is an untyped functional language with high-level features such as higher-order functions, higher-order mutable references, fork-based concurrency, and garbage collection. Due to these high-level features, programs written in HeapLang are reminiscent of those written in high-level programming languages with message passing like Go or Erlang.

Since HeapLang is an untyped language, safety of a program is not obtained by establishing a typing judgement, but by proving a Hoare triple in the Iris/Actris logic. Hinrichsen *et al.* [HLKB21] show how logical relations in Actris can be used to define and prove sound a session type system for HeapLang extended with message passing.

We proceed by describing HeapLang’s syntax (§ 2.1) and operational semantics (§ 2.2). We then present HeapLang’s standard library for spin locks (§ 2.3). We use this lock library to implement channels (§ 2.4), and to write programs that combine message passing with lock-based concurrency (§ 7).

2.1. Syntax. The syntax of HeapLang is as follows:

$$\begin{aligned}
 v \in \text{Val} &::= () \mid i \mid b \mid \ell \mid \text{rec } f \ x := e \mid & (i \in \mathbb{Z}, b \in \mathbb{B}, \ell \in \text{Loc}) \\
 &(v_1, v_2) \mid \text{inj}_1 v \mid \text{inj}_2 v \mid \dots \\
 e \in \text{Expr} &::= v \mid x \mid e_1 \ e_2 \mid (e_1, e_2) \mid \text{fst } e \mid \text{snd } e \mid \\
 &\text{if } e_1 \text{ then } e_2 \text{ else } e_3 \mid \text{inj}_1 e \mid \text{inj}_2 e \mid \\
 &(\text{match } e_1 \text{ with } (\text{inj}_1 x) \Rightarrow e_2 \mid (\text{inj}_2 x) \Rightarrow e_3 \text{ end}) \mid \\
 &\text{ref } e \mid !e \mid e_1 \leftarrow e_2 \mid & (\text{Mutable state}) \\
 &\text{fork } \{e\} \mid \text{CAS } e_1 \ e_2 \ e_3 \mid \dots & (\text{Concurrency})
 \end{aligned}$$

We elide the standard boolean and arithmetic operators such as equality, addition, subtraction, and multiplication. We define various notions as syntactic sugar (*i.e.*, as definitions in the meta language by use of $\triangleq$):

$$\begin{aligned}
 \lambda x. e &\triangleq \text{rec } _ \ x := e & e_1; e_2 &\triangleq \text{let } _ := e_1 \text{ in } e_2 \\
 \text{let } x := e_1 \text{ in } e_2 &\triangleq (\lambda x. e_2) \ e_1 & \text{skipN} &\triangleq \text{rec } go \ x := \text{if } 0 < x \\
 & & &\text{then } go \ (x - 1) \text{ else } ()
 \end{aligned}$$

We use $_$ as the anonymous binder that is not used in the body of the binding expression. The **skipN** operation, which performs a given number of no-op program steps, is used in the implementation of channels (§ 2.4) for proof-related reasons (explained in § 9.5). We often write definitions as $f \ x_1 \cdots x_n := e$ rather than $f \triangleq \text{rec } f \ x_1 \cdots x_n := e$. For example, we write $\text{skipN } x := \text{if } 0 < x \text{ then skipN } (x - 1) \text{ else } ()$.

HeapLang includes the usual operations for ML-style references. New references can be allocated using **ref** e , dereferenced using $!e$, and updated using $e_1 \leftarrow e_2$. Concurrency is supported via **fork** $\{e\}$, which spawns a new thread e that is executed in the background. The language also supports atomic operations like compare-and-set (**CAS**), which are used to implement lock-free data structures and synchronisation primitives, such as the locks (§ 2.3). HeapLang is garbage collected and thus does not have a deallocation operation.

Call-by-value evaluation contexts:

$$\begin{aligned}
K \in \text{Ctx} ::= & \bullet \mid e \ K \mid K \ v \mid (e_1, K) \mid (K, v_2) \mid \text{fst} \ (K) \mid \text{snd} \ (K) \mid \\
& \text{if} \ K \ \text{then} \ e_1 \ \text{else} \ e_2 \mid \text{inj}_1 \ (K) \mid \text{inj}_2 \ (K) \mid \\
& (\text{match} \ K \ \text{with} \ (\text{inj}_1 \ x) \Rightarrow e_2 \mid (\text{inj}_2 \ x) \Rightarrow e_3 \ \text{end}) \mid \\
& \text{ref} \ (K) \mid !K \mid e \leftarrow K \mid K \leftarrow v \mid \quad (\text{Mutable state}) \\
& \text{CAS} \ e_1 \ e_2 \ K \mid \text{CAS} \ e_1 \ K \ v_2 \mid \text{CAS} \ K \ v_1 \ v_2 \mid \dots \quad (\text{Concurrency})
\end{aligned}$$

Head reductions of HeapLang:

$$\begin{aligned}
((\text{rec} \ f \ x := e)(v); \sigma) & \rightarrow_h (e[v/x][(\text{rec} \ f \ x := e)/f]; \sigma; []) \\
(\text{fst} \ (v_1, v_2); \sigma) & \rightarrow_h (v_1; \sigma; []) \\
(\text{snd} \ (v_1, v_2); \sigma) & \rightarrow_h (v_2; \sigma; []) \\
(\text{if} \ \text{true} \ \text{then} \ e_1 \ \text{else} \ e_2; \sigma) & \rightarrow_h (e_1; \sigma; []) \\
(\text{if} \ \text{false} \ \text{then} \ e_1 \ \text{else} \ e_2; \sigma) & \rightarrow_h (e_2; \sigma; []) \\
\left(\begin{array}{l} \text{match} \ (\text{inj}_i \ v) \ \text{with} \\ \quad (\text{inj}_1 \ x) \Rightarrow e_1 \\ \quad \mid (\text{inj}_2 \ x) \Rightarrow e_2 \\ \text{end} \end{array} ; \sigma \right) & \rightarrow_h (e_i[v/x]; \sigma; []) \quad \text{if } i \in \{1, 2\} \\
(\text{ref} \ v; \sigma) & \rightarrow_h (\ell; \sigma[\ell \leftarrow v]; []) \quad \text{if } \sigma(\ell) = \perp \\
(!\ell; \sigma[\ell \leftarrow v]) & \rightarrow_h (v; \sigma[\ell \leftarrow v]; []) \\
(\ell \leftarrow w; \sigma[\ell \leftarrow v]) & \rightarrow_h ((); \sigma[\ell \leftarrow w]; []) \\
(\text{CAS} \ \ell \ v' \ w; \sigma[\ell \leftarrow v]) & \rightarrow_h (\text{true}; \sigma[\ell \leftarrow w]; []) \quad \text{if } v = v' \\
(\text{CAS} \ \ell \ v' \ w; \sigma[\ell \leftarrow v]) & \rightarrow_h (\text{false}; \sigma[\ell \leftarrow v]; []) \quad \text{if } v \neq v' \\
(\text{fork} \ \{e\}; \sigma) & \rightarrow_h ((); \sigma; [e])
\end{aligned}$$

Thread-local and threadpool reductions of HeapLang:

$$\frac{e_1; \sigma_1 \rightarrow_h e_2; \sigma_2; \vec{e}}{K[e_1]; \sigma_1 \rightarrow_{\text{tl}} K[e_2]; \sigma_2; \vec{e}} \qquad \frac{e_1; \sigma_1 \rightarrow_{\text{tl}} e_2; \sigma_2; \vec{e}}{T \cdot [e_1] \cdot T'; \sigma_1 \rightarrow_{\text{tp}} T \cdot [e_2] \cdot T' \cdot \vec{e}; \sigma_2}$$

Figure 1: The operational semantics of HeapLang.

2.2. Operational semantics. The small-step operational semantics of HeapLang is presented in Figure 1. The type of program states **State** is defined as:

$$\sigma \in \text{State} \triangleq \text{Loc} \xrightarrow{\text{fin}} \text{Val}$$

That is, program states are finite partial maps from allocated locations to their stored values.

The *head reduction* $(e_1; \sigma_1 \rightarrow_h e_2; \sigma_2; \vec{e})$ describes how an expression $e_1 \in \text{Expr}$ in an initial program state $\sigma_1 \in \text{State}$ reduces to a new expression $e_2 \in \text{Expr}$ in a possibly updated program state $\sigma_2 \in \text{State}$. Additionally, it keeps track of a list of newly spawned threads $\vec{e} \in \text{List Expr}$. The reduction rule $(\text{fork} \ \{e\}; \sigma) \rightarrow_h ((); \sigma; [e])$ describes how a new thread e is spawned by adding it to the list of newly spawned threads $[e]$. Conversely, the list of newly spawned threads is empty for all of the other reduction rules.

```

new_lock () := ref false
try_acquire lk := CAS lk false true
acquire lk := if (try_acquire lk) then () else acquire lk
release lk := lk ← false

```

Figure 2: Implementation of locks in HeapLang.

The *thread-local reduction* $(e_1; \sigma_1 \rightarrow_{\text{tl}} e_2; \sigma_2; \vec{e})$ lifts the head reduction to whole expressions. It decomposes the initial expression e_1 into $K[e'_1]$, where K is a *call-by-value evaluation context* [FH92] and a head expression e'_1 . The head expression e'_1 is then reduced, using $(e'_1; \sigma_1 \rightarrow_{\text{h}} e'_2; \sigma_2; \vec{e})$, and the final expression e_2 is set to $K[e'_2]$. Evaluation contexts (shown in Figure 1) provide a deterministic reduction order of sub-expressions. HeapLang reduces right-to-left, meaning that in expressions such as $e_1 \leftarrow e_2$ the expression e_2 reduces before e_1 . This is determined by the corresponding evaluation contexts $e \leftarrow K$ and $K \leftarrow v$, which state that we only evaluate sub-expressions of the target location, once the term to store is a value. More precisely, we would initially get $(e_1 \leftarrow \bullet)[e_2]$. If e_2 reduces to a value v_2 the context syntax dictates that the hole then moves to e_1 yielding $(\bullet \leftarrow v_2)[e_1]$. If e_1 reduces to a value v_1 we finally end up with the expression $v_1 \leftarrow v_2$, as there is no context syntax where both constituents are values, and this expression can be reduced using a standard head reduction.

Finally, the *threadpool reduction* $(\vec{e}_1; \sigma_1 \rightarrow_{\text{tp}} \vec{e}_2; \sigma_2)$ is the top-level reduction relation that describes the interleaving of threads. It describes how a concurrently running list of threads $\vec{e}_1$, in an initial program state σ_1 , reduce to a new list of threads $\vec{e}_2$ in an updated program state σ_2 . At each step a thread e_1 is picked non-deterministically from $\vec{e}_1$ and reduced one step to e_2 via the thread-local reduction $(e_1; \sigma_1 \rightarrow_{\text{tl}} e_2; \sigma_2; \vec{e})$. The final list of threads $\vec{e}_2$ is obtained from $\vec{e}_1$ by replacing the expression e_1 with e_2 and appending the list $\vec{e}$ of newly spawned threads to the end.

We refer the interested reader to Iris Development Team [Iri21, docs/heap_lang.md] for more details on the semantics of HeapLang, and to Jung *et al.* [JKJ⁺18, §6.1] for details on the language-parametric aspects of Iris.

2.3. Implementation of locks. Using HeapLang it is possible to implement various kinds of locks/mutexes. We consider the simplest kind of lock—a spin lock—whose implementation from the HeapLang standard library is shown in Figure 2.

A spin lock implemented using a reference to a boolean, which is `false` if the lock is unlocked, and `true` if the lock is locked. The `new_lock ()` operation creates a new lock lk , which is initially unlocked (*i.e.*, `false`). The operation `acquire lk` will atomically (using compare-and-set) take the lock, or loop if the lock is already taken. The `release lk` operation releases the lock so that it may be acquired by other threads.

2.4. Implementation of channels. Following the literature on asynchronous session types, the message-passing semantics of our channels is *binary* (communication is between two parties), *asynchronous* (sending messages does not block), *bidirectional* (messages can be in transit in both directions simultaneously), *reliable* (messages are never dropped), and *order preserving* (messages always arrive in the order that they were sent).

```

new_chan () := let (l, r, lk) := (lnil (), lnil (), new_lock ()) in
  ((l, r, lk), (r, l, lk))

send c v := let (l, r, lk) := c in
  acquire lk;
  lsnoc l v;
  skipN (llength r);
  release lk

try_recv c := let (l, r, lk) := c in
  acquire lk;
  let ret := (if (lisnil r) then (inj1 ()) else (inj2 (lpop r))) in
  release lk; ret

recv c := match (try_recv c) with
  | inj1 () => recv c
  | inj2 v => v
end

```

Figure 3: Implementation of bidirectional channels in HeapLang.

The implementation of our channels in HeapLang is displayed in Figure 3. It uses locks (§ 2.3) and a linked list library. This list library provides functions for creating a empty list (`lnil`), testing if a list is empty (`lisnil`), computing the length of a list (`llength`), adding an element to the back (`lsnoc`), and popping an element of front (`lpop`). The last two functions mutate the list, instead of creating a copy. The implementation of the list library is standard, and hence elided.

Intuitively, the channels can be thought of as a pair of buffers ($\vec{v}_1, \vec{v}_2$) of unbounded size. The `new_chan` () operation creates a new channel whose buffers are empty, and returns a tuple of endpoints (c_1, c_2). Bidirectionality is obtained by having one endpoint receive from the others send buffer and *vice versa*. As such, the `send ci v` operation enqueues the value v in its own buffer, *i.e.*, $\vec{v}_i$, and the `recv ci` operation dequeues a value from the other buffer, *i.e.*, from $\vec{v}_2$ if $i = 1$ and from $\vec{v}_1$ if $i = 2$. The message passing is asynchronous, as `send c v` will always reduce, while `recv c` will loop as long as the receiving buffer is empty.

More specifically, the `new_chan` function creates new channels by allocating two empty mutable linked lists l and r using `lnil` (), along with a lock lk using `new_lock` (), and returns the tuples (l, r, lk) and (r, l, lk) , where the order of the linked lists l and r determines the side of the endpoints. We refer to the list in the left position as the endpoint's own buffer, and the list in the right position as the other endpoint's buffer.

The `send` function sends a value v over a given channel endpoint (l, r, lk) , by enqueueing it in the l buffer. The function operates in an atomic fashion by first acquiring the lock via `acquire lk`, thereby entering the critical section, after which the value is enqueued (*i.e.*, appended to the end) of the endpoint's own buffer using the function `lsnoc l v`. The `skipN (llength r)` instruction is a no-op that is inserted to aid the proof. We come back to the reason why this instruction is needed in § 9.5.

The `recv` function receives a value over a channel endpoint (l, r, lk) , by dequeuing the first value in the r buffer. It does so by performing a loop that repeatedly calls the helper function `try_recv`. This helper function attempts to receive a value atomically, and

fails if there is no value in the other endpoint's buffer. The function `try_recv` acquires the lock with `acquire lk`, and then checks whether the other endpoint's buffer is empty using `lisnil r`. If it is empty, nothing is returned (*i.e.*, `inj1 ()`), while otherwise the value is dequeued and returned (*i.e.*, `inj2 (lpop r)`).

Throughout the paper, we often use a combined operation for starting a thread and creating a channel between the parent and child thread:

```
start f := let (c, c') := new_chan () in fork {f c'} ; c
```

3. THE IRIS LOGIC

We give a brief introduction to the features of Iris that play an important role in Actris: its support for basic separation logic (§3.1), higher-order impredicative separation logic (§3.2), guarded recursion and step-indexing (§3.3), and Iris's adequacy theorem (§3.4). This section does not present new material, so readers that are already familiar with Iris can skip it. An extensive overview of Iris can be found in [JKJ⁺18], and a tutorial-style introduction can be found in [BB20].

3.1. Basic separation logic. Propositions in separation logic describe ownership of resources, and can thus intuitively be thought of as predicates over resources. The propositions of Iris $P, Q \in \text{iProp}$ range over an extensible set of resources, which includes the program state. Iris is a higher-order separation logic, so it has the usual logical connectives such as conjunction ($P \wedge Q$), implication ($P \Rightarrow Q$), universal ($\forall x:\tau. P$) and existential ($\exists x:\tau. P$) quantification, as well as the connectives of separation logic:

- The *points-to connective* ($\ell \mapsto v$) asserts exclusive resource ownership of a location $\ell \in \text{Loc}$ in the program state, stating that it holds the value $v \in \text{Val}$.
- The *separating conjunction* ($P * Q$) states that P and Q holds for disjoint sets of resources.
- The *separating implication* ($P \multimap Q$) states that by giving up ownership of the resources described by P , we obtain ownership of the resources described by Q . Separating implication is used similarly to implication since (P entails $Q \multimap R$) iff ($P * Q$ entails R).
- The *Hoare triple* $\{P\} e \{w. Q\}$ states that if the initial program state satisfies the precondition P , then (1) the expression e is safe (*i.e.*, does not go wrong), and, (2) if e reduces to a value v , then the final program state satisfies the postcondition $Q[v/w]$. We often omit the binder w in the postcondition if the result is the unit value $()$.

We say that an Iris proposition P is *valid* iff it holds for all resources, *i.e.*, P is valid iff True entails P . Note that $P \multimap Q$ is valid iff P entails Q , so we often use the separating implication ($\multimap$) in place of entailment. For readability, we use inference-style rules to denote separation logic rules ($P_1 * \dots * P_n \multimap Q$) as:

$$\frac{P_1 \quad \dots \quad P_n}{Q}$$

Iris is an affine separation logic, which means that propositions are upwards closed in the resources, *i.e.*, $P * Q$ entails P (rule **AFFINE**). Affinity matches up with the use of a garbage-collected programming language—one can simply dispose of an unused points-to connective $\ell \mapsto v$ using rule **AFFINE** when a location ℓ is no longer referenced.

While many propositions of separation logic assert exclusive ownership of resources (*e.g.*, $\ell \mapsto v$), others do not (*e.g.*, $t = u$). Propositions that do not assert exclusive ownership

Grammar:

$$\begin{aligned}
\tau, \sigma &::= x \mid 0 \mid 1 \mid \mathbb{B} \mid \mathbb{N} \mid \mathbb{Z} \mid \text{Type} \mid \forall x : \tau. \sigma \mid \\
&\quad \text{Loc} \mid \text{Val} \mid \text{Expr} \mid \text{iProp} \mid \text{List } \tau \mid \dots \\
t, u, P, Q &::= x \mid \lambda x : \tau. t \mid t(u) \mid t(\tau) \mid & \text{(Polymorphic lambda-calculus)} \\
&\quad \text{True} \mid \text{False} \mid P \wedge Q \mid P \vee Q \mid P \Rightarrow Q \mid & \text{(Propositional logic)} \\
&\quad \forall x : \tau. P \mid \exists x : \tau. P \mid t = u \mid & \text{(Higher-order logic with equality)} \\
&\quad P * Q \mid P \multimap Q \mid \ell \mapsto v \mid \{P\} e \{v. Q\} \mid & \text{(Separation logic)} \\
&\quad \mu x : \tau. t \mid \triangleright P \mid \dots & \text{(Guarded recursion and step indexing)}
\end{aligned}$$

Basic affine separation logic:

$$\begin{array}{c}
\text{AFFINE} \\
\frac{P * Q}{P} \\
\\
\text{HT-FRAME} \\
\frac{\{P\} e \{w. Q\}}{\{P * R\} e \{w. Q * R\}} \\
\\
\text{HT-VAL} \\
\frac{}{\{\text{True}\} v \{w. w = v\}} \\
\\
\text{HT-FORK} \\
\frac{\{P\} e \{\text{True}\}}{\{P\} \text{fork } \{e\} \{w. w = ()\}} \\
\\
\text{HT-BIND} \\
\frac{\{P\} e \{v. Q\} \quad \forall v. \{Q\} K[v] \{w. R\}}{\{P\} K[e] \{w. R\}} \quad K \text{ a call-by-value evaluation context}
\end{array}$$

Heap manipulation:

$$\begin{array}{c}
\text{HT-ALLOC} \quad \text{HT-LOAD} \quad \text{HT-STORE} \\
\{\text{True}\} \text{ref } v \{\ell. \ell \mapsto v\} \quad \{\ell \mapsto v\} !\ell \{w. (w = v) * \ell \mapsto v\} \quad \{\ell \mapsto v\} \ell \leftarrow w \{\ell \mapsto w\}
\end{array}$$

Guarded recursion and step indexing:

$$\begin{array}{c}
\text{HT-REC} \quad \triangleright\text{-INTRO} \quad \triangleright\text{-MONO} \quad \text{LÖB} \\
\frac{\{P\} e[v/x][\text{rec } f \ x := e/f] \{w. Q\}}{\{\triangleright P\} (\text{rec } f \ x := e) v \{w. Q\}} \quad \frac{P}{\triangleright P} \quad \frac{P \multimap Q}{\triangleright P \multimap \triangleright Q} \quad \frac{\text{Löb}}{\triangleright P \Rightarrow P} \\
\\
\mu\text{-UNFOLD} \\
(\mu x. t) = t[\mu x. t/x]
\end{array}$$

Figure 4: The grammar and a selection of rules of Iris.

enjoy some useful laws. Separation conjunction ($P * Q$) is logically equivalent to regular conjunction ($P \wedge Q$) if at least one conjunct does not assert exclusive ownership, and separating implication ($P \multimap Q$) is logically equivalent to regular implication ($P \Rightarrow Q$) if the premise P does not assert exclusive ownership.¹ For example, $(t = u) * Q$ and $(t = u) \wedge Q$ are logically equivalent. Since separating conjunction/implication is omnipresent in Iris, we prefer the use of separating conjunction/implication over regular conjunction/implication if both can be used. This is also the convention used in the Iris Coq development.

Iris's notion of resources is not limited to locations in the program state (*i.e.*, $\ell \mapsto v$), but can be extended with user-defined *ghost* resources. We use ghost resources to define Actris's connective $c \mapsto \text{prot}$ for exclusive ownership of the channel endpoint c with protocol prot (§9), and to reason about programs with non-trivial sharing (§7).

¹Formally, these equivalences hold for the class of *persistent* propositions, see [JKJ⁺18, §2.3].

The rules for Hoare triples are mostly standard, but it is worth pointing out the rule for HT-BIND. This rule enables reductions of an expression e , in some evaluation context K , based on the precedence enforced by the evaluation contexts presented in § 2.2.

3.2. Higher-order impredicative separation logic. The Iris logic is:

- *Higher-order*: Using Iris’s quantifiers $\forall x : \tau. P$ and $\exists x : \tau. P$ it is not only possible to quantify over first-order types (like $\mathbb{Z}$ and $\text{List } \mathbb{Z}$), but over any type, including functions (like $\mathbb{Z} \rightarrow \mathbb{Z}$), higher-order functions (like $(\mathbb{Z} \rightarrow \mathbb{Z}) \rightarrow \mathbb{Z}$), polymorphic functions (like $\forall T. \text{List } T \rightarrow \mathbb{N}$), Iris propositions (iProp), and Iris predicates (like $\mathbb{Z} \rightarrow \text{iProp}$).
- *Impredicative*: Iris’s logical connectives can be nested arbitrarily. Notably, $\forall P : \text{iProp}. Q$ is an Iris proposition, and not an Iris proposition in a higher universe. Similarly, Hoare triples $\{P\} e \{v. Q\}$ and other Iris connectives like $\text{is_lock } lk \ R$ for lock ownership (§ 7) are first-class Iris propositions themselves.

As we will see in this paper, Actris expands on Iris’s support for higher-order impredicative separation logic by allowing the variables $\vec{x} : \vec{\tau}$ in the dependent separation protocols $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ to range over any type (including Actris’s type of protocols iProto), and the proposition P to contain any Iris/Actris connective (including the Actris connective $c \mapsto \text{prot}$ for channel ownership). This is particularly useful to reason about message-passing programs that transfer functions (§ 5.2) and channels (§ 5.5).

To define (pure) functions and predicates used in program specifications, Iris embeds the polymorphic lambda calculus. In the Coq development of Iris, this lambda calculus is obtained via a shallow embedding, and thus comprises the usual Coq data types and functions.² We should stress that Iris’s lambda calculus is different from our programming language (HeapLang)—the former is typed and pure, whereas the latter is untyped and impure. Consequently, there are two kinds of lambda abstraction ($\lambda x : \tau. t$ for Iris and $\lambda x. e$ for HeapLang). It should be clear from context which of the lambda abstractions is used.

Figure 4 includes a subset of the Iris grammar. The typing judgement is mostly standard and can be derived from the use of meta variables—we use the meta variables P and Q for propositions (type iProp), the meta variable v for values (type Val), and the meta variables t and u for general terms of any type. Similar to Coq, $\lambda x : \tau. t$ is used for both term and type abstraction, and we write $\tau \rightarrow \sigma$ for $\forall x : \tau. \sigma$ if x is free in σ .

3.3. Guarded recursion and step-indexing. Iris is step-indexed [AM01; Ahm04], meaning that propositions are indexed by a natural number—referred to as the *step-index*—which is used to stratify a number of semantically cyclic constructs and reasoning principles. Iris employs the logical account of step-indexing [AMRV07; DAB11] where the step-index is implicit, and internalised in the logic through the *later modality* ($\triangleright$) [Nak00]. Actris and Iris use step-indexing as follows:

- The principle of Löb induction (rule LÖB) is used to reason about (among others) recursive functions. When proving P , Löb induction lets us assume that a proposition holds *later*, denoted $\triangleright P$. The proposition $\triangleright P$ is strictly weaker than P , since P entails $\triangleright P$ (rule $\triangleright$ -INTRO), while the reverse does not hold. The later modality ($\triangleright$) can be eliminated by taking a program step, which is formalised by the Iris proof rule HT-REC. In Actris we use Löb induction to reason about infinite protocols (§ 6.4).

²Coq, Iris, and Actris have a predicative **Type** hierarchy, while propositions are impredicative. For brevity’s sake, we omit details about predicativity of **Type**, as they are standard.

- The guarded recursion operator ($\mu x : \tau. t$) lets us construct recursive predicates without a restriction on the variance of x in t . Instead, the variable x should be *guarded*, which means that it should appear under a *contractive* term construct. The prime example of a contractive construct is the later modality ($\triangleright$). The rule μ -UNFOLD says that $\mu x : \tau. t$ is in fact a fixpoint of t . Actris's dependent separation protocols $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ are contractive in the tail argument $prot$, and thereby make it possible to use Iris's guarded recursion operator to define recursive protocols (§5.4).
- Iris's support for higher-order ghost state [JKBD16] is used to provide a model of Actris in Iris (§9). Additionally, higher-order ghost state is used by Iris to obtain impredicative invariants [SB14], which in turn are used to prove the specification of locks [HAN08] used in §7.

3.4. Adequacy of Iris. The adequacy theorem of Iris connects the derivation of Hoare triples to the operational semantics of the programming language. A closed proof of a Hoare triple gives rise to safety and postcondition validity. By safety we mean that the program cannot go wrong, *e.g.*, by resolving an illegal function application (*e.g.*, `true` + 42), or accessing an invalid location (*i.e.*, $!\ell$ with $\ell \notin \text{dom}(\sigma)$). Safety is defined formally as:

$$\begin{aligned} \text{safe } e \triangleq \forall \sigma, T, \sigma'. ([e]; \sigma \rightarrow_{\text{tp}}^* T; \sigma') \\ \text{implies } \forall e' \in T. (e' \in \text{Val}) \text{ or } \\ (\exists e'', \sigma'', \vec{e}. e'; \sigma' \rightarrow_{\text{tl}} e''; \sigma''; \vec{e}) \end{aligned}$$

This definition is not concerned with whether a program terminates (total correctness).

Postcondition validity means that if the main thread terminates with a value v , then the postcondition holds for that value. This is defined formally as:

$$\begin{aligned} \text{post_valid } (e, \varphi) \triangleq \forall \sigma, v, T, \sigma'. (e; \sigma \rightarrow_{\text{tp}}^* [v] \cdot T; \sigma') \\ \text{implies } (\varphi \ v) \end{aligned}$$

Theorem 3.1 (Adequacy of Iris). *Let $\varphi \in \text{Val} \rightarrow \text{Prop}$ be a meta-level (*i.e.*, Coq) predicate over values and suppose $\{\text{True}\} e \{v. \varphi \ v\}$ is derivable in Iris, then $\text{safe } e$ and $\text{post_valid } (e, \varphi)$.*

4. THE ACTRIS LOGIC

This section describes the core features of Actris 1.0: its dependent separation protocols mechanism (§4.1), proof rules (§4.2), and its adequacy result (§4.3). Actris inherits all features of Iris, which is achieved by defining Actris as an embedded logic in Iris. This means that all of Actris's primitive constructs are defined in Iris, and all of Actris's primitive proof rules are in fact lemmas in Iris. We show how Actris is embedded in Iris in §9.

4.1. Dependent separation protocols. The key feature of Actris is its session-type like dependent separation protocols mechanism. Dependent separation protocols $prot$ are streams of $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ constructors that are either infinite or finite. The finite streams are ultimately terminated by an **end** constructor. The value v denotes the message that is being sent (!) or received (?), the Iris proposition P denotes the ownership that is transferred along the message, and $prot$ denotes the protocol that describes the subsequent messages. The logical variables $\vec{x}:\vec{\tau}$ can be used to bind variables in v , P , and

Grammar:

$$\begin{aligned}
\tau, \sigma &::= \dots \mid \text{iProto} \mid \dots \\
t, u, P, Q, \text{prot} &::= \dots \mid !\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot} \mid ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot} \mid \text{end} \mid \\
&\quad \overline{\text{prot}} \mid \text{prot}_1 \cdot \text{prot}_2 \mid c \multimap \text{prot} \mid \dots
\end{aligned}$$

Dependent separation protocols:

$$\begin{aligned}
\overline{!\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}} &= ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \overline{\text{prot}} & \overline{\text{end}} &= \text{end} \\
\overline{?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}} &= !\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \overline{\text{prot}} & \overline{\overline{\text{prot}}} &= \text{prot} \\
(!\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}_1) \cdot \text{prot}_2 &= !\vec{x}:\vec{\tau}\langle v \rangle\{P\}. (\text{prot}_1 \cdot \text{prot}_2) & \text{prot} \cdot \text{end} &= \text{prot} \\
(?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}_1) \cdot \text{prot}_2 &= ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. (\text{prot}_1 \cdot \text{prot}_2) & \text{end} \cdot \text{prot} &= \text{prot} \\
\text{prot}_1 \cdot (\text{prot}_2 \cdot \text{prot}_3) &= (\text{prot}_1 \cdot \text{prot}_2) \cdot \text{prot}_3 & \overline{\text{prot}_1 \cdot \text{prot}_2} &= \overline{\text{prot}_1} \cdot \overline{\text{prot}_2}
\end{aligned}$$

Message passing:

HT-NEW

$$\{\text{True}\} \text{new_chan } () \{w. \exists c_1, c_2. w = (c_1, c_2) * c_1 \multimap \text{prot} * c_2 \multimap \overline{\text{prot}}\}$$

HT-SEND

$$\{c \multimap !\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot} * P[\vec{t}/\vec{x}]\} \text{send } c (v[\vec{t}/\vec{x}]) \{c \multimap \text{prot}[\vec{t}/\vec{x}]\}$$

HT-RECV

$$\{c \multimap ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}\} \text{recv } c \{w. \exists \vec{y}. w = v[\vec{y}/\vec{x}] * c \multimap \text{prot}[\vec{y}/\vec{x}] * P[\vec{y}/\vec{x}]\}$$

Figure 5: The primitive constructs and proof rules of Actris 1.0.

prot. For example, the following dependent separation protocols expresses that a pair of a boolean and an integer reference whose value is at least 10 is sent:³

$$!(b : \mathbb{B}) (\ell : \text{Loc}) (i : \mathbb{N}) \langle (b, \ell) \rangle \{\ell \mapsto i * 10 < i\}. \text{prot}$$

We often omit the proposition $\{P\}$, which simply means it is **True**.

Apart from the constructors for dependent separation protocols, Actris provides two primitive operations, $\overline{\text{prot}}$ and $\text{prot}_1 \cdot \text{prot}_2$. The $\overline{\text{prot}}$ operator denotes the *dual* of a protocol. Similar to conventional session types, it transforms the protocol by changing all sends (!) into receives (?), and *vice versa*. Taking the dual twice thus results in the original protocol. The operator $\text{prot}_1 \cdot \text{prot}_2$ *appends* the protocols prot_1 and prot_2 , which is achieved by substituting any **end** in prot_1 with prot_2 .

Channel endpoints are ascribed with dependent separation protocols using the channel endpoint ownership connective $c \multimap \text{prot}$, which captures unique ownership of the channel endpoint c and states that the endpoint follows the protocol prot .

4.2. Actris's proof rules for message passing. Actris provides proof rules for the three message passing operations **new_chan**, **send**, and **recv** (see § 2.4 for the definition of these operations). The rule HT-NEW allows ascribing any protocol to newly created channels using

³Note that $\ell \mapsto i * 10 < i$ is logically equivalent to $\ell \mapsto i \wedge 10 < i$ as $10 < i$ does not describe ownership. As discussed in § 3.1, we prefer the version with separation conjunction.

`new_chan ()`, obtaining ownership of $c_1 \mapsto \text{prot}$ and $c_2 \mapsto \overline{\text{prot}}$ for the respective endpoints. The duality of the protocol guarantees that any receive (?) is matched with a send (!) by the dual endpoint, which is crucial for establishing safety.

The rule HT-SEND for `send c w` requires the head of the dependent separation protocol of c to be a send (!) constructor, and the value w that is sent to match up with the ascribed value. To send a message w , we need to give up ownership of $c \mapsto !\vec{x}:\vec{\tau}\langle v \rangle\{P\}.\text{prot}$, pick an appropriate instantiation $\vec{t}$ for the variables $\vec{x}:\vec{\tau}$ so that $w = v[\vec{t}/\vec{x}]$, give up ownership of the associated resources $P[\vec{t}/\vec{x}]$, and finally regain ownership of the protocol tail $c \mapsto \text{prot}[\vec{t}/\vec{x}]$.

The rule HT-RECV for `recv c` is essentially dual to the rule HT-SEND. We need to give up ownership of $c \mapsto ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}.\text{prot}$, and in return acquire the resources $P[\vec{y}/\vec{x}]$, the return value w where $w = v[\vec{y}/\vec{x}]$, and finally the ownership of the protocol tail $c \mapsto \text{prot}[\vec{y}/\vec{x}]$, where $\vec{y}$ is some instantiation of the protocol variables.

Finally, we derive the following specification for the `start` construct from Actris's rule HT-NEW and Iris's rule HT-FORK:

$$\frac{\text{HT-START} \quad \forall c_2. \{c_2 \mapsto \overline{\text{prot}}\} f \ c_2 \ \{\text{True}\}}{\{\text{True}\} \ \text{start} \ f \ \{c_1. c_1 \mapsto \text{prot}\}}$$

4.3. Adequacy of Actris. By virtue of being an extension of Iris, Actris inherits Iris's adequacy theorem (§ 3.4), which says that a closed proof of a Hoare triple gives rise to safety (programs cannot go wrong) and postcondition validity. In Actris this means that the implementation of the message passing operations (§ 2.4) cannot go wrong, and that transferred messages cannot cause the program to go wrong down the line.

Many conventional session-type systems additionally ensure deadlock freedom—which means that program execution cannot result in a state where all threads are waiting on a message to be sent. Deadlock freedom is ensured through a linear type system and combining thread and channel creation into a `start` primitive. Actris is affine (instead of linear), has a `fork` and `new_chan` primitive (instead of a `start` primitive), and supports locks for channel sharing. Actris thus provides more flexibility in terms of what programs can be written and verified (there exist programs that are deadlock free, but cannot be type-checked using conventional session types, while they can be verified using Actris). On the flip side, using Actris one can prove Hoare triples for programs that deadlock, for example:

$$\{\text{True}\} \ \text{let} \ (c, c') ::= \text{new_chan} \ () \ \text{in} \ \text{recv} \ c \ \{\text{True}\}$$

Indeed, in our operational semantics programs such as the above are safe. The semantics of both lock acquisition (`acquire`) and message reception (`recv`) is that the thread loops until it succeeds. Loops are considered safe in Iris (and thus also Actris), as the threads in question will continue to take steps, although they will never terminate.

5. A TOUR OF ACTRIS

This section demonstrates the core features of Actris. We introduce and iteratively extend a simple channel-based merge sort algorithm to demonstrate the main features of Actris (§ 5.1–§ 5.6). Note that as the point of the sorting algorithms is to showcase the features of Actris, they are intentionally kept simple and no effort has been made to make them efficient (*e.g.*, to avoid spawning threads for small jobs).

```

sort_service cmp c :=
  let l := recv c in
  if |l| ≤ 1 then send c () else
  let l' := lsplit l in
  let c1 := start (sort_service cmp) in
  let c2 := start (sort_service cmp) in
  send c1 l; send c2 l';
  recv c1; recv c2;
  lmerge cmp l l'; send c ()

sort_client cmp l :=
  let c :=
    start (sort_service cmp) in
  send c l;
  recv c

```

Figure 6: A channel-based merge sort algorithm (the code for `lmerge` and `lsplit` is standard and thus elided).

5.1. Basic protocols. We first prove functional correctness of a simple channel-based merge sort algorithm, whose code is shown in Figure 6. The function `sort_client cmp l` takes a comparison function `cmp` and a linked list `l` that will be sorted. The function mutates the linked list `l`, so it returns a unit value `()` when done. The bulk of the work is done by the `sort_service cmp c` function, which takes a channel endpoint `c` over which it receives a linked list, and over which it sends back `()` to inform the sender that the list has been sorted. The function `sort_service` is implemented as follows. If the received list is an empty or singleton list, which both are trivially sorted, the function immediately sends back `()`. Otherwise, the list is split into two partitions using `lsplit l`, which mutates the list `l` to contain the first partition, while returning `l'` containing the second partition. These partitions are recursively sorted using two newly started instances of `sort_service`. The results of the processes are then requested and merged using `lmerge cmp l l'`, which mutates the list `l` to contain the merged list. Finally, the unit value `()` is sent back along the original channel endpoint `c`.

In order to verify the correctness of the sorting algorithm we first need a specification for the comparison function `cmp`, which must satisfy the following specification:

$$\begin{aligned}
\text{cmp_spec } (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) (cmp : \text{Val}) \triangleq \\
& (\forall x_1 x_2. R x_1 x_2 \vee R x_2 x_1) * \\
& (\forall x_1 x_2 v_1 v_2. \{I x_1 v_1 * I x_2 v_2\} \text{ cmp } v_1 v_2 \{r. r = R x_1 x_2 * I x_1 v_1 * I x_2 v_2\})
\end{aligned}$$

This definition is polymorphic in type T . Here, R is a total relation in type T , and I is an interpretation predicate that relates language values to elements of type T . While the relation R dictates the ordering, the interpretation predicate I allows for flexibility about what is ordered. Setting I to *e.g.*, $\lambda x v. v \mapsto x$ orders references by what they point to in memory, rather than the memory address itself. To specify how lists are laid out in memory we use the following notation:

$$\ell \mapsto_I^{\text{list}} \vec{x} \triangleq \begin{cases} \ell \mapsto \text{inl } () & \text{if } \vec{x} = \epsilon \\ \exists v_1 \ell_2. \ell \mapsto \text{inr } (v_1, \ell_2) * I x_1 v_1 * \ell_2 \mapsto_I^{\text{list}} \vec{x}_2 & \text{if } \vec{x} = [x_1] \cdot \vec{x}_2 \end{cases}$$

The channel endpoint `c` adheres to the following dependent separation protocol:

$$\begin{aligned}
\text{sort_prot } (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \\
! (\vec{x} : \text{List } T) (\ell : \text{Loc}) \langle \ell \rangle \{ \ell \mapsto_I^{\text{list}} \vec{x} \}. ? \vec{y} \langle () \rangle \{ \ell \mapsto_I^{\text{list}} \vec{y} * \text{sorted_of}_R \vec{y} \vec{x} \}. \text{end}
\end{aligned}$$

```

sort_servicefunc c :=      sort_clientfunc cmp l :=
  let cmp := recv c in    let c := start sort_servicefunc in
  sort_service cmp c      send c cmp; send c l; recv c

```

Figure 7: A version of the sort service that receives the comparison function over the channel.

The protocol describes the interaction of first sending a linked list, and then receiving a unit value () once the list is sorted. The predicate $\text{sorted_of}_R \vec{y} \vec{x}$ is true iff $\vec{y}$ is a sorted version of $\vec{x}$ with respect to the relation R . We prove the following specifications of the service and the client:

$$\begin{array}{ll}
\{ \text{cmp_spec } I \ R \ \text{cmp} * c \mapsto \overline{\text{sort_prot } I \ R} \cdot \text{prot} \} & \{ \text{cmp_spec } I \ R \ \text{cmp} * \ell \xrightarrow{\text{list}}_I \vec{x} \} \\
\text{sort_service } \text{cmp } c & \text{sort_client } \text{cmp } \ell \\
\{ c \mapsto \text{prot} \} & \{ \exists \vec{y}. \text{sorted_of}_R \vec{y} \vec{x} * \ell \xrightarrow{\text{list}}_I \vec{y} \}
\end{array}$$

There are two important things to note about these specifications. First, the protocol $\text{sort_prot } I \ R$ is written from the point of view of the client. As such, the precondition for sort_service requires that c follows the dual $\overline{\text{sort_prot } I \ R}$. Second, the pre- and postcondition of sort_service are generalised to have an arbitrary protocol prot appended at the end. It is important to write specifications this way, so they can be embedded in other protocols. We will see examples of such an embedding in § 5.4 and § 5.5.

The proof of these specifications is almost entirely performed by symbolic execution using the rules HT-NEW, HT-SEND, HT-RECV, and the standard separation logic rules.

Now that we have proven Hoare triples for sort_service and sort_client , we can use them to prove Hoare triples of other programs that use these functions. Recall that if we use them to prove a Hoare triple of a closed program, we obtain safety and postcondition validity by virtue of Actris’s adequacy theorem (§ 3.4).

5.2. Transferring functions. The channel-based sort_service from the previous section (Figure 6) is parametric on a comparison function. To demonstrate Actris’s support for reasoning about functions transferred over channels, we verify the correctness of the function $\text{sort_service}_{\text{func}} \ c$ in Figure 7. This function takes a channel endpoint c , over which it receives the comparison function cmp (instead of via a function argument), followed by the list to sort. Similar to the service in § 5.1, it mutates the list, and sends back () when done. To verify this program, we extend the protocol sort_prot from § 5.1 as follows:

$$\begin{aligned}
\text{sort_prot}_{\text{func}} &\triangleq !(T : \text{Type}) \ (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) \ (R : T \rightarrow T \rightarrow \mathbb{B}) \ (\text{cmp} : \text{Val}) \\
&\quad \langle \text{cmp} \rangle \{ \text{cmp_spec } I \ R \ \text{cmp} \}. \text{sort_prot } I \ R
\end{aligned}$$

The new protocol specifies that we first send a comparison function cmp . It includes binders for the polymorphic type T , the interpretation predicate I , and the relation R . The specifications are much the same as before, with the proofs being similar besides the addition of a symbolic execution step to resolve the sending and receiving of the comparison function:

$$\begin{array}{ll}
\{ c \mapsto \overline{\text{sort_prot}_{\text{func}}} \cdot \text{prot} \} & \{ \text{cmp_spec } I \ R \ \text{cmp} * \ell \xrightarrow{\text{list}}_I \vec{x} \} \\
\text{sort_service}_{\text{func}} \ c & \text{sort_client}_{\text{func}} \ \text{cmp } \ell \\
\{ c \mapsto \text{prot} \} & \{ \exists \vec{y}. \ell \xrightarrow{\text{list}}_I \vec{y} * \text{sorted_of}_R \vec{y} \vec{x} \}
\end{array}$$

```

sort_servicerec cmp c :=
  branch c with
    left  ⇒ sort_service cmp c;
           sort_servicerec cmp c
    | right ⇒ ()
  end

sort_clientrec cmp l :=
  let c := start (sort_servicerec cmp) in
  liter (λl'. select c left; send c l'; recv c) l;
  select c right

```

Figure 8: A recursive version of the sort service that can perform multiple jobs in sequence (the code for the function `liter`, which applies a function to each element of the list, is standard and has been elided).

5.3. Choice. Branching communication is commonly modelled using the *choice* session types $\&$ for branching and $\oplus$ for selection. We show that corresponding dependent separation protocols can readily be encoded in Actris. At the level of the programming language, the instructions for choice are encoded by sending and receiving a boolean value that is matched using an if-then-else construct:

$$\text{select } e \ e' \triangleq \text{send } e \ e'$$

$$\text{branch } e \text{ with } \text{left} \Rightarrow e_1 \mid \text{right} \Rightarrow e_2 \text{ end} \triangleq \text{if } \text{recv } e \text{ then } e_1 \text{ else } e_2$$

The instructions are syntactic sugar, *i.e.*, defined in the meta language (using $\triangleq$), which effectively means that the arguments are evaluated lazily. We define syntactic sugar `left` $\triangleq$ `true` and `right` $\triangleq$ `false` to be used together with `select` for readability's sake.

Due to the higher-order nature of Actris, the usual protocol specifications for choice from session types can be encoded as regular logical branching within the protocols:

$$\text{prot}_1 \{Q_1\} \oplus \{Q_2\} \text{ prot}_2 \triangleq ! (b : \mathbb{B}) \langle b \rangle \{ \text{if } b \text{ then } Q_1 \text{ else } Q_2 \}. \text{if } b \text{ then } \text{prot}_1 \text{ else } \text{prot}_2$$

$$\text{prot}_1 \{Q_1\} \& \{Q_2\} \text{ prot}_2 \triangleq ? (b : \mathbb{B}) \langle b \rangle \{ \text{if } b \text{ then } Q_1 \text{ else } Q_2 \}. \text{if } b \text{ then } \text{prot}_1 \text{ else } \text{prot}_2$$

We often omit the conditions Q_1 and Q_2 , which simply means that they are `True`. The following rules can be directly derived from the rules HT-SEND and HT-RECV:

$$\begin{array}{c} \text{HT-SELECT} \\ \left\{ \begin{array}{l} c \mapsto \text{prot}_1 \{Q_1\} \oplus \{Q_2\} \text{ prot}_2 * \\ \text{if } b \text{ then } Q_1 \text{ else } Q_2 \end{array} \right\} \text{select } c \ b \{ c \mapsto \text{if } b \text{ then } \text{prot}_1 \text{ else } \text{prot}_2 \} \end{array}$$

$$\begin{array}{c} \text{HT-BRANCH} \\ \frac{\{P * Q_1 * c \mapsto \text{prot}_1\} e_1 \{v. R\} \quad \{P * Q_2 * c \mapsto \text{prot}_2\} e_2 \{v. R\}}{\{P * c \mapsto \text{prot}_1 \{Q_1\} \& \{Q_2\} \text{ prot}_2\} \text{branch } c \text{ with } \text{left} \Rightarrow e_1 \mid \text{right} \Rightarrow e_2 \text{ end } \{v. R\}} \end{array}$$

Apart from branching on boolean values, dependent separation protocols can be used to encode choice on any enumeration type (*e.g.*, lists, natural numbers, days of the week, *etc.*). These encodings follow the same scheme.

5.4. Recursive protocols. We now use choice and recursion to verify the correctness of a sorting service that supports performing multiple sorting jobs in sequence. The code of the sorting service `sort_servicerec` and a possible client `sort_clientrec` are displayed in Figure 8. The service `sort_servicerec cmp c` takes a comparison function `cmp` and a channel endpoint `c`, and returns `()`. It contains a loop in which choice is used to either

terminate the service, or to sort an individual list using the channel-based merge sort algorithm `sort_service` from §5.1. The client `sort_clientrec cmp l` takes a comparison function `cmp` and a nested linked list of linked lists `l`, and returns (). It starts a single instance of the service at channel endpoint `c`, and then sequentially sends requests to sort each inner linked list `l'` in `l`. Finally, the client selects the terminating branch to end the communication with the service. A protocol for interacting with the sorting service can be defined as follows:

$$\text{sort_prot}_{\text{rec}} (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \\ \mu(\text{rec} : \text{iProto}). (\text{sort_prot } I \ R \cdot \text{rec}) \oplus \text{end}$$

The protocol uses the choice operator $\oplus$ to specify that the client may either request the service to perform a sorting job, or terminate communication with the service. After the job has been finished the protocol proceeds recursively.

We use Iris’s operator $\mu x : \tau. t$ for guarded recursion (§3.3) to define recursive protocols. It is important to recall that—as is usual in logics with guarded recursion—the variable x should appear under a *contractive* term construct in the body t of $\mu x : \tau. t$. In our protocol, the recursive variable rec appears under the argument of $\oplus$, which is defined in terms of $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$, which, similarly to $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$, is contractive in the tail protocol prot . We can then prove the following specifications of the service and the client:

$$\left\{ \begin{array}{l} \text{cmp_spec } I \ R \ \text{cmp} \ * \\ c \mapsto \text{sort_prot}_{\text{rec}} \ I \ R \cdot \text{prot} \\ \text{sort_service}_{\text{rec}} \ \text{cmp} \ c \\ c \mapsto \text{prot} \end{array} \right\} \quad \left\{ \begin{array}{l} \text{cmp_spec } I \ R \ \text{cmp} \ * \ \ell \mapsto_J^{\text{list}} \vec{x} \\ \text{sort_client}_{\text{rec}} \ \text{cmp} \ \ell \\ \left\{ \exists \vec{y}. |\vec{y}| = |\vec{x}| * \ell \mapsto_J \vec{y} * (\forall i < |\vec{x}|. \text{sorted_of}_R \ \vec{y}_i \ \vec{x}_i) \right\} \end{array} \right\}$$

We let $J \triangleq \lambda \ell' \vec{y}. \ell' \mapsto_I^{\text{list}} \vec{y}$ to express that ℓ points to a list of lists $\vec{x}$. The proof of the service follows naturally by symbolic execution using the induction hypothesis (obtained from LÖB), the rules HT-BRANCH and HT-SELECT, and the specification of `sort_service`. Note that we rely on the specification of `sort_service` having an arbitrary protocol as its suffix.

It is worth pointing out that protocols in Actris provide a lot of flexibility. Using just minor changes, we can extend the protocol to support transferring a comparison function over the channel, like the extension made in `sort_clientfunc`, or in a way such that a different comparison function can be used for each sorting job.

5.5. Higher-order protocols. Higher-order communication is a common feature within communication protocols, and particularly the session-types community—it is the concept of transferring a channel endpoint over a channel, often called delegation. Due to the impredicativity of dependent separation protocols in Actris, higher-order reasoning about programs with delegation is readily available. The protocols $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ can simply refer to the channel endpoint ownership $c \mapsto \text{prot}'$ in the proposition P .

An example of a program that uses delegation is the `sort_servicedel` variant of the recursive sorting service in Figure 9, which allows multiple sorting jobs to be performed in parallel. The function `sort_servicedel cmp c` takes a comparison function `cmp`, a channel endpoint `c`, and returns (). Using the channel endpoint `c`, a client can request the service to start a new inner sorting service `c'`, which the service delegates over channel endpoint `c`.

Similar to the client in §5.4, the client `sort_clientdel cmp l` takes a comparison function `cmp` and a nested linked list of linked lists `l`, and returns (). The client starts a connection `c` to the service, and for each inner list `l'`, it acquires a delegated channel endpoint `c'`, over

```

sort_servicedel cmp c :=
  branch c with
    left ⇒
      let c' :=
        start (sort_service cmp) in
        send c c';
        sort_servicedel cmp c
    | right ⇒ ()
  end

sort_clientdel cmp l :=
  let c := start (sort_servicedel cmp) in
  let k := lnil () in
  liter (λl'. select c left;
        let c' := recv c in
        send c' l'; lcons c' k) l
  select c right;
  liter recv k

```

Figure 9: A recursive version of the sort service that uses delegation to perform multiple jobs in parallel (the code for the function `lcons`, which pushes an element to the head of a list, has been elided).

which it sends the inner list l' that should be sorted. The client keeps track of all channels to delegated services in a linked list k so that it can wait for all of them to finish (using `liter recv`).

A protocol for the delegation service can be defined as follows, denoting that the client can select whether to acquire a connection to a new delegated service or to terminate:

$$\text{sort_prot}_{\text{del}} (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \\ \mu(\text{rec} : \text{iProto}). (? (c : \text{Val}). \langle c \rangle \{c \mapsto \text{sort_prot } I \ R\}. \text{rec}) \oplus \text{end}$$

We can then prove the following specifications of the service and the client:

$$\left\{ \begin{array}{l} \text{cmp_spec } I \ R \ \text{cmp} * \\ c \mapsto \text{sort_prot}_{\text{del}} \ I \ R \cdot \text{prot} \\ \text{sort_service}_{\text{del}} \ \text{cmp} \ c \\ c \mapsto \text{prot} \end{array} \right\} \quad \left\{ \begin{array}{l} \text{cmp_spec } I \ R \ \text{cmp} * \ell \xrightarrow{\text{list}}_J \vec{x} \\ \text{sort_client}_{\text{del}} \ \text{cmp} \ \ell \\ \exists \vec{y}. |\vec{y}| = |\vec{x}| * \ell \xrightarrow{\text{list}}_J \vec{y} * (\forall i < |\vec{x}|. \text{sorted_of}_R \ \vec{y}_i \ \vec{x}_i) \end{array} \right\}$$

As before, we let $J \triangleq \lambda \ell'. \vec{y}. \ell' \xrightarrow{\text{list}}_I \vec{y}$ to express that ℓ points to a list of lists $\vec{x}$. Once again the proofs are straightforward, as they are simply a combination of recursive reasoning combined with the application of Actris's rules for channels.

5.6. Dependent protocols. The protocols we have seen so far have only made limited use of Actris's support for recursion. We now demonstrate Actris's support for dependent protocols, which make it possible to keep track of the history of what messages have been sent and received. We demonstrate this feature by considering a fine-grained version of the channel-based merge-sort service as shown in Figure 10. Like previous versions, the function `sort_servicefg cmp c` takes a comparison function `cmp` and a channel endpoint `c`, and returns `()`. However, unlike previous versions, the input list should be transferred element by element over the channel endpoint `c` to the service, and when done, the service sends back the sorted list element by element. We use choice to indicate whether the whole list has been sent (`right`) or another element remains to be sent (`left`).

The structure of `sort_servicefg` is somewhat similar to the coarse-grained merge-sort algorithm that we have seen before. The base cases of the empty or the singleton list are handled initially. This is achieved by waiting for at least two values before starting the recursive sub-services c_1 and c_2 . In the base cases the values are sent back immediately,

```

sort_servicefg cmp c :=
  branch c with
    right ⇒ select c right
  | left ⇒
    let x1 := recv c in
    branch c with
      right ⇒ select c left; send c x1;
              select c right
    | left ⇒
      let x2 := recv c in
      let c1 := start (sort_servicefg cmp) in
      let c2 := start (sort_servicefg cmp) in
      select c1 left; send c1 x1;
      select c2 left; send c2 x2;
      splitfg c c1 c2; mergefg cmp c c1 c2
    end
  end
end

splitfg c c1 c2 :=
  branch c with
    right ⇒ select c1 right;
            select c2 right
  | left ⇒
    let x := recv c in
    select c1 left; send c1 x;
    splitfg c c2 c1
  end

mergefg cmp c c1 c2 :=
  branch c1 with
    right ⇒ assert false
  | left ⇒
    let x := recv c1 in
    mergefgaux cmp c x c1 c2
  end

mergefgaux cmp c x c1 c2 :=
  branch c2 with
    right ⇒ select c left; send c x1;
            transfer c1 c
  | left ⇒
    let y := recv c2 in
    if cmp x y then
      select c left; send c x;
      mergefgaux cmp c y c2 c1
    else
      select c left; send c y;
      mergefgaux cmp c x c1 c2
    end
  end

sort_clientfg cmp l :=
  let c :=
    start (sort_servicefg cmp) in
  send_all c l; recv_all c l

```

Figure 10: A fine-grained version of the sort service that transfers elements one by one (the code for the functions **transfer**, **send_all**, and **recv_all** has been elided).

as they are trivially sorted. The inductive case is handled by starting two sub-services at the channel endpoints c_1 and c_2 . First, each of the channel endpoints are sent one of the two initially received elements. The remaining elements are then received by the parent service on c , and forwarded to the sub-services alternatingly on c_1 and c_2 , using the function $\text{split}_{fg} c c_1 c_2$. Once the **right** flag is received, the split_{fg} function terminates, and the algorithm moves to the second phase.

In the second phase, the function $\text{merge}_{fg} cmp c c_1 c_2$ is used to merge the stream of elements returned by the sub-services on c_1 and c_2 and forwards them to the parent service on c . It initially acquires the first element x from the first sub-service on c_1 , which it passes to the auxiliary function merge_{fg}^{aux} as the current largest value. The auxiliary function $\text{merge}_{fg}^{aux} cmp c x c_1 c_2$ recursively requests a value y from the sub-service from which the current largest value was not acquired from (initially c_2). It then compares x and y using the comparison function cmp , and forwards the smallest element on c . This is repeated until the **right** flag is received from either sub-service, after which the remaining values of the other sub-service are forwarded to the parent service on c using **transfer** $c_1 c$.

The interface of the client `sort_clientfg cmp l` is similar to the one from § 5.1 and 5.2. It takes a comparison function `cmp` and a linked lists `l`, sorts the linked list `l`, and returns `()` when done. The client sorts the list `l` by sending its elements to the sort service using the `send_all c l` function (which mutates the list `l` by removing all of its values and sending them over the channel `c`), and puts the received values back into the linked list using the `recv_all c l` function (which also mutates the list `l`). A suitable protocol for proving functional correctness of the fine-grained sorting service is as follows:

$$\begin{aligned}
& \text{sort_prot}_{fg} (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \text{sort_prot}_{fg}^{\text{head}} I R \epsilon \\
& \text{sort_prot}_{fg}^{\text{head}} (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \mu(\text{rec} : \text{List } T \rightarrow \text{iProto}). \\
& \quad \lambda \vec{x}. (! (x : T) (v : \text{Val}) \langle v \rangle \{ I x v \}. \text{rec } (\vec{x} \cdot [x])) \oplus \text{sort_prot}_{fg}^{\text{tail}} I R \vec{x} \epsilon \\
& \text{sort_prot}_{fg}^{\text{tail}} (I : T \rightarrow \text{Val} \rightarrow \text{iProp}) (R : T \rightarrow T \rightarrow \mathbb{B}) \triangleq \mu(\text{rec} : \text{List } T \rightarrow \text{List } T \rightarrow \text{iProto}). \\
& \quad \lambda \vec{x} \vec{y}. (? (y : T) (v : \text{Val}) \langle v \rangle \{ (\forall i < |\vec{y}|. R \vec{y}_i y) * I y v \}. \text{rec } \vec{x} (\vec{y} \cdot [y])) \ \&_{\{\vec{x} \equiv_p \vec{y}\}} \text{end}
\end{aligned}$$

The protocol is split into two phases `sort_protfghead` and `sort_protfgtail`, mimicking the behaviour of the program. The `sort_protfghead` phase is indexed by the values $\vec{x}$ that have been sent so far. The protocol describes that one can either send another value and proceed recursively, or stop, which moves the protocol to the next phase.

The `sort_protfgtail` phase is dependent on the list of values $\vec{x}$ received in the first phase, and the list of values $\vec{y}$ returned so far. The condition $(\forall i < |\vec{y}|. R \vec{y}_i y)$ states that the received element is larger than any of the elements that have previously been returned, which maintains the invariant that the stream of received elements is sorted. When the `right` flag is received $\vec{x} \equiv_p \vec{y}$ shows that the received values $\vec{y}$ are a permutation of the ones $\vec{x}$ that were sent, making sure that all of the sent elements have been accounted for.

We can then prove top-level specifications for the service and client that are similar to the coarse-grained version of the channel-based merge sort:

$$\begin{array}{ll}
\{ \text{cmp_spec } I R \text{ cmp} * c \mapsto \overline{\text{sort_prot}_{fg}} I R \cdot \text{prot} \} & \{ \text{cmp_spec } I R \text{ cmp} * \ell \stackrel{\text{list}}{\mapsto}_I \vec{x} \} \\
\text{sort_prot}_{fg} c & \text{sort_client}_{fg} \text{ cmp } \ell \\
\{ c \mapsto \text{prot} \} & \{ \exists \vec{y}. \ell \stackrel{\text{list}}{\mapsto}_I \vec{y} * \text{sorted_of}_R \vec{y} \vec{x} \}
\end{array}$$

Proving these specifications requires one to pick appropriate specifications for the auxiliary functions to capture the required invariants with regard to sorting. After having picked these specifications, the parts of the proofs that involve communication are mostly straightforward, but require a number of trivial auxiliary results about sorting and permutations.

6. SUBPROTOCOLS

This section describes **Actris 2.0**, which extends Actris 1.0—as presented in the conference version of this paper [HBK20]—with *subprotocols*, inspired by asynchronous subtyping of session types [MYH09; MY15]. The intention of both of these relations is to capture protocol-preserving changes, that allow for some internal flexibility of how an endpoint fulfills a protocol, while being indistinguishable by the other endpoint. In particular, subprotocols have two key features. First, they exploit the asynchronous semantics of channels by relaxing the notion of duality, thereby making it possible to prove functional correctness of a larger class of programs. Second, they give rise to a more extensional approach to reasoning about

dependent separation protocols, as we can work up to the subprotocol relation rather than equality, thereby providing more flexibility in the design and reuse of protocols.

We first introduce Actris 2.0’s subprotocol relation and its proof rules (§6.1). These should (similar to the Actris 1.0 logic, presented in §4) be considered to be primitives of Actris; in §9.2 we define and prove them in Iris. We then show how subprotocols can be employed to prove a mapper service, which handles requests one at a time, while its client may send multiple requests up front (§6.2). Next, we demonstrate how the subprotocol relation allows for the composition of slightly differing protocols, by composing a list reversal service whose protocol is based on a list predicate that does not carry ownership, with a client whose protocol is based on a list predicate that does carry ownership (§6.3). Finally, we show that the subprotocol relation is coinductive, and, when combined with Löb induction, can be used to reason about recursive protocols (§6.4).

6.1. The subprotocol relation. The dependent separation protocols of channel endpoints are picked on channel creation (using the rule HT-NEW shown in Figure 5), which then determines how the channel endpoints should interact. To ensure safe communication, Actris adapts the notion of duality from session types, which requires every send (!) of one endpoint to be paired with a receive (?) for the other endpoint, and *vice versa*. However, working with a channel’s protocol and its dual is more restrictive than strictly necessary. Some variations from the original protocol preserve the externally observed interaction, as the other endpoint is agnostic to the variations in question, which will be made clear momentarily. We capture some of these so-called protocol-preserving variations via a new notion—the *subprotocol relation*—denoted as follows:

$$prot_1 \sqsubseteq prot_2$$

The subprotocol relation describes that protocol $prot_1$ is *stronger* than $prot_2$, or conversely, that protocol $prot_2$ is *weaker* than $prot_1$. More specifically, this means that $prot_2$ can be used *in place of* $prot_1$ whenever such a protocol is expected during verification. This property is captured by the following monotonicity rule for channel ownership:

$$\frac{c \multimap prot_1 \quad prot_1 \sqsubseteq prot_2}{c \multimap prot_2}$$

The subprotocol relation is inspired by asynchronous subtyping for session types [MYH09; MY15], which allows (1) sending subtypes (contravariance), (2) receiving supertypes (covariance), and (3) swapping sends ahead of receives. These variations preserve the protocol, as (1) the originally expected type that is to be sent can be derived from the subtype, (2) the originally expected type to be received can be derived from the supertype, and (3) sends do not block because channels are buffered in both directions, so messages can be enqueued ahead of time. These variations, including the swapping property, are generalised to dependent separation protocols using the following proof rules:

$$\begin{array}{c} \text{⊆-SEND-MONO'} \\ \frac{\forall \vec{x}:\vec{\tau}. P_2 \multimap P_1 \quad \forall \vec{x}:\vec{\tau}. prot_1 \sqsubseteq prot_2}{!\vec{x}:\vec{\tau} \langle v \rangle \{P_1\}. prot_1 \sqsubseteq !\vec{x}:\vec{\tau} \langle v \rangle \{P_2\}. prot_2} \end{array} \quad \begin{array}{c} \text{⊆-RECV-MONO'} \\ \frac{\forall \vec{x}:\vec{\tau}. P_1 \multimap P_2 \quad \forall \vec{x}:\vec{\tau}. prot_1 \sqsubseteq prot_2}{?\vec{x}:\vec{\tau} \langle v \rangle \{P_1\}. prot_1 \sqsubseteq ?\vec{x}:\vec{\tau} \langle v \rangle \{P_2\}. prot_2} \end{array}$$

$$\begin{array}{c} \text{⊆-SWAP'} \\ ?\vec{x}:\vec{\tau} \langle v \rangle \{P\}. !\vec{y}:\vec{\sigma} \langle w \rangle \{Q\}. prot \sqsubseteq !\vec{y}:\vec{\sigma} \langle w \rangle \{Q\}. ?\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot \end{array}$$

The rules $\sqsubseteq\text{-SEND-MONO}'$ and $\sqsubseteq\text{-RECV-MONO}'$ use *separation implication* $P \multimap Q$ —which states that ownership of Q can be obtained by giving up ownership of P —to mimic the contra- and covariance of session subtyping. The rule $\sqsubseteq\text{-SWAP}'$ states that sends can be swapped ahead of receives. To be well-formed, this rule has the implicit side condition that $\vec{x}:\vec{\tau}$ does not bind into w and Q , and that $\vec{y}:\vec{\sigma}$ does not bind into v and P .

To give an intuition behind the protocol-consistent changes that the above rules capture, consider the following subprotocol derivation:

$$\begin{array}{ll} ?(i : \mathbb{Z}) \langle i \rangle \{i < 42\}. ! (j : \mathbb{Z}) \langle j \rangle \{j > 42\}. prot & \sqsubseteq\text{-SEND-MONO}' \\ \sqsubseteq ?(i : \mathbb{Z}) \langle i \rangle \{i < 42\}. ! (j : \mathbb{Z}) \langle j \rangle \{j > 50\}. prot & \sqsubseteq\text{-RECV-MONO}' \\ \sqsubseteq ?(i : \mathbb{Z}) \langle i \rangle \{i < 40\}. ! (j : \mathbb{Z}) \langle j \rangle \{j > 50\}. prot & \sqsubseteq\text{-SWAP}' \\ \sqsubseteq ! (j : \mathbb{Z}) \langle j \rangle \{j > 50\}. ?(i : \mathbb{Z}) \langle i \rangle \{i < 40\}. prot & \end{array}$$

Here, we first strengthen the proposition of the send (by increasing the bound from $j > 42$ to $j > 50$), then weaken the proposition of the receive (by reducing the bound from $i < 42$ to $i < 40$), and finally swap the send ahead of the receive.

While the aforementioned rules cover the intuition behind Actris's subprotocol relation, Actris's actual subprotocol rules provide a number of additional features:

- (1) They can be used to manipulate the logical variables $\vec{x}:\vec{\tau}$ that appear in protocols.
- (2) They can be used to transfer ownership of resources in and out of messages.
- (3) They can be used to reason about recursive protocols defined using Löb induction.

The full set of primitive rules for subprotocols is shown in Figure 11. The first four rules account for logical variable manipulation and resource transfer: Rules $\sqsubseteq\text{-SEND-OUT}$ and $\sqsubseteq\text{-RECV-OUT}$ generalise over the logical variables $\vec{x}:\vec{\tau}$ and transfer ownership of P out of the weaker sending protocol $! \vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$, and stronger receiving protocol $? \vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$, respectively. Rule $\sqsubseteq\text{-SEND-IN}$ weakens a sending protocol $! \vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$ by instantiating the logical variables $\vec{x}:\vec{\tau}$ and transferring ownership of $P[\vec{t}/\vec{x}]$ into the protocol. Dually, the rule $\sqsubseteq\text{-RECV-IN}$ strengthens a receiving protocol $? \vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$ by instantiating the logical variables $\vec{x}:\vec{\tau}$ and transferring ownership of $P[\vec{t}/\vec{x}]$ into the protocol.

To demonstrate the intuition behind these rules consider the following proof of the subprotocol relation presented in § 1.3, where we transfer ownership of $\ell'_1 \mapsto 20$ into a protocol, while instantiating the logical variable ℓ_1 with ℓ'_1 :

$$\begin{array}{l} \frac{\ell'_1 \mapsto 20 * \ell_2 \mapsto 22 \multimap \ell'_1 \mapsto 20 * \ell_2 \mapsto 22}{\ell'_1 \mapsto 20 * \ell_2 \mapsto 22 \multimap} \sqsubseteq\text{-SEND-IN} \\ \frac{\ell'_1 \mapsto 20 * \ell_2 \mapsto 22 \multimap \quad !(\ell_1, \ell_2 : \text{Loc}) \langle (\ell_1, \ell_2) \rangle \{ \ell_1 \mapsto 20 * \ell_2 \mapsto 22 \}. prot \sqsubseteq ! \langle (\ell'_1, \ell_2) \rangle. prot}{\ell'_1 \mapsto 20 \multimap} \sqsubseteq\text{-SEND-OUT} \\ \ell'_1 \mapsto 20 \multimap \quad !(\ell_1, \ell_2 : \text{Loc}) \langle (\ell_1, \ell_2) \rangle \{ \ell_1 \mapsto 20 * \ell_2 \mapsto 22 \}. prot \sqsubseteq \\ !(\ell_2 : \text{Loc}) \langle (\ell'_1, \ell_2) \rangle \{ \ell_2 \mapsto 22 \}. prot \end{array}$$

We first use rule $\sqsubseteq\text{-SEND-OUT}$ to generalise over the logical variable ℓ_2 and transfer ownership of $\ell_2 \mapsto 22$ out of the weaker protocol (*i.e.*, the send on the RHS), and then use $\sqsubseteq\text{-SEND-IN}$ to instantiate the logical variables ℓ'_1 and ℓ_2 and transfer ownership of $\ell'_1 \mapsto 20$ and $\ell_2 \mapsto 22$ into the stronger protocol (*i.e.*, the send on the LHS).

The rules for monotonicity ($\sqsubseteq\text{-SEND-MONO}$ and $\sqsubseteq\text{-RECV-MONO}$) and swapping ($\sqsubseteq\text{-SWAP}$) in Figure 11 differ in two aspects from the rules for monotonicity ($\sqsubseteq\text{-SEND-MONO}'$ and $\sqsubseteq\text{-RECV-MONO}'$) and swapping ($\sqsubseteq\text{-SWAP}'$) that we have seen in the beginning of this section. First, the actual rules only apply to protocols whose head does not have logical variables

Grammar:

$$t, u, P, Q, prot ::= \dots \mid prot_1 \sqsubseteq prot_2 \mid \dots$$

Logical variable manipulation and resource transfer:

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-SEND-OUT}} \\
\frac{\forall \vec{x}:\vec{\tau}. P \multimap (prot_1 \sqsubseteq !\langle v \rangle. prot_2)}{prot_1 \sqsubseteq !\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot_2} \quad prot_1 \neq \texttt{end}
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-SEND-IN}} \\
\frac{P[\vec{t}/\vec{x}]}{!\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot \sqsubseteq !\langle v[\vec{t}/\vec{x}] \rangle. prot[\vec{t}/\vec{x}]}
\end{array}$$

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-RECV-OUT}} \\
\frac{\forall \vec{x}:\vec{\tau}. P \multimap (? \langle v \rangle. prot_1 \sqsubseteq prot_2)}{?\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot_1 \sqsubseteq prot_2} \quad prot_2 \neq \texttt{end}
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-RECV-IN}} \\
\frac{P[\vec{t}/\vec{x}]}{?\langle v[\vec{t}/\vec{x}] \rangle. prot[\vec{t}/\vec{x}] \sqsubseteq ?\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot}
\end{array}$$

Monotonicity and swapping:

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-SEND-MONO}} \\
\frac{\triangleright(prot_1 \sqsubseteq prot_2)}{!\langle v \rangle. prot_1 \sqsubseteq !\langle v \rangle. prot_2}
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-RECV-MONO}} \\
\frac{\triangleright(prot_1 \sqsubseteq prot_2)}{?\langle v \rangle. prot_1 \sqsubseteq ?\langle v \rangle. prot_2}
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-SWAP}} \\
?\langle v \rangle. !\langle w \rangle. prot \sqsubseteq !\langle w \rangle. ?\langle v \rangle. prot
\end{array}$$

Reflexivity and transitivity:

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-REFL}} \\
prot \sqsubseteq prot
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-TRANS}} \\
\frac{prot_1 \sqsubseteq prot_2 \quad prot_2 \sqsubseteq prot_3}{prot_1 \sqsubseteq prot_3}
\end{array}$$

Dual and append:

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-DUAL}} \\
\frac{prot_2 \sqsubseteq prot_1}{\overline{prot_1} \sqsubseteq \overline{prot_2}}
\end{array}
\qquad
\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-APPEND}} \\
\frac{prot_1 \sqsubseteq prot_2 \quad prot_3 \sqsubseteq prot_4}{prot_1 \cdot prot_3 \sqsubseteq prot_2 \cdot prot_4}
\end{array}$$

Channel ownership:

$$\begin{array}{c}
\text{\texttt{\(\sqsubseteq\)-CHAN-MONO}} \\
\frac{c \multimap prot_1 \quad prot_1 \sqsubseteq prot_2}{c \multimap prot_2}
\end{array}$$

Figure 11: The grammar and primitive rules of Actris 2.0 for subprotocols.

$\vec{x}:\vec{\tau}$ and resources P , *i.e.*, protocols of the shape $!\langle v \rangle. prot$ or $?\langle v \rangle. prot$, instead of those of the shape $!\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$ or $?\vec{x}:\vec{\tau} \langle v \rangle \{P\}. prot$. While this restriction might seem to make the rules more restrictive, the more general rules for monotonicity ($\sqsubseteq$ -SEND-MONO' and $\sqsubseteq$ -RECV-MONO') and swapping ($\sqsubseteq$ -SWAP') are derivable from these simpler rules. This is done using the rules for logical variable manipulation and resource transfer. Second, the actual rules for monotonicity have a later modality ($\triangleright$) in their premise. The later modality makes these rules stronger (by $\triangleright$ -INTRO we have that P entails $\triangleright P$), and thereby internalizes its coinductive nature into the Actris logic so LÖB induction can be used to prove subprotocol relations for recursive protocols (§6.4).

```

mapper_service f_v c :=
  branch c with
    left  => let x := recv c in
              let y := f_v x in
              send c y;
              mapper_service f_v c
    | right => ()
  end

mapper_client f_v l :=
  let c := start (mapper_service f_v) in
  let n := |l| in
  send_all c l;
  recvN c l n;
  select c right;

```

Figure 12: A mapper service whose verification relies on swapping (the code for the functions `send_all` and `recvN` has been elided).

The remaining rules in Figure 11 express that the subprotocol relation is reflexive ($\sqsubseteq$ -REFL) and transitive ($\sqsubseteq$ -TRANS), as well as that the dual operation is anti-monotone ($\sqsubseteq$ -DUAL) and the append operation is monotone ($\sqsubseteq$ -APPEND).

Let us consider the following subprotocol relation to provide some further insight into the expressivity of our rules, (where logical variables are omitted for simplicity):

$$!\langle v \rangle \{P\}. ?\langle w \rangle \{Q\}. prot \sqsubseteq !\langle v \rangle \{P * R\}. ?\langle w \rangle \{Q * R\}. prot$$

Here we extend the protocol $!\langle v \rangle \{P\}. ?\langle w \rangle \{Q\}. prot$ with a *frame* R . The proposition R describes resources that can be sent along with the originally expected resources P , and which are reacquired along with the resources Q that are sent back. We demonstrate the usefulness of this notion of framing at the protocol level in §6.3.

The above subprotocol relation mimics the frame rule of separation logic (HT-FRAME), which makes it possible to apply specifications while maintaining a *frame* of resources R :

$$\frac{\{P\} e \{w. Q\}}{\{P * R\} e \{w. Q * R\}}$$

The frame-like subprotocol relation is proven as follows:

$Q * R \multimap$	$?\langle w \rangle. prot \sqsubseteq ?\langle w \rangle \{Q * R\}. prot$	$\sqsubseteq$ -RECV-IN
$R \multimap$	$?\langle w \rangle \{Q\}. prot \sqsubseteq ?\langle w \rangle \{Q * R\}. prot$	$\sqsubseteq$ -RECV-OUT
$R \multimap$	$!\langle v \rangle. ?\langle w \rangle \{Q\}. prot \sqsubseteq !\langle v \rangle. ?\langle w \rangle \{Q * R\}. prot$	$\sqsubseteq$ -SEND-MONO, $\triangleright$ -INTRO
$P * R \multimap !\langle v \rangle \{P\}. ?\langle w \rangle \{Q\}. prot$	$\sqsubseteq !\langle v \rangle. ?\langle w \rangle \{Q * R\}. prot$	$\sqsubseteq$ -SEND-IN, $\sqsubseteq$ -TRANS
	$!\langle v \rangle \{P\}. ?\langle w \rangle \{Q\}. prot \sqsubseteq !\langle v \rangle \{P * R\}. ?\langle w \rangle \{Q * R\}. prot$	$\sqsubseteq$ -SEND-OUT

We use rule $\sqsubseteq$ -SEND-OUT to transfer P and the frame R out of the weaker protocol (*i.e.*, the send on the RHS), and then use rule $\sqsubseteq$ -SEND-IN to transfer P into the stronger protocol (*i.e.*, the send on the LHS), leaving us with a context in which we still own the frame R . We then use rule $\sqsubseteq$ -SEND-MONO to proceed with the receiving part of the protocol in a dual fashion—we use rule $\sqsubseteq$ -RECV-OUT to transfer out Q of the stronger protocol (*i.e.*, the receive on the LHS), and use rule $\sqsubseteq$ -RECV-IN to transfer Q and the frame R into the weaker protocol (*i.e.*, the receive on the RHS).

6.2. Swapping. Subprotocols make it possible to verify message-passing programs whose order of sends and receives does not match up w.r.t. duality. As an example of such a program, let us consider the mapper service and client in Figure 12. The service `mapper_service` $f_v c$ is a loop, which iteratively receives an element over channel endpoint c , maps the function f_v over that element, and sends the resulting value back. Conversely, the client `mapper_client` $f_v l$ sends all of the elements of the list l up front, and only requests the mapped results back once all elements have been sent. Since the service interleaves the sends and receives, while the client does not, the dependent separation protocols for the service and client cannot be dual of each other. However, the communication between the service and client is in fact safe as messages are buffered. We now show that using subprotocols we can prove that this is indeed the case. We define the protocol based on the communication where sends and receives are interleaved:

$$\text{mapper_prot } (I_T : T \rightarrow \text{Val} \rightarrow \text{iProp}) (I_U : U \rightarrow \text{Val} \rightarrow \text{iProp}) (f : T \rightarrow U) \triangleq \\ \mu(\text{rec} : \text{iProto}). (! (x : T) (v : \text{Val}) \langle v \rangle \{I_T x v\}. ?(w : \text{Val}) \langle w \rangle \{I_U (f x) w\}. \text{rec}) \oplus \text{end}$$

The protocol is parameterised by representation predicates I_T and I_U that relate HeapLang values to elements of type T and U in the Iris/Actris logic, and a function $f : T \rightarrow U$ in Iris/Actris that specifies the behaviour of the HeapLang function f_v . The connection between f and f_v is formalised as:

$$\text{f_spec } (I_T : T \rightarrow \text{Val} \rightarrow \text{iProp}) (I_U : U \rightarrow \text{Val} \rightarrow \text{iProp}) (f : T \rightarrow U) (f_v : \text{Val}) \triangleq \\ \forall x v. \{I_T x v\} f_v v \{w. I_U (f x) w\}$$

Since `mapper_prot` describes an interleaved sequence of transactions, `mapper_service` can be readily verified against the protocol `mapper_prot` using just the symbolic execution rules of Actris 1.0 as presented in § 4.2. However, to verify `mapper_client` against the protocol `mapper_prot`, we need to weaken the protocol using the rules for subprotocols of Actris 2.0. Given a list of n elements, the subprotocol relation (together with an intermediate step) that describes this weakening is:

$$\begin{aligned} & \text{mapper_prot } I_T I_U f \\ \sqsubseteq & ! \langle \text{left} \rangle. ! (x_1 : T) (v_1 : \text{Val}) \langle v_1 \rangle \{I_T x_1 v_1\}. & n \text{ times } \mu\text{-UNFOLD and} \\ & ?(y_1 : U) \langle y_1 \rangle \{I_U (f x_1) y_1\}. \dots & \text{weaken } \oplus \text{ into } ! \langle \text{left} \rangle \\ & ! \langle \text{left} \rangle. ! (x_n : T) (v_n : \text{Val}) \langle v_n \rangle \{I_T x_n v_n\}. \\ & ?(y_n : U) \langle y_n \rangle \{I_U (f x_n) y_n\}. \\ & \text{mapper_prot } I_T I_U f \\ \sqsubseteq & ! \langle \text{left} \rangle. ! (x_1 : T) (v_1 : \text{Val}) \langle v_1 \rangle \{I_T x_1 v_1\}. \dots & n \text{ times } \sqsubseteq\text{-SWAP}' \\ & ! \langle \text{left} \rangle. ! (x_n : T) (v_n : \text{Val}) \langle v_n \rangle \{I_T x_n v_n\}. \\ & ?(y_1 : U) \langle y_1 \rangle \{I_U (f x_1) y_1\}. \dots \\ & ?(y_n : U) \langle y_n \rangle \{I_U (f x_n) y_n\}. \\ & \text{mapper_prot } I_T I_U f \end{aligned}$$

Both steps are proven by induction on n . In the first step, we unfold the recursive protocol n times using $\mu\text{-UNFOLD}$, and use the derived rule $(\text{prot}_1 \oplus \text{prot}_2) \sqsubseteq ! \langle \text{left} \rangle. \text{prot}_1$ to weaken the choices to the left choice `left`. Recall from § 5.3 that $\oplus$ is defined in terms of the send protocol $!$. This allows us to prove the derived rule $(\text{prot}_1 \oplus \text{prot}_2) \sqsubseteq ! \langle \text{left} \rangle. \text{prot}_1$ using $\sqsubseteq\text{-SEND-OUT}$ and $\sqsubseteq\text{-SEND-IN}$. The second step involves swapping all sends ahead of the receives using the rule $\sqsubseteq\text{-SWAP}'$.

The weakened protocol that we have obtained follows the behaviour of the client, making its verification straightforward using Actris's rules for symbolic execution. Concretely, we

```

list_rev_service c :=      list_rev_client l :=
  let l := recv c in      let c := start list_rev_service in
  lreverse l; send c ()   send c l; recv c

```

Figure 13: A list reversing service (the code for the function `lreverse` has been elided).

prove the following specifications for the service and the client:

$$\begin{array}{ll}
\{f_spec\ I_T\ I_U\ f\ f_v\ *c \mapsto \overline{\text{mapper_prot}\ I_T\ I_U\ f \cdot prot}\} & \{f_spec\ I_T\ I_U\ f\ f_v\ * \ell \mapsto_{I_T}^{\text{list}} \vec{x}\} \\
\text{mapper_service}\ f_v\ c & \text{mapper_client}\ f_v\ \ell \\
\{c \mapsto prot\} & \{\ell \mapsto_{I_U}^{\text{list}} \text{map}\ f\ \vec{x}\}
\end{array}$$

6.3. Protocol compositionality. An essential feature of separation logic is the ability to compose specifications of different libraries, so that each library can be defined and verified once against its own specification, while being used in the context of slightly differing specifications and proofs of other libraries. To achieve a similar property for our dependent separation protocols we would similarly like to be able to compose compatible protocols.

A key ingredient that enables such compositionality in traditional separation logic is the frame rule (HT-FRAME). In § 6.1 we demonstrated how subprotocols allow for similar framing in our protocols. In this section we give a more detailed example of such framing in our protocols by considering the service `list_rev_service` c in Figure 13, which receives a linked list over channel endpoint c , reverses it, and sends it back over c .

To specify this service, we could use a protocol similar to the sorting service in § 5.1, defined in terms of the representation predicate $\ell \mapsto_{I_T}^{\text{list}} \vec{x}$ for linked lists:

$$\text{list_rev_prot}_{I_T} \triangleq !(\ell : \text{Loc})(\vec{x} : \text{List } T) \langle \ell \rangle \{ \ell \mapsto_{I_T}^{\text{list}} \vec{x} \}. ? \langle () \rangle \{ \ell \mapsto_{I_T}^{\text{list}} \text{reverse } \vec{x} \}. \text{end}$$

Although it is possible to verify the service against the protocol $\overline{\text{list_rev_prot}_{I_T}}$, this approach is not quite satisfactory. Unlike the sorting service, the reversal service does not access the list elements, but only changes the structure of the list. Hence, there is no need to keep track of the ownership of the elements through the predicate I_T . A self-contained and simpler protocol for this service would instead be the following:

$$\text{list_rev_prot} \triangleq !(\ell : \text{Loc})(\vec{v} : \text{List Val}) \langle \ell \rangle \{ \ell \mapsto \vec{v} \}. ? \langle () \rangle \{ \ell \mapsto \text{reverse } \vec{v} \}. \text{end}$$

Here, $\ell \mapsto \vec{v}$ is a version of the list representation predicate that does not keep track of the resources of the elements, but only describes the structure of the list. It is defined as:

$$\ell \mapsto \vec{v} \triangleq \begin{cases} \ell \mapsto \text{inl } () & \text{if } \vec{v} = \epsilon \\ \exists \ell_2. \ell \mapsto \text{inr } (v_1, \ell_2) * \ell_2 \mapsto \vec{v}_2 & \text{if } \vec{v} = [v_1] \cdot \vec{v}_2 \end{cases}$$

However, once we have verified the service against the simple protocol, the proof of a client might prefer to interact with the list reversal service through the general protocol $\text{list_rev_prot}_{I_T}$. Doing so can be achieved by proving the subprotocol relation $\text{list_rev_prot} \sqsubseteq \text{list_rev_prot}_{I_T}$. To prove this subprotocol relation, we first establish the following relation between the two versions of the list representation predicate:

$$\ell \mapsto_{I_T}^{\text{list}} \vec{x} ** (\exists \vec{v}. \ell \mapsto \vec{v} * \bigstar_{(x,v) \in (\vec{x}, \vec{v})} I_T\ x\ v) \quad (\text{LIST-REL})$$

Here, $\star_{(x,v) \in (\vec{x}, \vec{v})}$ is the pairwise iterated separation conjunction over two lists of equal length, and $\star\star$ is a bi-directional separation implication. The above result thus states that $\ell \xrightarrow{\text{list}}_{I_T} \vec{x}$ can be split into two parts, ownership of the links of the list $\ell \xrightarrow{\text{list}} \vec{v}$, and a range of interpretation predicates I_T for each element of the list, and *vice versa*. With this result at hand, the proof of the desired subprotocol relation is carried out as follows:

$$\begin{aligned}
& \text{list_rev_prot} \\
&= !(\ell : \text{Loc})(\vec{v} : \text{List Val}) \langle \ell \rangle \{ \ell \xrightarrow{\text{list}} \vec{v} \}. ? \langle () \rangle \{ \ell \xrightarrow{\text{list}} \text{reverse } \vec{v} \}. \text{end} \\
&\sqsubseteq !(\ell : \text{Loc})(\vec{v} : \text{List Val})(\vec{x} : \text{List T}) \langle \ell \rangle \left\{ \ell \xrightarrow{\text{list}} \vec{v} * \star_{(x,v) \in (\vec{x}, \vec{v})} . I_T x v \right\}. \\
&\quad ? \langle () \rangle \left\{ \ell \xrightarrow{\text{list}} (\text{reverse } \vec{v}) * \star_{(x,v) \in (\vec{x}, \vec{v})} . I_T x v \right\}. \text{end} \\
&\sqsubseteq !(\ell : \text{Loc})(\vec{x} : \text{List T}) \langle \ell \rangle \{ \ell \xrightarrow{\text{list}}_{I_T} \vec{x} \}. ? \langle () \rangle \{ \ell \xrightarrow{\text{list}}_{I_T} \text{reverse } \vec{x} \}. \text{end} \\
&= \text{list_rev_prot}_{I_T}
\end{aligned}$$

We first frame the range of interpretation predicates owned by the list $\star_{(x,v) \in (\vec{x}, \vec{v})} . I_T x v$, using an approach similar to the frame example in §6.1, and then use LIST-REL to combine it with $\ell \xrightarrow{\text{list}} \vec{v}$ and $\ell \xrightarrow{\text{list}} \text{reverse } \vec{v}$ for the sending and receiving step, to turn them into $\ell \xrightarrow{\text{list}}_{I_T} \vec{x}$ and $\ell \xrightarrow{\text{list}}_{I_T} \text{reverse } \vec{x}$, respectively. Note that the logical variable $\vec{v}$ is changed into $\vec{x}$, using the subprotocol rules for logical variable manipulation. With this subprotocol relation at hand, it is possible to prove the following specifications for the service and client:

$$\begin{array}{ll}
\{c \mapsto \overline{\text{list_rev_prot}} \cdot \text{prot}\} & \{ \ell \xrightarrow{\text{list}}_{I_T} \vec{x} \} \\
\text{list_rev_service } c & \text{list_rev_client } \ell \\
\{c \mapsto \text{prot}\} & \{ \ell \xrightarrow{\text{list}}_{I_T} \text{reverse } \vec{x} \}
\end{array}$$

6.4. Subprotocols and recursion. We conclude this section by showing how subprotocol relations involving recursive protocols can be proven using LÖB induction. Recall from §5 that the principle of LÖB induction is as follows:

$$\frac{\triangleright P \Rightarrow P}{P}$$

By letting P be $\text{prot}_1 \sqsubseteq \text{prot}_2$, we can prove $\text{prot}_1 \sqsubseteq \text{prot}_2$ using the induction hypothesis $\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2)$. The later modality ($\triangleright$) ensures that the induction hypothesis is not used immediately, but a monotonicity rule for send ($\sqsubseteq$ -SEND-MONO) or receive ($\sqsubseteq$ -RECV-MONO) is applied first. This is done typically after unfolding the recursion operator using μ -UNFOLD. The monotonicity rules $\sqsubseteq$ -SEND-MONO or $\sqsubseteq$ -RECV-MONO contain a later modality ($\triangleright$) in their premise, which makes it possible to strip off the later of the induction hypotheses (by rule $\triangleright$ -MONO for monotonicity of $\triangleright$).

Our approach for proving subprotocol relations using LÖB induction is similar to the approach of Brandt and Henglein [BH98] for proving subtyping relations for recursive types using coinduction. Brand and Henglein [BH98] however have a syntactic restriction on proofs to ensure that the induction hypothesis is not used immediately (*i.e.*, is used in a *contractive* fashion), while we use the later modality ($\triangleright$) of Iris to achieve that.

To demonstrate how our approach works, we prove $\text{prot}_1 \sqsubseteq \text{prot}_2$, where:

$$\begin{aligned}
\text{prot}_1 &\triangleq \mu(\text{rec} : \text{iProto}). (\text{list_rev_prot} \cdot \text{rec}) \oplus \text{end} \\
\text{prot}_2 &\triangleq \mu(\text{rec} : \text{iProto}). (\text{list_rev_prot}_{I_T} \cdot \text{rec}) \oplus \text{end}
\end{aligned}$$

Here, list_rev_prot and $\text{list_rev_prot}_{I_T}$ are the protocols from § 6.3, for which we have already proven $\text{list_rev_prot} \sqsubseteq \text{list_rev_prot}_{I_T}$. The proof of $\text{prot}_1 \sqsubseteq \text{prot}_2$ is as follows:

$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$\text{prot}_1 \sqsubseteq \text{prot}_2$	$\sqsubseteq\text{-APPEND}$
$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$\text{list_rev_prot} \cdot \text{prot}_1 \sqsubseteq \text{list_rev_prot}_{I_T} \cdot \text{prot}_2$	$\triangleright\text{-MONO}$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap$	$\triangleright(\text{list_rev_prot} \cdot \text{prot}_1 \sqsubseteq \text{list_rev_prot}_{I_T} \cdot \text{prot}_2)$	$\oplus\text{-MONO}$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap (\text{list_rev_prot} \cdot \text{prot}_1) \oplus \text{end} \sqsubseteq (\text{list_rev_prot}_{I_T} \cdot \text{prot}_2) \oplus \text{end}$		$\mu\text{-UNFOLD}$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap$	$\text{prot}_1 \sqsubseteq \text{prot}_2$	LÖB
$\text{prot}_1 \sqsubseteq \text{prot}_2$		

The proof starts with the LÖB rule, followed by unfolding the recursive types with $\mu\text{-UNFOLD}$. We then proceed with the following derived rule for monotonicity of selection ($\oplus$):

$$\frac{\oplus\text{-MONO} \quad \triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2 \wedge \text{prot}_3 \sqsubseteq \text{prot}_4)}{(\text{prot}_1 \oplus \text{prot}_3) \sqsubseteq (\text{prot}_2 \oplus \text{prot}_4)}$$

Due to the regular conjunction in the premise, the same resources can be used to prove both branches of $\oplus$. This is sound because only one branch of $\oplus$ will be chosen. The rule $\oplus\text{-MONO}$ follows from $\sqsubseteq\text{-SEND-MONO}$ as selection ($\oplus$) is defined in terms of send (!).

We continue the main proof with monotonicity of the later modality ($\triangleright\text{-MONO}$), which lets us strip off the later of the induction hypothesis $\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2)$. We then use $\sqsubseteq\text{-APPEND}$, along with $\text{list_rev_prot} \sqsubseteq \text{list_rev_prot}_{I_T}$, which we have proven in § 6.3. The remaining proof obligation $\text{prot}_1 \sqsubseteq \text{prot}_2$ follows from the induction hypothesis.

While the protocols in the prior examples are similar in structure, our approach scales to protocols for which that is not the case. For example, consider $\text{prot}_1 \sqsubseteq \text{prot}_2$, where:

$$\begin{aligned} \text{prot}_1 &\triangleq \mu(\text{rec} : \text{iProto}). ! (x : \mathbb{Z}) \langle x \rangle. ? \langle x + 2 \rangle. \text{rec} \\ \text{prot}_2 &\triangleq \mu(\text{rec} : \text{iProto}). ! (x : \mathbb{Z}) \langle x \rangle. ! (y : \mathbb{Z}) \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{rec} \end{aligned}$$

Intuitively, these protocols are related, as we can unfold the body of prot_1 twice, the body of prot_2 once, and swap the second receive over the first send. The proof is as follows:

$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$\text{prot}_1 \sqsubseteq \text{prot}_2$	$\sqsubseteq\text{-RECV-MONO}, \triangleright\text{-INTRO}$
$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_1 \sqsubseteq$ $? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_2$	$\sqsubseteq\text{-SEND-MONO}', \triangleright\text{-INTRO}$
$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$! y \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_1 \sqsubseteq$ $! y \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_2$	$\sqsubseteq\text{-SWAP}', \sqsubseteq\text{-TRANS}$
$\text{prot}_1 \sqsubseteq \text{prot}_2 \multimap$	$? \langle x + 2 \rangle. ! y \langle y \rangle. ? \langle y + 2 \rangle. \text{prot}_1 \sqsubseteq$ $! y \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_2$	$\triangleright\text{-MONO}$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap$	$\triangleright(? \langle x + 2 \rangle. ! y \langle y \rangle. ? \langle y + 2 \rangle. \text{prot}_1 \sqsubseteq$ $! y \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_2)$	$\sqsubseteq\text{-SEND-MONO}'$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap ! x \langle x \rangle. ? \langle x + 2 \rangle. ! y \langle y \rangle. ? \langle y + 2 \rangle. \text{prot}_1 \sqsubseteq$ $! x \langle x \rangle. ! y \langle y \rangle. ? \langle x + 2 \rangle. ? \langle y + 2 \rangle. \text{prot}_2$		$\mu\text{-UNFOLD}$
$\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2) \multimap$	$\text{prot}_1 \sqsubseteq \text{prot}_2$	LÖB
$\text{prot}_1 \sqsubseteq \text{prot}_2$		

Grammar:

$$t, u, P, Q, \text{prot} ::= \dots \mid \text{is_lock } lk \ P \mid \dots$$
Locks:

$\{R\} \text{new_lock } () \{lk.\text{is_lock } lk \ R\}$	(HT-NEW-LOCK)
$\{\text{is_lock } lk \ R\} \text{acquire } lk \ \{R\}$	(HT-ACQUIRE)
$\{\text{is_lock } lk \ R * R\} \text{release } lk \ \{\text{True}\}$	(HT-RELEASE)
$\text{is_lock } lk \ R \multimap \text{is_lock } lk \ R * \text{is_lock } lk \ R$	(LOCK-DUP)

Figure 14: The grammar and rules of locks in Iris.

After we use $\sqsubseteq$ -SEND-MONO' for the first time, we strip off the later of the induction hypothesis $\triangleright(\text{prot}_1 \sqsubseteq \text{prot}_2)$, using $\triangleright$ -MONO. Subsequently, when we use $\sqsubseteq$ -SEND-MONO' and $\sqsubseteq$ -RECV-MONO, there are no more lateres to strip. We therefore instead introduce the lateres using $\triangleright$ -INTRO before applying the appropriate subprotocol monotonicity rule.

7. MANIFEST SHARING VIA LOCKS

Since dependent separation protocols and the connective $c \multimap \text{prot}$ for ownership of protocols are first-class objects of the Actris logic, they can be used like any other logical connective. This means that protocols can be combined with any other mechanism that Actris inherits from Iris. In particular, they can be combined with Iris's generic invariant and ghost state mechanism, and can be used in combination with Iris's abstractions for reasoning about other concurrency connectives like locks, barriers, lock-free data structures, *etc.*

In this section we demonstrate how dependent separation protocols can be combined with lock-based concurrency. This combination allows us to prove functional correctness of programs that make use of the notion of *manifest sharing* [BP17; BTP19], where channel endpoints are shared between multiple parties. Instead of having to extend Actris, we make use of the locks and ghost state that Actris inherits from Iris. We present the basic idea with a simple introductory example of sharing a channel endpoint between two parties (§7.1). We then consider a more challenging example of a channel-based load-balancing mapper (§7.2).

7.1. Locks and ghost state. As presented in §2, HeapLang includes a lock library, with the operations `new_lock ()`, `acquire lk`, and `release lk`. The operations satisfy the separation logic specifications shown in Figure 14.

The specifications for locks make use of the representation predicate `is_lock lk R`, which expresses that a lock `lk` guards the resources described by the proposition `R`. When creating a new lock one has to give up ownership of `R`, and in turn, obtains the representation predicate `is_lock lk R` (HT-NEW-LOCK). The representation predicate can then be freely duplicated so it can be shared between multiple threads (LOCK-DUP). When entering a critical section using `acquire lk`, a thread gets exclusive ownership of `R` (HT-ACQUIRE), which has to be given up when releasing the lock using `release lk` (HT-RELEASE). The resources `R` that are protected by the lock are therefore invariant in-between any of the critical sections. The lock can only ever be acquired by one thread at a time, as `acquire lk` will loop until the

```

prog_lock := let c := start (λc. let lk := new_lock () in
                             fork {acquire lk; send c 21; release lk};
                             acquire lk; send c 21; release lk) in
recv c + recv c

```

Figure 15: A sample program that combines locks and channels to achieve manifest sharing.

$$\begin{aligned}
\text{True} &\Rightarrow \exists \gamma. \text{auth}_\gamma 0 && (\text{AUTH-INIT}) \\
\text{auth}_\gamma n &\Rightarrow \text{contrib}_\gamma * \text{auth}_\gamma (1 + n) && (\text{AUTH-ALLOC}) \\
\text{auth}_\gamma (1 + n) * \text{contrib}_\gamma &\Rightarrow \text{auth}_\gamma n && (\text{AUTH-DEALLOC}) \\
\text{auth}_\gamma n * \text{contrib}_\gamma &\multimap n > 0 && (\text{AUTH-CONTRIB-POS})
\end{aligned}$$

Figure 16: The authoritative contribution ghost theory.

lock is released. The HT-ACQUIRE rule reflects this, as the exclusive resources R are only obtained once the function terminates, *i.e.*, when the lock is available.

To show how locks can be used, consider the program `prog_lock` in Figure 15. This program uses a lock to share a channel endpoint between two threads, which each send the integer 21 to the main thread. The following dependent protocol specifies the expected interaction from the point of view of the main thread:

$$\text{lock_prot} \triangleq \mu(\text{rec} : \mathbb{N} \rightarrow \text{iProto}). \lambda n. \text{if } (n = 0) \text{ then end else } ?\langle 21 \rangle. \text{rec } (n - 1)$$

Here, n denotes the number of messages that should be exchanged. In the example program, n is initially 2. Since $c \mapsto \overline{\text{lock_prot}} \, n$ is an exclusive resource, we need a lock to share it between the threads that send 21. For this we will use the following lock invariant:

$$\text{is_lock } lk \, (\exists n. \text{auth}_\gamma n * c \mapsto \overline{\text{lock_prot}} \, n)$$

The natural number n is existentially quantified since it changes whenever a message is exchanged. To keep track of the number of exchanges that each thread is allowed to make we then need to tie the number n to some local resource. We achieve this by using Iris's *ghost theory* mechanism for creating user-defined ghost state [JSS⁺15; JKJ⁺18]. In particular, we define two logical connectives $\text{auth}_\gamma n$ and contrib_γ using Iris.⁴

The $\text{auth}_\gamma n$ fragment can be thought of as an authority that keeps track of the number of ongoing contributions n , while each contrib_γ is a token that witnesses that a contribution is still in progress. This intuition is made precise by the rules in Figure 16. The rule AUTH-INIT expresses that an authority $\text{auth}_\gamma 0$ can always be created, capturing that 0 contributions are initially in progress. A fresh ghost identifier γ is given, which is conceptually similar to how we obtain fresh locations for newly allocated references on the physical heap. Using the rules AUTH-ALLOC and AUTH-DEALLOC, we can allocate and deallocate contrib_γ tokens as long as the number n of ongoing contributions in $\text{auth}_\gamma n$ is updated accordingly. The rule AUTH-CONTRIB-POS expresses that ownership of a token contrib_γ implies that the count n of $\text{auth}_\gamma n$ must be positive.

⁴Defining a ghost theory in Iris involves picking an appropriate *resource algebra* with which one can define a set of abstract predicates (here $\text{auth}_\gamma n$ and contrib_γ). The details of resource algebras are beyond the scope of this paper and can be found in Jung *et al.* [JKJ⁺18].

Most of the rules in Figure 16 involve Iris's *view shift* connective $\Rightarrow$ for performing ghost updates. This is made precise by the structural rules VS-CSQ and VS-FRAME, which establish the connection between $\Rightarrow$ and Iris's Hoare triples:

$$\frac{\text{VS-CSQ} \quad P \Rightarrow P' \quad \{P'\} e \{v. Q'\} \quad \forall v. Q' \Rightarrow Q}{\{P\} e \{v. Q\}} \quad \frac{\text{VS-FRAME} \quad P \Rightarrow Q}{P * R \Rightarrow Q * R}$$

With the ghost theory in place, we can now prove suitable specifications for the program. The specification of the top-level program is shown on the right, while the left Hoare triple shows the auxiliary specification of both threads that send the integer 21:

$$\frac{\{ \text{contrib}_\gamma * \text{is_lock } lk \ (\exists n. \text{auth}_\gamma n * c \mapsto \overline{\text{lock_prot } n}) \} \quad \{ \text{True} \}}{\text{acquire } lk; \text{ send } c \ 21; \text{ release } lk \quad \text{prog_lock} \quad \{ v. v = 42 \}}$$

We use rule HT-NEW to assign protocol `lock_prot 2` to the channel. To establish the initial lock invariant, we use the rules AUTH-INIT and AUTH-ALLOC to create the authority `authγ 2` and two `contribγ` tokens. The `contribγ` tokens play a crucial role in the proofs of the sending threads to establish that the existentially quantified variable n is positive (using AUTH-CONTRIB-POS). Knowing $n > 0$, these threads can establish that the protocol `lock_prot  $n$` has not terminated yet (*i.e.*, is not **end**). This is needed to use the rule HT-SEND to prove the correctness of sending 21, and thereby advancing the protocol from `lock_prot  $n$` to `lock_prot ( $n - 1$ )`. Subsequently, the sending threads can deallocate the token `contribγ` (using AUTH-DEALLOC) to decrement the n of `authγ  $n$` accordingly to restore the lock invariant.

7.2. A channel-based load-balancing mapper. This section demonstrates a more interesting use of manifest sharing. We show how Actris can be used to verify functional correctness of a channel-based load-balancing mapper that maps the HeapLang function f_v over a list. Our channel-based mapper consists of one client that distributes the work, and a number of workers that perform the function f_v on individual elements of the list. To enable communication between the client and the workers, we make use of a single channel. One endpoint is used by the client to distribute the work between the workers, while the other endpoint is shared between all workers to request and return work from the client. The implementation of the workers `par_mapper_worker  $f_v$   $lk$   $c$` , which can be found in Figure 17, consists of a loop over three phases:

- (1) The worker notifies the client that it wants to perform work (using `select  $c$  left`), after which it is then notified (using `branch`) whether there is more work or all elements have been mapped. If there is more work, the worker receives an element x that needs to be mapped. Otherwise, the worker will terminate.
- (2) The worker maps the function f_v on x .
- (3) The worker notifies the client that it wants to send back a result (using `select  $c$  right`), and subsequently sends back the result y of mapping f_v on x .

The first and last phases are in a critical section guarded by a lock lk since they involve interaction over a shared channel endpoint. As the sharing behaviour is encapsulated by the worker, we omit the code of the client for brevity's sake.⁵

⁵The entire code is present in the accompanied Coq development [HBK21].

```

par_mapper_worker fv lk c :=
  acquire lk;
  select c left;
  branch c with
    right ⇒ release lk
  | left ⇒ let x := recv c in release lk;      (* acquire work *)
          let y := fv x in                    (* map it *)
          acquire lk;
          select c right; send c y;          (* send it back *)
          release lk;
  par_mapper_worker fv lk c
end

```

Figure 17: A worker of the channel-based mapper service.

$$\begin{aligned}
\text{True} &\Rightarrow \exists \gamma. \text{auth}_\gamma 0 \emptyset && (\text{AUTHM-INIT}) \\
\text{auth}_\gamma n X &\Rightarrow \text{auth}_\gamma (1 + n) X * \text{contrib}_\gamma \emptyset && (\text{AUTHM-ALLOC}) \\
\text{auth}_\gamma n X * \text{contrib}_\gamma \emptyset &\Rightarrow \text{auth}_\gamma (n - 1) X && (\text{AUTHM-DEALLOC}) \\
\text{auth}_\gamma n X * \text{contrib}_\gamma Y &\Rightarrow \text{auth}_\gamma n (X \uplus Z) * \text{contrib}_\gamma (Y \uplus Z) && (\text{AUTHM-ADD}) \\
Z \subseteq Y * \text{auth}_\gamma n X * \text{contrib}_\gamma Y &\Rightarrow \text{auth}_\gamma n (X \setminus Z) * \text{contrib}_\gamma (Y \setminus Z) && (\text{AUTHM-REMOVE}) \\
\text{auth}_\gamma n X * \text{contrib}_\gamma Y &\multimap n > 0 * Y \subseteq X && (\text{AUTHM-CONTRIB-AGREE}) \\
\text{auth}_\gamma 1 X * \text{contrib}_\gamma Y &\multimap Y = X && (\text{AUTHM-CONTRIB-AGREE1})
\end{aligned}$$

Figure 18: The authoritative contribution ghost theory extended with multisets.

A protocol that describes the interaction from the client's point of view is as follows:

```

par_mapper_prot (IT : T → Val → iProp) (IU : U → Val → iProp) (f : T → List U) ≜
  μ(rec : ℕ → MultiSet T → iProto). λn X.
    if n = 0 then end else
      (! (x : T) (v : Val) ⟨v⟩ {IT x v}. rec n (X ⊔ {x})) ⊕ rec (n - 1) X
      {(n=1) * (X=∅)} & {True}
      ?(x : T) (ℓ : Loc) ⟨ℓ⟩ {x ∈ X * ℓ  $\xrightarrow{\text{map}}$  IU (f x)}. rec n (X \ {x})

```

Similarly to `mapper_prot` from §6.2, the protocol is parameterised by representation predicates I_T and I_U , and a function $f : T \rightarrow \text{List } U$ in the Iris/Actris logic that will be related to f_v through a `f_spec` specification. Similar to the protocol `lock_prot` from §7.1, the protocol `par_mapper_prot` is indexed by the number of remaining workers n . On top of that, it carries a multiset X describing the values currently being processed by all the workers. The multiset X is used to make sure that the returned results are in fact the result of mapping the function f . The condition $(n = 1) * (X = \emptyset)$ on the branching operator ($\&$) expresses that the last worker may only request more work if there are no ongoing jobs.

To accommodate sharing of the channel endpoint between all workers using a lock invariant, we extend the authoritative contribution ghost theory from §7.1. We do this by adding multisets X and Y to the connectives $\text{auth}_\gamma n X$ and $\text{contrib}_\gamma Y$. These multisets

keep track of the values held by the workers. The rules for the ghost theory extended with multisets are shown in Figure 18. The rules `AUTHM-INIT`, `AUTHM-ALLOC` and `AUTHM-DEALLOC` are straightforward generalisations of the ones we have seen before. The new rules `AUTHM-ADD` and `AUTHM-REMOVE` determine that the multiset Y of `contribγ Y` can be updated as long as it is done in accordance with the multiset X of `authγ n X`. Finally, the `AUTHM-CONTRIB-AGREE` rule expresses that the multiset Y of `contribγ Y` must be a subset of the multiset X of `authγ n X`, while the stricter rule `AUTHM-CONTRIB-AGREE1` asserts equality between X and Y when only one contribution remains.

We then prove the following specifications of `par_mapper_worker` and a possible top-level client `par_mapper_client` that uses n workers to map f_v over the linked list ℓ :

$$\left\{ \begin{array}{l} \text{f_spec } I_T \ I_U \ f \ f_v * \text{contrib}_\gamma \emptyset * \\ \text{is_lock } lk \left(\frac{\exists n \ X. \text{auth}_\gamma n \ X *}{c \mapsto \text{par_mapper_prot } I_T \ I_U \ f \ n \ X} \right) \end{array} \right\} \quad \left\{ \begin{array}{l} \text{f_spec } I_T \ I_U \ f \ f_v * \\ 0 < n * \ell \xrightarrow{\text{list}}_{I_T} \vec{x} \end{array} \right\}$$

$$\text{par_mapper_worker } f_v \ lk \ c \quad \text{par_mapper_client } n \ f_v \ \ell$$

$$\{\text{True}\} \quad \{\exists \vec{y}. \vec{y} \equiv_p \text{flatMap } f \ \vec{x} * \ell \xrightarrow{\text{list}}_{I_U} \vec{y}\}$$

The lock invariant and specification of `par_mapper_worker` are similar to those used in the simple example in §7.1. The specification of `par_mapper_client` $n \ f_v \ \ell$ simply states that the resulting linked list contains a permutation of performing the map at the level of the logic. To specify that, we make use of `flatMap` : $(T \rightarrow \text{List } U) \rightarrow (\text{List } T \rightarrow \text{List } U)$, whose definition is standard.

The proof of the client involves allocating the channel with the protocol `par_mapper_prot`, with the initial number of workers n . Subsequently, we use the rules `AUTHM-INIT` and `AUTHM-ALLOC` to create the authority `authγ n ∅` and n tokens `contribγ ∅`, which allow us to establish the lock invariant and to distribute the tokens among the mappers. The proof of the mapper proceeds as usual. After acquiring the lock, the mapper obtains ownership of the lock invariant. Since the worker owns the token `contribγ ∅`, it knows that the number of remaining workers n is positive, which allows it to conclude that the protocol has not terminated (*i.e.*, is not **end**). After using the rules for channels, the rules `AUTHM-ADD` and `AUTHM-REMOVE` are used to update the authority, which is needed to reestablish the lock invariant so the lock can be released.

8. CASE STUDY: MAP-REDUCE

As a means of demonstrating the use of Actris for verifying more realistic programs, we present a proof of functional correctness of a simple channel-based load-balancing implementation of the map-reduce model by Dean and Ghemawat [DG04].

Since Actris is not concerned with distributed systems over networks, we consider a version of map-reduce that delegates the work over forked-off threads on a single machine. This means that we do not consider mechanics like handling the failure, restarting, and rescheduling of nodes that a version that operates on a network has to consider.

In order to implement and verify our map-reduce version we make use of the implementation and verification of the fine-grained channel-based merge sort algorithm (§5.6) and the channel-based load-balancing mapper (§7.2). As such, our map-reduce implementation is mostly a suitable client that glues together communication with these services. The purpose of this section is to give a high-level description of the implementation. The actual code and proofs can be found in the accompanied Coq development [HBK21].

8.1. A functional specification of map-reduce. The purpose of the map-reduce model is to transform an input set of type $\text{List } T$ into an output set of type $\text{List } V$ using two functions f (often called “map”) and g (often called “reduce”):

$$f : T \rightarrow \text{List } (K * U) \quad g : (K * \text{List } U) \rightarrow \text{List } V$$

An implementation of map-reduce performs the transformation in three steps:

- (1) First, the function f is applied to each element of the input set. This results in lists of key/value pairs which are then flattened using a `flatMap` operation (an operation that takes a list of lists and appends all nested lists):

$$\text{flatMap } f : \text{List } T \rightarrow \text{List } (K * U)$$

- (2) Second, the resulting lists of key/value pairs are grouped together by their key (this step is often called “shuffling”):

$$\text{group} : \text{List } (K * U) \rightarrow \text{List } (K * \text{List } U)$$

- (3) Finally, the grouped key/value pairs are passed on to the g function, after which the results are flattened to aggregate the results. This is done using a `flatMap` operation:

$$\text{flatMap } g : \text{List } (K * \text{List } U) \rightarrow \text{List } V$$

The complete functionality of map-reduce is equivalent to applying the following `map_reduce` function on the entire data set:

$$\text{map_reduce} : \text{List } T \rightarrow \text{List } V \triangleq (\text{flatMap } g) \circ \text{group} \circ (\text{flatMap } f)$$

A standard instance of map-reduce is counting word occurrences, where we let $T \triangleq K \triangleq \text{String}$ and $U \triangleq \mathbb{N}$ and $V \triangleq \text{String} * \mathbb{N}$ with:

$$\begin{aligned} f : \text{String} &\rightarrow \text{List } (\text{String} * \mathbb{N}) \triangleq \lambda x. [(x, 1)] \\ g : (\text{String} * \text{List } \mathbb{N}) &\rightarrow \text{List } (\text{String} * \mathbb{N}) \triangleq \lambda (k, \vec{n}). [(k, \sum_{i < |\vec{n}|} \vec{n}_i)] \end{aligned}$$

8.2. Implementation of map-reduce. The general distributed model of map-reduce is achieved by delegating the phases of mapping, shuffling, and reducing, over a number of worker nodes (*e.g.*, nodes of a cluster or individual CPUs). To perform the computation in a delegated way, there is some work involved in coordinating the jobs over these worker nodes, which is usually done as follows:

- (1) Split the input data into chunks and delegate these chunks to worker nodes, that each apply the “map” function f to their given data in parallel. We call these nodes the “mappers”.
- (2) Collect the complete set of mapped results and “shuffle” them, *i.e.*, group them by key. The grouping is commonly implemented using a parallel sorting algorithm.
- (3) Split the shuffled data into chunks and delegate these chunks to worker nodes that each apply the “reduce” function g to their given data in parallel. We call these nodes the “reducers”.
- (4) Collect and aggregate the complete set of result of the reducers.

Our variant of the map-reduce model is defined as a function `map_reducev n m fv gv ℓ` in `HeapLang`, which coordinates the work for performing map-reduce on a linked list ℓ between n mappers applying the `HeapLang` “map” function f_v , and m reducers applying the `HeapLang` “reduce” function g_v . To make the implementation more interesting, we prevent

storing intermediate values locally by forwarding/returning them immediately as they are available/requested. The global structure is as follows:

- (1) Start n instances of the load-balancing `par_mapper_worker` from §7, parameterised with the f_v function, acting as the mappers. Additionally start an instance of `sort_servicefg` from §5.6, parameterised by a concrete comparison function on the keys, corresponding to $\lambda(k_1, -) (k_2, -). k_1 < k_2$. Note that the type of keys are restricted to be integers for brevity's sake.
- (2) Perform a loop that handles communication with the mappers. If a mapper requests work, pop a value from the input list. If a mapper returns work, forward it to the sorting service. This process is repeated until all inputs have been mapped and forwarded.
- (3) Start m instances of the `par_mapper_worker`, parameterised by g_v , acting as the reducers.
- (4) Perform a loop that handles communication with the mappers. If a mapper requests work, group elements returned by the sort service. If a mapper returns work, aggregate the returned value in a the linked list. Grouped elements are created by requesting and aggregating elements from the sorter until the key changes.

The aggregated linked list then contains the fully mapped input set upon completion.

8.3. Functional correctness of map-reduce. The specification of the map-reduce program that we prove is as follows:

$$\begin{aligned} & \{0 < n * 0 < m * \text{f_spec } I_T \text{ } I_{\mathbb{Z}*U} f f_v * \text{f_spec } I_{\mathbb{Z}*List} U \text{ } I_V g g_v * \ell \xrightarrow{\text{list}}_{I_T} \vec{x}\} \\ & \quad \text{map_reduce}_v n m f_v g_v \ell \\ & \{ \exists \vec{z}. \vec{z} \equiv_p \text{map_reduce } f g \vec{x} * \ell \xrightarrow{\text{list}}_{I_V} \vec{z} \} \end{aligned}$$

The `f_spec` predicates (as introduced in §6.2) establish a connection between the functions f and g in Iris/Actris and the functions f_v and g_v in HeapLang. These make use of the various interpretation predicates I_T , $I_{\mathbb{Z}*U}$, $I_{\mathbb{Z}*List} U$, and I_V for the types in question. Lastly, the $\ell \xrightarrow{\text{list}}_{I_T} \vec{x}$ predicate determines that the input is a linked list of the initial type T . The postcondition asserts that the result $\vec{z}$ is a permutation of the original linked list $\vec{x}$ applied to the functional specification `map_reduce` of map-reduce from §8.1.

9. THE MODEL OF ACTRIS

We construct a model of Actris as a shallow embedding in the Iris framework [KJB⁺17; JSS⁺15; JKBD16; JKJ⁺18]. This means that the type `iProto` of dependent separation protocols, the subprotocol relation $prot_1 \sqsubseteq prot_2$, and the connective $c \multimap prot$ for the channel ownership, are definitions in Iris, and the Actris proof rules are lemmas about these definitions in Iris.

In this section we describe the relevant aspects of the model of Actris. We model the type `iProto` of dependent separation protocols as the solution of a recursive domain equation, and describe how the operations for dual and composition are defined (§9.1). We then define the subprotocol relation $prot_1 \sqsubseteq prot_2$ and prove its proof rules as lemmas (§9.2). To connect protocols to the endpoint channel buffers in the semantics we define the *protocol consistency relation*, which ensures that a pair of protocols is consistent with the messages in their associated buffers (§9.3). On top of the protocol consistency relation, we define the *Actris ghost theory* for dependent separation protocols (§9.4), which forms the key ingredient for defining the connective $c \multimap prot$ for channel ownership (§9.5) that links protocols to the

semantics of channels (§2.4). We then show how adequacy follows from the embedding in Iris (§9.6). Finally, we show how to solve the recursive domain equation for the type `iProto` of dependent separation protocols (§9.7).

9.1. The model of dependent separation protocols. To construct a model of dependent separation protocols, we first need to determine what they mean semantically. The challenging part involves the constructors $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$, whose (higher-order and impredicative) logical variables $\vec{x}:\vec{\tau}$ bind into the communicated value v , the transferred resources P , and the tail protocol $prot$. We model these constructors as predicates over the communicated value and the tail protocol. To describe the transferred resources P , we model these protocols as Iris predicates (functions to `iProp`). This gives rise to the following recursive domain equation:

$$\begin{aligned} \text{action} &::= \text{send} \mid \text{recv} \\ \text{iProto} &\cong 1 + (\text{action} \times (\text{Val} \rightarrow \blacktriangleright \text{iProto} \rightarrow \text{iProp})) \end{aligned}$$

The left part of the sum type (the unit type 1) indicates that the protocol has terminated, while the right part describes a message that is exchanged, expressed as an Iris predicate. Since the recursive occurrence of `iProto` in the predicate appears in negative position, we guard it using Iris's *type-level later* ($\blacktriangleright$) operator (whose only constructor is $\text{next} : T \rightarrow \blacktriangleright T$). The exact way the solution is constructed is detailed in §9.7. For now, we assume a solution exists, and define the dependent separation protocols constructors as:

$$\begin{aligned} \text{end} &\triangleq \text{inj}_1 () \\ !\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot &\triangleq \text{inj}_2 (\text{send}, \lambda w prot'. \exists \vec{x}:\vec{\tau}. (v = w) * P * (prot' = \text{next } prot)) \\ ?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot &\triangleq \text{inj}_2 (\text{recv}, \lambda w prot'. \exists \vec{x}:\vec{\tau}. (v = w) * P * (prot' = \text{next } prot)) \end{aligned}$$

The definitions of $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ make use of the (higher-order and impredicative) existential quantifiers of Iris to constrain the actual message w and tail $prot'$ so that they agree with the message v and tail $prot$ prescribed by the protocol.

Recursive protocols. Iris's guarded recursion operator $\mu x. t$ requires the recursion variable x to appear under a *contractive* term construct in t . Hence, to use Iris's recursion operator to construct recursive protocols, it is essential that the protocols $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ are contractive in the tail $prot$. To show why this is the case, let us first define what it means for a function $f : T \rightarrow U$ to be contractive:

$$\forall x, y. \triangleright (x = y) \Rightarrow f x = f y$$

Examples of contractive functions are the later modality $\triangleright : \text{iProp} \rightarrow \text{iProp}$ and the constructor $\text{next} : T \rightarrow \blacktriangleright T$. The protocols $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.prot$ are defined so that $prot$ appears below a next , and hence we can prove that they are contractive in $prot$.

Operations. With these definitions at hand, the dual $\overline{(-)}$ and append $(- \cdot -)$ operations are defined using Iris's guarded recursion operator ($\mu x. t$):

$$\overline{(-)} \triangleq \mu_{\text{rec}}. \lambda prot. \begin{cases} \text{inj}_1 () & \text{if } prot = \text{inj}_1 () \\ \text{inj}_2 (\overline{a}, \lambda w prot'. \exists prot''. & \text{if } prot = \text{inj}_2 (a, \Phi) \\ \quad \Phi w (\text{next } prot'') * & \\ \quad prot' = \text{next } (\text{rec } prot'')) & \end{cases}$$

$$(\cdot \cdot \text{prot}_2) \triangleq \mu \text{rec}. \lambda \text{prot}_1. \begin{cases} \text{prot}_2 & \text{if } \text{prot}_1 = \text{inj}_1 () \\ \text{inj}_2 (a, \lambda w \text{ prot}' . \exists \text{prot}'' . & \text{if } \text{prot}_1 = \text{inj}_2 (a, \Phi) \\ \quad \Phi w (\text{next } \text{prot}'') * & \\ \quad \text{prot}' = \text{next } (\text{rec } \text{prot}'')) & \end{cases}$$

In the above definitions, we let $\overline{\text{send}} \triangleq \text{recv}$ and $\overline{\text{recv}} = \text{send}$.

The base cases of both definitions are as expected. In the recursive cases, we construct a new predicate, given the original predicate Φ . In these new predicates, we quantify over an original tail protocol prot'' such that $\Phi w (\text{next } \text{prot}'')$ holds, and unify the new tail protocol prot' with the result of the recursive call $\text{rec } \text{prot}''$.

The equational rules for dual $\overline{(\cdot)}$ and append $(\cdot \cdot \cdot)$ from Figure 5 are proven as lemmas in Iris using LÖB induction. This is possible as the recursive call $\text{rec } \text{prot}''$ appears below a next constructor—since the next constructor is contractive, we can strip-off the later from the induction hypothesis when proving the equality for the tail.

Difference from the conference version. In the conference version of this paper [HBK20], we described two versions of the recursive domain equation for dependent separation protocols: an “ideal” version (as used in this paper), where iProto appears in negative position, and an “alternative” version, where iProto appears in positive position. At that time, we were unable to construct a solution of the “ideal” version, so we used the “alternative” version. In §9.7 we show how we are now able to solve the “ideal” version.

In the conference version of this paper, the proposition P appeared under a later modality in the definitions of the protocols $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$, making these protocols contractive in P . This choice was motivated by the ability to construct recursive protocols like $\mu \text{rec}. ! (c : \text{Val}) \langle c \rangle \{c \mapsto \text{prot}\}. \text{prot}'$, where the payload refers to the recursion variable rec . In the current version (without the later modality) we can still construct such protocols, because $c \mapsto \text{prot}$ is contractive in prot . We removed the later modality because it is incompatible with the rules $\sqsubseteq\text{-SEND-OUT}$ and $\sqsubseteq\text{-RECV-OUT}$ for subprotocols.

9.2. The model of the subprotocol relation. We now model the subprotocol relation $\text{prot}_1 \sqsubseteq \text{prot}_2$ from §6. For legibility, we present it in the style of an inference system through its constructors, whereas it is formally defined using Iris’s guarded recursion operator $(\mu x. t)$:

$$\begin{array}{c} \text{inj}_1 () \sqsubseteq \text{inj}_1 () \\[10pt] \frac{\forall v, \text{prot}_2. \Phi_2 v (\text{next } \text{prot}_2) \multimap \exists \text{prot}_1. \Phi_1 v (\text{next } \text{prot}_1) * \triangleright (\text{prot}_1 \sqsubseteq \text{prot}_2)}{\text{inj}_2 (\text{send}, \Phi_1) \sqsubseteq \text{inj}_2 (\text{send}, \Phi_2)} \quad \frac{\forall v, \text{prot}_1. \Phi_1 v (\text{next } \text{prot}_1) \multimap \exists \text{prot}_2. \Phi_2 v (\text{next } \text{prot}_2) * \triangleright (\text{prot}_1 \sqsubseteq \text{prot}_2)}{\text{inj}_2 (\text{recv}, \Phi_1) \sqsubseteq \text{inj}_2 (\text{recv}, \Phi_2)} \\[10pt] \frac{\forall v_1, v_2, \text{prot}_1, \text{prot}_2. (\Phi_1 v_1 (\text{next } \text{prot}_1) * \Phi_2 v_2 (\text{next } \text{prot}_2)) \multimap \exists \text{prot}. \triangleright (\text{prot}_1 \sqsubseteq ! \langle v_2 \rangle. \text{prot}) * \triangleright (? \langle v_1 \rangle. \text{prot} \sqsubseteq \text{prot}_2)}{\text{inj}_2 (\text{recv}, \Phi_1) \sqsubseteq \text{inj}_2 (\text{send}, \Phi_2)} \end{array}$$

To be a well-formed guarded recursion definition, every recursive occurrence of $\sqsubseteq$ is guarded by the later modality $\triangleright$. Aside from later being required for well-formedness, these lateres make it possible to reason about the subprotocol relation using LÖB induction; both to prove the subprotocol rules from Figure 11 as lemmas, and for Actris users to reason about

recursive protocols as shown in § 6.4. The relation is defined in a syntax directed fashion (*i.e.*, there are no overlapping rules), and therefore all constructors need to be defined so that they are closed under monotonicity and transitivity.

The first constructor states that terminating protocols (**end** $\triangleq$ **inj**₁ ()) are related. The other constructors concern the protocols $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}$, which are modelled as **inj**₂ (**send**, Φ) and **inj**₂ (**recv**, Φ), where $\Phi : \text{Val} \rightarrow \blacktriangleright \text{iProto} \rightarrow \text{iProp}$ is a predicate over the communicated value and tail protocol. While the actual constructors are somewhat intimidating because they are defined in terms of these predicates in the model, they essentially correspond to the following high-level versions:

$$\frac{\forall \vec{y} : \vec{\sigma}. P_2 \multimap \exists \vec{x} : \vec{\tau}. (v_1 = v_2) * P_1 * \triangleright (\text{prot}_1 \sqsubseteq \text{prot}_2)}{! \vec{x} : \vec{\tau} \langle v_1 \rangle \{P_1\}. \text{prot}_1 \sqsubseteq ! \vec{y} : \vec{\sigma} \langle v_2 \rangle \{P_2\}. \text{prot}_2}$$

$$\frac{\forall \vec{x} : \vec{\tau}. P_1 \multimap \exists \vec{y} : \vec{\sigma}. (v_1 = v_2) * P_2 * \triangleright (\text{prot}_1 \sqsubseteq \text{prot}_2)}{? \vec{x} : \vec{\tau} \langle v_1 \rangle \{P_1\}. \text{prot}_1 \sqsubseteq ? \vec{y} : \vec{\sigma} \langle v_2 \rangle \{P_2\}. \text{prot}_2}$$

$$\frac{\forall \vec{x} : \vec{\tau}, \vec{y} : \vec{\sigma}. (P_1 * P_2) \multimap \exists \text{prot}. \triangleright (\text{prot}_1 \sqsubseteq ! \langle v_2 \rangle. \text{prot}) * \triangleright (? \langle v_1 \rangle. \text{prot} \sqsubseteq \text{prot}_2)}{? \vec{x} : \vec{\tau} \langle v_1 \rangle \{P_1\}. \text{prot}_1 \sqsubseteq ! \vec{y} : \vec{\sigma} \langle v_2 \rangle \{P_2\}. \text{prot}_2}$$

To obtain syntax directed rules, the first rule combines $\sqsubseteq$ -SEND-OUT, $\sqsubseteq$ -SEND-IN, and $\sqsubseteq$ -SEND-MONO, and dually, the second rule combines $\sqsubseteq$ -RECV-OUT, $\sqsubseteq$ -RECV-IN, and $\sqsubseteq$ -RECV-MONO. The third rule combines $\sqsubseteq$ -RECV-OUT, $\sqsubseteq$ -SEND-OUT and $\sqsubseteq$ -SWAP and bakes in transitivity, instead of asserting that prot_1 and prot_2 are equal to $! \langle v_2 \rangle. \text{prot}$ and $? \langle v_1 \rangle. \text{prot}$, respectively.

The rules from the beginning of this section are defined by generalising the high-level rules to arbitrary predicates. For example, rule **inj**₂ (**send**, Φ_1) $\sqsubseteq$ **inj**₂ (**send**, Φ_2) requires that for any value v and tail protocol prot_2 that are allowed by the predicate Φ_2 , there is a stronger tail protocol prot_1 (*i.e.*, where $\text{prot}_1 \sqsubseteq \text{prot}_2$), so that the same value v and stronger tail protocol prot_1 are allowed by the predicate Φ_1 .

The rules in Figure 11 on page 28 are proven as lemmas. Those for logical variable and resource manipulation ($\sqsubseteq$ -SEND-OUT, $\sqsubseteq$ -SEND-IN, $\sqsubseteq$ -RECV-OUT and $\sqsubseteq$ -RECV-IN), monotonicity ($\sqsubseteq$ -SEND-MONO and $\sqsubseteq$ -RECV-MONO), and swapping ($\sqsubseteq$ -SWAP) follow almost immediately from the definition, whereas those for reflexivity ($\sqsubseteq$ -REFL), transitivity ($\sqsubseteq$ -TRANS), and the dual and append operator ($\sqsubseteq$ -DUAL and $\sqsubseteq$ -APPEND) are proven using LÖB induction.

9.3. Protocol consistency. To connect dependent separation protocols to the semantics of channels in § 9.5, we define the *protocol consistency relation* **prot_consistent** $\vec{v}_1 \vec{v}_2 \text{prot}_1 \text{prot}_2$, which expresses that protocols prot_1 and prot_2 are *consistent* w.r.t. channel buffers containing values $\vec{v}_1$ and $\vec{v}_2$. The consistency relation is defined as:

$$\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{prot}_1 \text{prot}_2 \triangleq \exists \text{prot}. \\ (\langle ? \langle \vec{v}_{2.1} \rangle \dots ? \langle \vec{v}_{2.|\vec{v}_2|} \rangle. \text{prot} \sqsubseteq \text{prot}_1 \rangle * \langle ? \langle \vec{v}_{1.1} \rangle \dots ? \langle \vec{v}_{1.|\vec{v}_1|} \rangle. \overline{\text{prot}} \sqsubseteq \text{prot}_2 \rangle)$$

Intuitively, **prot_consistent** $\vec{v}_1 \vec{v}_2 \text{prot}_1 \text{prot}_2$ ensures that for all messages $\vec{v}_1 = \vec{v}_{1.1} \dots \vec{v}_{1.|\vec{v}_1|}$ in transit from the endpoint described by prot_1 to the endpoint described by prot_2 , the protocol prot_2 is expecting to receive these message in order (and *vice versa* for $\vec{v}_2$), after which the remaining protocols prot and $\overline{\text{prot}}$ are dual. To account for weakening we close the consistency relation under subprotocols (by using $\sqsubseteq$ instead of equality). Closure under the subprotocol relation additionally implicitly captures ownership of the quantifiers and

$$\begin{aligned}
\text{True} &\Rightarrow \exists \gamma. (\gamma \mapsto_{\bullet} \text{prot}) * (\gamma \mapsto_{\circ} \text{prot}) && (\text{HO-GHOST-ALLOC}) \\
(\gamma \mapsto_{\bullet} \text{prot}) * (\gamma \mapsto_{\circ} \text{prot}') &\Rightarrow (\gamma \mapsto_{\bullet} \text{prot}'') * (\gamma \mapsto_{\circ} \text{prot}'') && (\text{HO-GHOST-UPDATE}) \\
(\gamma \mapsto_{\bullet} \text{prot}) * (\gamma \mapsto_{\circ} \text{prot}') &\multimap \triangleright (\text{prot} = \text{prot}') && (\text{HO-GHOST-AGREE})
\end{aligned}$$

Figure 19: Ghost theory for higher-order ghost variables in Iris.

resources associated with the messages $\vec{v}_1$ and $\vec{v}_2$. That is, since the subprotocol relations relate the protocol arguments prot_1 and prot_2 with protocols that specify no quantifiers or resources. More precisely, by the definition of the subprotocol relation (shown in § 9.2), a relation such as $? \langle v \rangle. \text{prot}_1 \sqsubseteq ?(\vec{x}:\vec{\tau}) \langle v \rangle \{P\}. \text{prot}_2$ is equivalent to a separation implication of the form $\text{True} \multimap \exists \vec{x}:\vec{\tau}. P * \triangleright \text{prot}_1 \sqsubseteq \text{prot}_2$, where the obligation True is trivial, meaning that it implicitly asserts ownership of P .

Finally, closure under the subprotocol relation gives that $\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2$ and $\text{prot}_1 \sqsubseteq \text{prot}'_1$ implies $\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}'_1 \text{ prot}_2$, and ensures that the consistency relation enjoys the following rules corresponding to creating a channel, sending a message, and receiving a message:

$$\begin{aligned}
&\text{prot_consistent } \epsilon \in \text{prot } \overline{\text{prot}} \\
&(\text{prot_consistent } \vec{v}_1 \vec{v}_2 (!\vec{x}:\vec{\tau} \langle v \rangle \{P\}. \text{prot}_1) \text{ prot}_2 * P[\vec{t}/\vec{x}]) \multimap \\
&\quad \triangleright^{|\vec{v}_2|} (\text{prot_consistent } (\vec{v}_1 \cdot [v[\vec{t}/\vec{x}]]) \vec{v}_2 \text{ prot}_1 \text{ prot}_2) \\
&\text{prot_consistent } \vec{v}_1 ([w] \cdot \vec{v}_2) (? \vec{x}:\vec{\tau} \langle v \rangle \{P\}. \text{prot}_1) \text{ prot}_2 \multimap \\
&\quad \exists \vec{y}. (w = v[\vec{y}/\vec{x}]) * P[\vec{y}/\vec{x}] * \triangleright (\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2)
\end{aligned}$$

The first rule states that dual protocols are consistent w.r.t. a pair of empty buffers. The second rule states that a protocol $!\vec{x}:\vec{\tau} \langle v \rangle \{P\}. \text{prot}_1$ can be advanced to prot_1 by giving up ownership of $P[\vec{t}/\vec{x}]$ and enqueueing the value $v[\vec{t}/\vec{x}]$ in the buffer $\vec{v}_1$. Dually, the third rule states that given a protocol $? \vec{x}:\vec{\tau} \langle v \rangle \{P\}. \text{prot}_1$ and a buffer that contains value w as its head, we learn that w is equal to $v[\vec{y}/\vec{x}]$, and that we can obtain ownership of $P[\vec{y}/\vec{x}]$ by advancing the protocol to prot_1 and dequeuing the value w from the buffer. Since the relation is symmetric, *i.e.*, if $\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2$ then $\text{prot_consistent } \vec{v}_2 \vec{v}_1 \text{ prot}_2 \text{ prot}_1$, we obtain similar rules for the protocol prot_2 on the right-hand side.

The last two rules are proven by case analysis on the subprotocol relation ($\sqsubseteq$) in the assumption. Since the subprotocol relation ($\sqsubseteq$) is defined using guarded recursion, we obtain a later modality ($\triangleright$) for each case analysis. To prove the first of the rules, we need to perform a number of case analyses equal to the size of the buffer $\vec{v}_2$, whereas for the second rule we need to perform just a single case analysis. These later modalities are eliminated through the `skipN` operation in the `send` operation, see § 9.5 for further discussion.

9.4. The Actris ghost theory. To provide a general interface for making Actris’s reasoning principles independent of HeapLang, we employ a standard ghost theory approach in Iris to compartmentalise channel ownership. In § 9.5 we define the connective $c \mapsto \text{prot}$ for channel endpoint ownership that links the ghost theory to the buffers of our implementation of channels in HeapLang.

The Actris ghost theory is similar in its interface to the ghost theory for contributions that we used in § 7. We define three new logical connectives—an authority $\text{prot_ctx } \chi \vec{v}_1 \vec{v}_2$,

$$\begin{aligned}
& \text{True} \Rightarrow \exists \chi. \text{prot_ctx } \chi \in \epsilon * \text{prot_own}_l \chi \text{ prot} * \text{prot_own}_r \chi \overline{\text{prot}} & (\text{PROTO-ALLOC}) \\
& \text{prot_ctx } \chi \vec{v}_1 \vec{v}_2 * \text{prot_own}_l \chi (!\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}) * P[\vec{t}/\vec{x}] \Rightarrow & (\text{PROTO-SEND-L}) \\
& \quad (\triangleright^{|\vec{v}_2|} \text{prot_ctx } \chi (\vec{v}_1 \cdot [v[\vec{t}/\vec{x}]]) \vec{v}_2) * \text{prot_own}_l \chi (\text{prot}[\vec{t}/\vec{x}]) \\
& \text{prot_ctx } \chi \vec{v}_1 \vec{v}_2 * \text{prot_own}_r \chi (!\vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}) * P[\vec{t}/\vec{x}] \Rightarrow & (\text{PROTO-SEND-R}) \\
& \quad (\triangleright^{|\vec{v}_1|} \text{prot_ctx } \chi \vec{v}_1 (\vec{v}_2 \cdot [v[\vec{t}/\vec{x}]]) * \text{prot_own}_r \chi (\text{prot}[\vec{t}/\vec{x}]) \\
& \text{prot_ctx } \chi \vec{v}_1 ([w] \cdot \vec{v}_2) * \text{prot_own}_l \chi (? \vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}) \Rightarrow & (\text{PROTO-RECV-L}) \\
& \quad \triangleright \exists \vec{y}. w = v[\vec{y}/\vec{x}] * P[\vec{y}/\vec{x}] * \text{prot_ctx } \chi \vec{v}_1 \vec{v}_2 * \text{prot_own}_l \chi (\text{prot}[\vec{y}/\vec{x}]) \\
& \text{prot_ctx } \chi ([w] \cdot \vec{v}_1) \vec{v}_2 * \text{prot_own}_r \chi (? \vec{x}:\vec{\tau}\langle v \rangle\{P\}. \text{prot}) \Rightarrow & (\text{PROTO-RECV-R}) \\
& \quad \triangleright \exists \vec{y}. w = v[\vec{y}/\vec{x}] * P[\vec{y}/\vec{x}] * \text{prot_ctx } \chi \vec{v}_1 \vec{v}_2 * \text{prot_own}_r \chi (\text{prot}[\vec{y}/\vec{x}]) \\
& \text{prot_own}_l \chi \text{ prot} * \text{prot} \sqsubseteq \text{prot}' \multimap \text{prot_own}_l \chi \text{ prot}' & (\text{PROTO-}\sqsubseteq\text{-L}) \\
& \text{prot_own}_r \chi \text{ prot} * \text{prot} \sqsubseteq \text{prot}' \multimap \text{prot_own}_r \chi \text{ prot}' & (\text{PROTO-}\sqsubseteq\text{-R})
\end{aligned}$$

Figure 20: The Actris ghost theory.

and tokens $\text{prot_own}_l \chi \text{ prot}_l$ and $\text{prot_own}_r \chi \text{ prot}_r$ —and prove rules about how they can be allocated, updated, and used. Similar to prior ghost theories, the identifier χ associates the connectives to each other. The $\text{prot_ctx } \chi \vec{v}_1 \vec{v}_2$ connective can be thought of as an authority that governs the global state of the buffers $\vec{v}_1$ and $\vec{v}_2$. The tokens $\text{prot_own}_l \chi \text{ prot}_l$ and $\text{prot_own}_r \chi \text{ prot}_r$ provide local views of the buffers state in terms of the protocols prot_l and prot_r . As we will see in § 9.5, the authority can be shared using a lock, while the tokens provide unique ownership of each endpoint.

To define the connectives of the Actris ghost theory we use Iris’s existing ghost theory for higher-order ghost variables, revolving around the two connectives $\gamma \mapsto_{\bullet} \text{prot}$ and $\gamma \mapsto_{\circ} \text{prot}'$, which we call the inner and outer fragments, respectively. As before, the γ is the ghost identifier that associates the connectives. The fragments can be thought of as two pieces of a single variable, which can only be updated in the presence of both fragments. As a result, we know that inner and outer fragment with the same ghost identifier γ always point to the same protocol prot . This is made precise by the rules as shown in Figure 19. In particular, higher-order ghost variables are allocated in pairs $\gamma \mapsto_{\bullet} \text{prot}$ and $\gamma \mapsto_{\circ} \text{prot}$ for an identical protocol prot (HO-GHOST-ALLOC), and they can only be updated together (HO-GHOST-UPDATE). This means that they will always hold the same protocol (HO-GHOST-AGREE). The subtle part of the higher-order ghost variables is that they involve ownership of a protocol of type iProto , which is defined in terms of Iris propositions iProp . Due to the dependency on iProp (which is covered in detail in § 9.1 and 9.7) the rule HO-GHOST-AGREE only gives the equality between the protocols under a later modality ($\triangleright$).

With Iris’s higher-order ghost variables at hand, we can define the Actris ghost theory connectives as:

$$\begin{aligned}
& \text{prot_ctx } (\gamma_1, \gamma_2) \vec{v}_1 \vec{v}_2 \triangleq \exists \text{prot}_1, \text{prot}_2. \gamma_1 \mapsto_{\bullet} \text{prot}_1 * \gamma_2 \mapsto_{\bullet} \text{prot}_2 * & \\
& \quad \triangleright \text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2 \\
& \text{prot_own}_l (\gamma_1, \gamma_2) \text{ prot}_l \triangleq \exists \text{prot}'_l. \gamma_1 \mapsto_{\circ} \text{prot}'_l * \triangleright (\text{prot}'_l \sqsubseteq \text{prot}_l) \\
& \text{prot_own}_r (\gamma_1, \gamma_2) \text{ prot}_r \triangleq \exists \text{prot}'_r. \gamma_2 \mapsto_{\circ} \text{prot}'_r * \triangleright (\text{prot}'_r \sqsubseteq \text{prot}_r)
\end{aligned}$$

Since we use two higher-order ghost variables, our identifiers $\chi ::= (\gamma_1, \gamma_2)$ are pairs of Iris ghost identifiers. The authority $\text{prot_ctx } (\gamma_1, \gamma_2) \vec{v}_1 \vec{v}_2$ asserts ownership of the inner fragments of the higher-order ghost variables $\gamma_1 \mapsto_{\bullet} \text{prot}_1$ and $\gamma_2 \mapsto_{\bullet} \text{prot}_2$ for some protocols prot_1 and prot_2 . It then asserts that the buffers $\vec{v}_1$ and $\vec{v}_2$ are consistent with respect to those protocols prot_1 and prot_2 (via $\text{prot_consistent } \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2$). The tokens $\text{prot_own}_l (\gamma_1, \gamma_2) \text{ prot}_l$ and $\text{prot_own}_r (\gamma_1, \gamma_2) \text{ prot}_r$ respectively assert ownership of the outer higher-order ghost variable fragments $\gamma_1 \mapsto_{\circ} \text{prot}'_l$ and $\gamma_2 \mapsto_{\circ} \text{prot}'_r$. Here prot'_l and prot'_r are protocols that are weaker than the protocol arguments prot_l and prot_r (via $\text{prot}'_l \sqsubseteq \text{prot}_l$ and $\text{prot}'_r \sqsubseteq \text{prot}_r$). The explicit weakening under the subprotocol relation may seem redundant, as weakening is already accounted for in prot_consistent . However, it allows us to weaken the protocols of the tokens without the presence of the authority as shown momentarily. The later modality ($\triangleright$) makes sure that $\text{prot_own}_l (\gamma_1, \gamma_2) \text{ prot}$ and $\text{prot_own}_r (\gamma_1, \gamma_2) \text{ prot}$ are contractive in prot .

With the definitions of the ghost theory connectives at hand, we prove the rules of the ghost theory presented in Figure 20. The rule PROTO-ALLOC corresponds to allocation of a buffer pair, the rules PROTO-SEND-L and PROTO-SEND-R correspond to sending a message, and the rules PROTO-RECV-L and PROTO-RECV-R correspond to receiving a message. Finally, the rules $\text{PROTO-}\sqsubseteq\text{-L}$ and $\text{PROTO-}\sqsubseteq\text{-R}$ captures that we can weaken the protocols of the tokens without the presence of the authority. The rules of Figure 20 are proven through a combination of the rules for higher-order ghost state from Figure 19, and the rules for the protocol consistency relation prot_consistent from §9.3.

9.5. The model of channel ownership. To link the physical contents of the bidirectional channel c to the Actris ghost theory we define the channel ownership connective as follows:

$$c \mapsto \text{prot} \triangleq \exists \chi, l, r, lk. \left(\begin{array}{l} (c = (l, r, lk) * \text{prot_own}_l \chi \text{ prot}) \vee \\ (c = (r, l, lk) * \text{prot_own}_r \chi \text{ prot}) \end{array} \right) * \\ \text{is_lock } lk \ (\exists \vec{v}_1 \vec{v}_2. l \xrightarrow{\text{list}} \vec{v}_1 * r \xrightarrow{\text{list}} \vec{v}_2 * \text{prot_ctx } \chi \vec{v}_1 \vec{v}_2)$$

The predicate states that the referenced channel endpoint c is either the left (l, r, lk) or the right (r, l, lk) side of a channel, and that we have exclusive ownership of the ghost token $\text{prot_own}_l \chi \text{ prot}$ or $\text{prot_own}_r \chi \text{ prot}$ for the corresponding side. Iris's lock representation predicate is_lock (previously presented in §7) is used to make sharing of the buffers possible. The lock invariant is governed by lock lk , and carries the ownership $l \xrightarrow{\text{list}} \vec{v}_1$ and $r \xrightarrow{\text{list}} \vec{v}_2$ of the mutable linked lists containing the channel buffers, as well as $\text{prot_ctx } \chi \vec{v}_1 \vec{v}_2$, which asserts protocol consistency of the buffers with respect to the protocols.

With the definition of the channel endpoint ownership along with the ghost theory and lock rules we then prove the channel rules HT-NEW , HT-SEND and HT-RECV from Figure 5. The proofs are carried out through symbolic execution to the point where the critical section is entered, after which the rules of the Actris ghost theory (Figure 20) are used to allocate or update the ghost state appropriately so that it matches the physical channel buffers.

The need for skip instructions. The rules PROTO-SEND-L and PROTO-SEND-R from Figure 20 contain a number of later modalities ($\triangleright$) proportional to the other endpoint's buffer. As explained in §9.3 these later modalities are the consequence of having to perform a number of case analyses on the subprotocol relation, which is defined using guarded recursion, and thus contains a later modality for each recursive unfolding.

To eliminate these later modalities, we instrument the code of the `send` function with the `skipN (llength r)` instruction, which performs a number of skips equal to the size of

the other endpoint's buffer r . The `skipN` instruction has the following specification:

$$\{\triangleright^n P\} \text{skipN } n \{P\}$$

Instrumentation with skip instructions is used often in work on step-indexing, see *e.g.*, [SSB16; GST⁺20]. Instrumentation is needed because current step-indexed logics like Iris unify physical/program steps and logical steps, *i.e.*, for each physical/program step at most one later can be eliminated from the hypotheses. In recent work by Svendsen *et al.* [SSB16], Matsushita and Jourdan [MJ20], and Spies *et al.* [SGG⁺21] more liberal versions of step-indexing have been proposed, but none of these versions of step-indexing have been integrated into the main Coq development of Iris and HeapLang.

9.6. Adequacy of Actris. Having constructed the model of Actris in Iris, we obtain the following main result, as first presented in §3.4:

Theorem 9.1 (Adequacy of Actris). *Let $\varphi \in \text{Val} \rightarrow \text{Prop}$ be a meta-level (*i.e.*, Coq) predicate over values and suppose $\{\text{True}\} e \{v. \varphi v\}$ is derivable in Iris, then `safe e` and `post_valid (e,  $\varphi$ )`.*

Since Actris is an internal logic embedded in Iris, the proof is an immediate consequence of Iris's adequacy theorem (Theorem 3.1).

9.7. Solving the recursive domain equation for protocols. Recall the recursive domain equation for dependent separation protocols from §9.1:

$$\text{iProto} \cong 1 + (\text{action} \times (\text{Val} \rightarrow \blacktriangleright \text{iProto} \rightarrow \text{iProp}))$$

This recursive domain equation shows that `iProto` depends on the type `iProp` of Iris propositions. To use types that depend on `iProp` as part of higher-order ghost state in Iris, such types need to be bi-functorial in `iProp`. Hence, this means that to construct `iProto`, in a way that it can be used in combination with the higher-order ghost variables in Figure 19, we need to solve the following recursive domain equation:

$$\text{iProto}(X^-, X^+) \cong 1 + (\text{action} \times (\text{Val} \rightarrow \blacktriangleright \text{iProto}(X^+, X^-) \rightarrow X^+))$$

Since the recursive occurrence of `iProto` appears in negative position, the polarity needs to be inverted for `iProto` to be bi-functorial.

The version of Iris's recursive domain equation solver based on [AR89; BST10] as mechanised in Iris's Coq development is not readily able to construct a solution of `iProto`(X^-, X^+). Concretely, the solver can only construct solutions of non-parameterised recursive domain equations. While a general construction for solving such recursive domain equations exists [BMSS12, § 7], that construction has not been mechanised in Coq. We circumvent this shortcoming by solving the following recursive domain equation instead, in which we unfold the recursion once by hand:

$$\begin{aligned} \text{iProto}_2(X^-, X^+) &\cong \\ &1 + \left(\text{action} \times (\text{Val} \rightarrow \blacktriangleright (1 + (\text{action} \times (\text{Val} \rightarrow \blacktriangleright \text{iProto}_2(X^-, X^+) \rightarrow X^-))) \rightarrow X^+) \right) \end{aligned}$$

Here, the polarity in the recursive occurrence is fixed, allowing us to solve `iProto`₂(X^-, X^+) using Iris's existing recursive domain equation solver. This is sufficient because a solution of `iProto`₂(X^-, X^+) is isomorphic to a solution of `iProto`(X^-, X^+).

	Notation on paper	Notation in Coq
Send	$!x_1 \dots x_n \langle v \rangle \{P\}. prot$	<code><! x_1 .. x_n> MSG v {{ P }}; prot</code>
Receive	$?x_1 \dots x_n \langle v \rangle \{P\}. prot$	<code><? x_1 .. x_n> MSG v {{ P }}; prot</code>
End	end	END
Dual	$\overline{prot}$	<code>iProto_dual prot</code>
Literals	<code>()</code> , 5, true	<code>#()</code> , <code>#5</code> , <code>#true</code>
Logical variables	$x, y, z, -$	<code>"x"</code> , <code>"y"</code> , <code>"z"</code> , <code><></code>
Types	1, $\mathbb{N}$, $\mathbb{Z}$	<code>()</code> , <code>nat</code> , <code>Z</code>

Figure 21: Overview of notations in the Actris Coq mechanisation.

10. COQ MECHANISATION

The definition of the Actris logic, its model, and the proofs of all examples in this paper have been fully mechanised using the Coq proof assistant [Coq20]. In this section we will elaborate on the mechanisation effort (§ 10.1), and go through the full proof of a message-passing program (§ 10.2) and a subprotocol relation (§ 10.3) showcasing the tactics for Actris. We display proofs and proof states taken directly from the Coq mechanisation, which differ in notation from the paper as shown in Figure 21.

10.1. Mechanisation effort. The mechanisation of Actris is built on top of the mechanisation of Iris [KJB⁺17; JKBD16; KJ⁺18]. To carry out proofs in separation logic, we use the MoSeL Proof Mode (formerly Iris Proof Mode) [KTB17; KJJ⁺18], which provides an embedded proof assistant for separation logic in Coq. Building Actris on top of the Iris and MoSeL framework in Coq has a number of tangible advantages:

- By defining channels on top of HeapLang, we do not have to define a full programming language semantics, and can reuse all of the program libraries and Coq machinery, including the tactics for symbolic execution of non message-passing programs.
- Since Actris is mechanised as an Iris library we get all of the features of Iris for free, such as the ghost state mechanisms for reasoning about concurrency.
- When proving the Actris proof rules, we can make use of the MoSeL Proof Mode to carry out proofs directly using separation logic, thus reasoning at a high level of abstraction.
- We can make use of the extendable nature of the MoSeL Proof Mode to define custom tactics for symbolic execution of message-passing programs.

These advantages made it possible to mechanise Actris, along with the examples of the paper, with a small Coq development of a total size of about 5000 lines of code (comments and whitespace included). The line count of the different components are shown in Figure 22.

10.2. Tactic support for session type-based reasoning. To carry out interactive Actris proofs using symbolic execution, we follow the methodology described in the original Iris Proof Mode paper [KTB17]. In particular, this means that the logic in Coq is presented in weakest precondition style rather than using Hoare triples. For handling **send** or **recv** we define the following tactics:

`wp_send (t1 .. tn) with "[H1 .. Hn]"    and    wp_recv (y1 .. yn) as "H".`

These tactics roughly perform the following actions:

- Find a **send** or **recv** in evaluation position of the program under consideration.

Component	Sections	~LOC
The Actris model	§ 9.1–§ 9.4	1500
Channel implementation and proof rules	§ 2.4 and 9.5	350
Tactics for symbolic execution	§ 10.2	500
Utilities (linked lists, permutations, <i>etc.</i>)	n.a.	450
Authoritative contribution ghost theory	§ 7	150
Recursive domain equation theory solver	§ 9.7	100
Examples:		
• Basic examples	§ 1 and 7.1	400
• Coarse-grained channel-based merge sort	§ 5.1–§ 5.5	250
• Fine-grained channel-based merge sort	§ 5.6	300
• Mapper with swapping	§ 6.2	400
• List reversal	§ 6.3	100
• Channel-based load-balancing mapper	§ 7.2	200
• Channel-based map-reduce	§ 8	300
Total		5000

Figure 22: Overview of components of the Actris Coq mechanisation.

- Find a corresponding $c \mapsto \text{prot}$ hypothesis in the separation logic context.
- Normalise the protocol prot using the rules for duals, composition, recursion, and swapping so it has a $!\vec{x}:\vec{\tau}\langle v\rangle\{P\}.\text{prot}$ or $?\vec{x}:\vec{\tau}\langle v\rangle\{P\}.\text{prot}$ construct in its head position.
- In case of wp_send , instantiate the variables $\vec{x}:\vec{\tau}$ using the terms $(\text{t1} \dots \text{tn})$, and create a goal for the proposition P with the hypotheses $[\text{H1} \dots \text{Hn}]$. Hypotheses prefixed with $\$$ will automatically be consumed to resolve a subgoal of P if possible. In case the terms $(\text{t1} \dots \text{tn})$ are omitted, an attempt is made to determine these using unification.
- In case of wp_recv , introduce the variables $\vec{x}:\vec{\tau}$ into the context by naming them $(\text{y1} \dots \text{yn})$, and create a hypothesis H for P .

The implementation of these tactics follows the approach by Krebbers *et al.* [KTB17]. The protocol normalisation is implemented via logic programming with type classes.

As an example we will go through a proof of the following program:

```

prog_ref_swap_loop := λ_. let c := start (rec go c' := let l := recv c' in
                                           l ← !l + 2;
                                           send c' (); go c') in
  let l1 := ref 18 in let l2 := ref 20 in
  send c l1; send c l2;
  recv c; recv c;
  !l1 + !l2

```

Here, the forked-off thread acts as a service that recursively receives locations, adds 2 to their stored number, and then sends back a flag indicating that the location has been updated. The main thread, acting like a client, first allocates two new references, to 18 and 20, respectively, which are both sent to the service after which the update flags are received. It finally dereferences the updated locations, and adds their values together, thus returning 42. To verify this program, we use the following recursive protocol:

$$\text{prot_ref_loop} \triangleq \mu(\text{rec} : \text{iProto}). !(\ell : \text{Loc})(x : \mathbb{Z}) \langle \ell \rangle \{ \ell \mapsto x \}. ?\langle () \rangle \{ \ell \mapsto x + 2 \}. \text{rec}$$

```

1 Lemma prog_ref_swap_loop_spec :  $\forall \Phi, \Phi \#42 \text{ -* WP prog\_ref\_swap\_loop \#() \{ \Phi \} }$ .
2 Proof.
3   iIntros ( $\Phi$ ) "H $\Phi$ ". wp_lam.
4   wp_apply (start_chan_spec prot_ref_loop); iIntros (c) "Hc".
5   - iLöb as "IH". wp_lam.
6     wp_recv (l x) as "Hl". wp_load. wp_store. wp_send with "[ $\$Hl$ ]".
7     do 2 wp_pure _. by iApply "IH".
8   - wp_alloc l1 as "Hl1". wp_alloc l2 as "Hl2".
9     wp_send with "[ $\$Hl1$ ]". wp_send with "[ $\$Hl2$ ]".
10    wp_recv as "Hl1". wp_recv as "Hl2".
11    wp_load. wp_load.
12    wp_pures. by iApply "H $\Phi$ ".
13 Qed.

```

Figure 23: Proof of message-passing program

The (forked-off) service follows the (dual of) the protocol exactly, while the main thread follows a weakened version. The recursion is unfolded twice, after which the second send has been swapped ahead of the first receive, allowing it to first send both values before receiving:

$$\begin{aligned}
\text{prot_ref_loop} \sqsubseteq & !(\ell_1 : \text{Loc})(x_1 : \mathbb{Z}) \langle \ell_1 \rangle \{ \ell_1 \mapsto x_1 \}. \\
& !(\ell_2 : \text{Loc})(x_2 : \mathbb{Z}) \langle \ell_2 \rangle \{ \ell_2 \mapsto x_2 \}. \\
& ?\langle () \rangle \{ \ell_1 \mapsto (x_1 + 2) \}. \\
& ?\langle () \rangle \{ \ell_2 \mapsto (x_2 + 2) \}. \text{prot_ref_loop}
\end{aligned}$$

The full Coq proof of the program is shown in Figure 23. The proven lemma is logically equivalent to the specification $\{\text{True}\} \text{prog_ref_swap_loop } () \{v.v = 42\}$, but is presented in weakest precondition style as is common in Iris in Coq. The initial proof state is:

```

-----*
 $\forall \Phi, \Phi \#42 \text{ -* WP prog\_ref\_swap\_loop \#() \{ \Phi \} }$ 

```

We start the proof on line 3 by introducing the postcondition Φ , and the hypothesis $H\Phi$: $\Phi \#42$, and then continue by evaluating the lambda expression with `wp_lam`. On line 4 we apply the specification `start_chan_spec`, which is the weakest precondition variant of HT-START for `start` by picking the expected protocol `prot_ref_loop`. This leaves us with two subgoals, separated by bullets “-”: one for the forked-off thread, and one for the main thread.

Proof of the forked-off thread. In the proof of the recursively-defined forked-off thread we use `iLöb as "IH"` for LÖB induction on line 5. This leaves us with the proof state:

```

"IH" :  $\triangleright (c \mapsto \text{iProto\_dual prot\_ref\_loop -*$ 
      WP (rec: "go" "c'" :=
        let: "l" := recv "c'" in
        "l" <- ! "l" + #2;;
        send "c'" #();; "go" "c'") c {_, True })
-----□
"Hc" :  $c \mapsto \text{iProto\_dual prot\_ref\_loop}$ 
-----*
WP (rec: "go" "c'" :=
  let: "l" := recv "c'" in
  "l" <- ! "l" + #2;;
  send "c'" #();; "go" "c'") c {_, True })

```

We now resolve the application of c to the recursive function with wp_lam . This lets us strip the later from the LÖB induction hypothesis, as the program has taken a step. The proof state is then as follows:

```
"IH" : c  $\mapsto$  iProto_dual prot_ref_loop -* WP prog_rec c {{ _, True }}
-----□
"Hc" : c  $\mapsto$  iProto_dual prot_ref_loop
-----*
WP let: "l" := recv c in
  "l" <- ! "l" + #2;;
  send c #();; prog_rec c {{ _, True }}
```

For brevity's sake we abbreviate the recursive code in "IH" as $\text{prog_rec } c$.

On line 6 we resolve the proof of the body of the recursive function. So far, the proof only used Iris's standard tactics, we now use the Actris tactic for receive $\text{wp_recv } (l \ x)$ as "H1", to resolve the receive in evaluation position, introducing the received logical variables l and x , along with the predicate of the protocol $l \mapsto \#x$ naming it H1. To do so, the protocol is normalised, unfolding the recursive definition once, as well as resolving the dualisation of the head, turning it into a receive as expected. This leads to the following proof state:

```
"IH" : c  $\mapsto$  iProto_dual prot_ref_loop -* WP prog_rec c {{ _, True }}
-----□
"H1" : l  $\mapsto$  #x
"Hc" : c  $\mapsto$  iProto_dual (<?> MSG #() {{ l  $\mapsto$  #(x + 2) }}; prot_ref_loop)
-----*
WP let: "l" := #l in
  "l" <- ! "l" + #2;;
  send c #();; prog_rec c {{ _, True }}
```

We then use the HeapLang tactics wp_load and wp_store to resolve the dereferencing and updating of the location:

```
"IH" : c  $\mapsto$  iProto_dual prot_ref_loop -* WP prog_rec c {{ _, True }}
-----□
"H1" : l  $\mapsto$  #(x + 2)
"Hc" : c  $\mapsto$  iProto_dual (<?> MSG #() {{ l  $\mapsto$  #(x + 2) }}; prot_ref_loop)
-----*
WP send c #();; prog_rec c {{ _, True }}
```

We then use the Actris tactic wp_send with " $[\$H1]$ " to resolve the send operation in evaluation position, by giving up the ownership of "H1". Again, the protocol is automatically normalised by resolving the dualisation of the receive (?) to obtain the send (!) as expected.

We finally close the proof of the forked-off thread on line 7. We first take two pure evaluation steps revolving the sequencing of operations with $\text{do } 2 \text{ wp_pure } _$ to reach the recursive call. This results in the proof state:

```
"IH" : c  $\mapsto$  iProto_dual prot_ref_loop -* WP prog_rec c {{ _, True }}
-----□
"Hc" : c  $\mapsto$  iProto_dual prot_ref_loop
-----*
WP prog_rec c {{ _, True }}
```

We then use by iApply "IH" to close the proof by using the LÖB induction hypothesis.

Proof of the main thread. The proof of the main thread follows similarly. On line 8 we use `wp_alloc 11` as "H11" and `wp_alloc 12` as "H12", to resolve the allocations of the new locations, binding the logical variables of the locations to 11 and 12, and adding hypotheses "H11" and "H12" for ownership of these locations to the separation logic proof context. The proof state is then:

```
"HΦ" : Φ #42
"Hc" : c ↦ prot_ref_loop
"H11" : 11 ↦ #18
"H12" : 12 ↦ #20
-----*
WP send c #11;; send c #12;; recv c;; recv c;; ! #11 + ! #12 {{ v, Φ v }}
```

On line 9, we resolve the first send operation with the Actris tactic `wp_send` with "[H11]", by giving up ownership of the location 11. Here, the protocol is normalised by unfolding the recursive definition, after which the head symbol is a send (!) as expected. The resulting proof state is as follows:

```
"HΦ" : Φ #42
"H12" : 12 ↦ #20
"Hc" : c ↦ (<?> MSG #() {{ 11 ↦ #(18 + 2) }}; prot_ref_loop)
-----*
WP send c #12;; recv c;; recv c;; ! #11 + ! #12 {{ v, Φ v }}
```

To resolve the second send operation, we need to weaken the protocol using swapping (rule $\sqsubseteq$ -SWAP'), which is taken care of automatically by the Actris tactic `wp_send` with "[H12]". The normalisation detects that the protocol has a receive (?) as a head symbol, and therefore attempts swapping. To do so it steps ahead of the receive (?), and unfolds the recursive definition, which results in a send (!) as the first symbol after the head. It then detects that there are no dependencies between the two, and can thus apply the swapping rule $\sqsubseteq$ -SWAP', moving the send (!) ahead of the receive (?). With the head symbol now being a send (!), the symbolic execution continues as normal, resulting in the proof state:

```
"HΦ" : Φ #42
"Hc" : c ↦ (<?> MSG #() {{ 11 ↦ #(18 + 2) }};
           <?> MSG #() {{ 12 ↦ #(20 + 2) }}; prot_ref_loop)
-----*
WP recv c;; recv c;; ! #11 + ! #12 {{ v, Φ v }}
```

On line 10 we then proceed as expected with `wp_recv` as "H11" and `wp_recv` as "H12", to resolve the receive operations, giving us back the updated point-to resources:

```
"HΦ" : Φ #42
"H11" : 11 ↦ #(18 + 2)
"H12" : 12 ↦ #(20 + 2)
"Hc" : c ↦ prot_ref_loop
-----*
WP ! #11 + ! #12 {{ v, Φ v }}
```

At line 11 we then continue by using `wp_load` twice to dereference the reacquired and updated locations, and then use trivial symbolic execution using `wp_pures` to resolve the remaining computations. On line 12 we finally close the proof by applying the hypothesis "HΦ" about the postcondition.

```

1 Lemma list_rev_subprot :
2   ⊢ (<!(l : loc) (vs : list val)> MSG #1 {{ llist l vs }};
3     <?> MSG #() {{ llist internal_eq l (reverse vs) }}; END) ⊆
4     (<!(l : loc) (xs : list T)> MSG #1 {{ llistI IT l xs }};
5       <?> MSG #() {{ llistI IT l (reverse xs) }}; END).
6 Proof.
7   iIntros (l xs) "H1".
8   iDestruct (H1r with "H1") as (vs) "[H1 HIT]".
9   iExists l, vs. iFrame "H1".
10  iModIntro. iIntros "H1".
11  iSplitL.
12  { rewrite big_sepL2_reverse_2. iApply H1r.
13    iExists (reverse vs). iFrame "H1 HIT". }
14  done.
15 Qed.

```

Figure 24: Proof of subprotocol relation

10.3. Tactic support for subprotocols. While the Actris tactics automatically apply the subprotocol rules during symbolic execution, as shown in § 10.2, we sometimes want to prove subprotocol relations as explicit lemmas. We have tactic support for such proofs as well. We extend the existing MoSeL tactics `iIntros`, `iExists`, `iFrame`, `iModIntro`, and `iSplitL`/`iSplitR` to automatically use the subprotocol rules to turn the goal into an equivalent goal where the regular Iris tactics apply.

- `iIntros (x1 .. xn) "H1 .. Hm"` transforms the subprotocol goal to begin with n universal quantification and m implications, using the rules $\sqsubseteq$ -SEND-OUT and $\sqsubseteq$ -RECV-OUT, and then introduces the quantifiers (naming them $x1 \dots xn$) into the Coq context, and the hypotheses (naming them $H1 \dots Hm$) into the separation logic context.
- `iExists (t1 .. tn)` transforms the subprotocol goal to start with n existential quantifiers, using the $\sqsubseteq$ -SEND-IN, $\sqsubseteq$ -RECV-IN and $\sqsubseteq$ -TRANS rules, and then instantiates these quantifiers with the terms $t1 \dots tn$ specified by the pattern.
- `iFrame "H"` transforms the subprotocol goal into a separating conjunction between the payload predicates of the head symbols of either protocol, using the rules $\sqsubseteq$ -SEND-IN and $\sqsubseteq$ -RECV-IN, and then tries to solve the payload predicate subgoal using "H".
- `iModIntro` transforms the subprotocol goal into a goal starting with a later modality ($\triangleright$), using the rules $\sqsubseteq$ -SEND-MONO and $\sqsubseteq$ -RECV-MONO, and then introduces that later by stripping off a later from any hypothesis in the separation logic context.
- `iSplitL`/`iSplitR` "H1 .. Hn" transforms the subprotocol goal into a separating conjunction between the payload predicates of the head symbols of either protocol, using the $\sqsubseteq$ -SEND-IN, $\sqsubseteq$ -RECV-IN and $\sqsubseteq$ -TRANS rules, and then creates two subgoals. For `iSplitL` the left subgoal is given the hypotheses $H1 \dots Hn$ from the separation logic context, while the right subgoal is given any remaining hypotheses, and *vice versa* for `iSplitR`.

The extensions of these tactics are implemented by defining custom type class instances that hook into the existing MoSeL tactics as described by Krebbers *et al.* [KTB17].

To demonstrate these tactics, we will go through a proof of the subprotocol relation for the list reversing service presented in § 6.3:

$$\begin{aligned}
& !(\ell : \text{Loc})(\vec{v} : \text{List Val}) \langle \ell \rangle \{ \ell \xrightarrow{\text{list}} \vec{v} \}. ? \langle () \rangle \{ \ell \xrightarrow{\text{list}} \text{reverse } \vec{v} \}. \text{end} \\
& \sqsubseteq !(\ell : \text{Loc})(\vec{x} : \text{List } T) \langle \ell \rangle \{ \ell \xrightarrow{\text{list}}_{IT} \vec{x} \}. ? \langle () \rangle \{ \ell \xrightarrow{\text{list}}_{IT} \text{reverse } \vec{x} \}. \text{end}
\end{aligned}$$

Recall that the following conversion between the list representation predicate with payload $\ell \xrightarrow{\text{list}}_{IT} \vec{x}$ and one without payload $\ell \xrightarrow{\text{list}} \vec{v}$ holds:

$$\text{Hlr} : \ell \xrightarrow{\text{list}}_{IT} \vec{x} \text{ ** } (\exists \vec{v}. \ell \xrightarrow{\text{list}} \vec{v} * \bigstar_{(x,v) \in (\vec{x}, \vec{v})} . IT \ x \ v)$$

The full Coq proof of the subprotocol relation is shown in Figure 24. The initial proof state is identical to the lemma statement. On line 7 we start by introducing the logical variables `l`, `xs` and the payload `llistI IT l xs` of the weaker protocol with the tactic `iIntros (l xs) "H1"`. This tactic will implicitly apply the rule $\sqsubseteq$ -SEND-OUT, so the goal starts with a universal quantification $\forall(l : \text{loc})(xs : \text{list } T). \text{llistI IT l xs} \text{ -* } \dots$, which is then introduced based on the regular Iris introduction pattern. This gives us:

```
"H1" : llistI IT l vs
-----*
(<! (l : loc) (vs : list val)> MSG #1 {{ llist l vs }};
 <?> MSG #() {{ llist l (reverse vs) }}; END)  $\sqsubseteq$ 
(<!> MSG #1; <?> MSG #() {{ llistI IT l (reverse xs) }}; END)
```

To obtain the payload predicate expected by the stronger protocol, we use the lemma `Hlr`, to derive `llist l vs` and `[*list] x;v  $\in$  xs;vs, IT x v` from `llistI l xs` with the tactic `iDestruct (Hlr with "H1") as (vs) "[H1 HIT]"` on line 8. The resulting proof state is:

```
"H1" : llist l vs
"HIT" : [*list] x;v  $\in$  xs;vs, IT x v
-----*
(<! (l : loc) (vs : list val)> MSG #1 {{ llist l vs }};
 <?> MSG #() {{ llist l (reverse vs) }}; END)  $\sqsubseteq$ 
(<!> MSG #1; <?> MSG #() {{ llistI IT l (reverse xs) }}; END)
```

At line 9 we instantiate the logical variables of the stronger protocol with the logical variables `l` and `vs` using `iExists l, vs`. This will implicitly apply the rules $\sqsubseteq$ -SEND-IN and $\sqsubseteq$ -TRANS, which makes the goal start with $\exists(l : \text{loc})(vs : \text{list val})$, so the existentials can be instantiated. To resolve the payload predicate obligation `llist l vs`, we use `iFrame "H1"`. This uses the rules $\sqsubseteq$ -SEND-IN and $\sqsubseteq$ -TRANS to turn the goal into `llist l vs * ...`, where the left subgoal is resolved using "H1". We then have the following remaining proof state:

```
"HIT" : [*list] x;v  $\in$  xs;vs, IT x v
-----*
(<!> MSG #1; <?> MSG #() {{ llist l (reverse vs) }}; END)  $\sqsubseteq$ 
(<!> MSG #1; <?> MSG #() {{ llistI IT l (reverse xs) }}; END)
```

As the head symbols of both protocols are sends (!) with no logical variables or payload predicates, we use `iModIntro` on line 10, which first applies $\sqsubseteq$ -SEND-MONO to step over the sends, and then introduces the later modality ($\triangleright$). This gives us the proof state:

```
"HIT" : [*list] x;v  $\in$  xs;vs, IT x v
-----*
(<?> MSG #() {{ llist l (reverse vs) }}; END)  $\sqsubseteq$ 
(<?> MSG #() {{ llistI IT l (reverse xs) }}; END)
```

On line 10, similarly to before, we use `iIntros "H1"`, to introduce the payload predicate, but this time we do it for the stronger protocol, as dictated by $\sqsubseteq$ -RECV-OUT:

```
"HIT" : [*list] x;v  $\in$  xs;vs, IT x v
"H1" : llist l (reverse vs)
-----*
```

```
(<?> MSG #() ; END)  $\sqsubseteq$ 
(<?> MSG #() {{ llistI IT 1 (reverse xs) }}; END)
```

To resolve the payload predicate of the weaker protocol, we use `iSplitL "H1 HIT"` on line 11, that first use $\sqsubseteq$ -RECV-IN and $\sqsubseteq$ -TRANS, to turn the goal into `llistI IT 1 (reverse xs) * ...`, and then use the goal splitting pattern of Iris, to give us two subgoals, where we use the hypotheses "H1" and "HIT" in the left subgoal. The first subgoal is then:

```
"HIT" : [* list] x;v  $\in$  xs;vs, IT x v
"H1" : llist 1 (reverse vs)
-----*
llistI IT 1 (reverse xs)
```

On line 12, we first use the lemma `H1r` in the right-to-left direction, and then rewrite the hypothesis "HIT" using a lemma from the Iris library with `rewrite big_sepL2_reverse_2`. We do this to obtain `[*list] x;v  $\in$  reverse xs;reverse vs, IT x v`, in order to match the proof goal. This gives the proof obligation:

```
"HIT" : [* list] x;v  $\in$  reverse xs;reverse vs, IT x v
"H1" : llist 1 (reverse vs)
-----*
 $\exists$  vs : list val, llist 1 vs * ([* list] x;v  $\in$  reverse xs;vs, IT x v)
```

We finally close the proof on line 13 with `iExists (reverse vs)`, followed by `iFrame "H1 HIT"`, as the goal matches the hypotheses exactly, when picking `reverse vs` as the existential quantification. We then move on to the second subgoal:

```
-----*
(<?> MSG #(); END)  $\sqsubseteq$  (<?> MSG #(); END)
```

We resolve this subgoal, on line 14, with the tactic `done`, which tries to close the proof, by automatically applying $\sqsubseteq$ -REFL.

11. RELATED WORK

This section elaborates on the relation to message passing in separation logic (§ 11.1) and process calculi (§ 11.2), session types (§ 11.3), session subtyping (§ 11.4), endpoint sharing (§ 11.5), and verification of map-reduce (§ 11.6).

11.1. Message passing and separation logic. Lozes and Villard [VLC09; LV12] present a logic for contract-based reasoning about programs in a small imperative language with bi-directional asynchronous channels. Contracts are represented by finite-state automata with labelled send or receive transitions, equipped with separation logic predicates. Similar to session types (and Actris), contracts have a notion of duality, but unlike Actris they do not support dependencies between messages. Their logic supports ownership transfer (including ownership transfer of channels, akin to delegation), session-type like choice, and a form of recursive contracts. Their language has a close operation for channel deallocation instead of being garbage collected. A restriction to structured concurrency (*i.e.*, par instead of fork-based), structured channel deallocation (*i.e.*, must close both endpoints together) and linear (instead of affine) logic ensures memory-leak freedom. A form of channel sharing is supported, which we further discuss in § 11.5.

Craciun *et al.* [CKC15] introduced *session logic*, a variant of separation logic that includes predicates for protocol specifications similar to ours. This work includes support for mutable state, ownership transfer (including ownership transfer of channels, akin to delegation), session-type like choice using a special type of disjunction operator on the protocol level, and a sketch of an approach to verify deadlock freedom of programs. Combined, these features allow them to verify interesting and non-trivial message-passing programs. Their logic as a whole is not higher-order, which means that sending functions over channels is not possible. Moreover, their logic does not support protocol-level logical variables that can connect the transferred message with the tail protocol. It is therefore not possible to model dependent protocols like we do in Actris. Their work includes a notion of subtyping as weakening and strengthening of the payload predicates, however they do not consider swapping, and do not allow manipulation of resources as a part of their subtyping relation. There also exists no support for other concurrency primitives such as locks, which by extension means that manifest sharing is not possible. In Actris we get this for free by building on top of Iris, and reusing its ghost state mechanism. Their work has not been mechanised in a proof assistant, but example programs can be checked using the HIP/SLEEK verifier.

The original Iris paper [JSS⁺15] includes a small message-passing language with channels that do not preserve message order. It was included to demonstrate that Iris is flexible enough to handle other concurrency models than standard shared-memory concurrency. Since the Hoare triples for send and receive reason about the entire channel buffer, protocol reasoning must be done via STSs or other forms of ghost state.

Hamin and Jacobs [HJ19] take an orthogonal direction and use separation logic to prove deadlock freedom of programs that communicate via message passing using a custom logic tailored to this purpose. They do not provide abstractions akin to our session-type based protocols. Instead one has to reason using invariants and ghost state explicitly.

Mansky *et al.* [MAN17] verify the functional correctness of a message-passing system written in C using the VST framework in Coq [App14]. While they do not verify message-passing programs like we do, they do verify that the implementation of their message-passing system is resilient to faulty behaviour in the presence of malicious senders and receivers.

Tassarotti *et al.* [TJH17] prove correctness and termination preservation of a compiler from a simple language with session types to a functional language with mutable state, where channels are implemented using references on the heap. This work is also done in Iris in Coq. The session types they consider are more like standard session types, which cannot express functional properties of messages, but only their types.

The Disel logic by Sergey *et al.* [SWT18] and the Aneris logic by Krogh-Jespersen *et al.* [KTO⁺20] can be used to reason about message-passing programs that work on network sockets. Channels can only be used to send strings, are not order preserving, and messages can be dropped but not duplicated. Since only strings are sent over channels complex data (such as functions) must be marshalled and unmarshalled in order to be sent over the network. Both Disel and Aneris therefore address a different problem than we do.

SteelCore [SRF⁺20] is a framework for concurrent separation logic embedded in the F^{*} language. SteelCore has been used to encode unidirectional synchronous channels that can be typed with protocols akin to session types. Their protocols are defined as a dependent sequence of value obligations with associated separation logic predicates, dictating what can be sent over the channel, including the transfer of ownership. Channels are first-class and can also be transferred (akin to delegation), but their protocols do not include higher-order

protocol-level logical variables, or subtyping. They postulated that their approach scales to bidirectional asynchronous communication, but left that for future work.

11.2. Separation logic and process calculi. Another approach to verify message-passing programs is to combine separation logic and process calculus. Neither of the approaches below support delegation or concurrency paradigms other than message passing.

Francalanza *et al.* [FRS11] use separation logic to verify programs written in a CCS-like language. Channels model memory location, which has the effect that their input-actions behave a lot like our updates of mutable state with variable substitutions updating the state. As a proof of concept they prove the correctness of an in-place quick-sort algorithm.

Oortwijn *et al.* [OBH16] use separation logic and the mCRL2 process calculus to model communication protocols. The logic itself operates on a high level of abstraction and deals exclusively with intraprocess communication where a fractional separation logic is used to distribute channel resources to concurrent threads. Protocols are extracted from code, but there is no formal connection between the specification logic and the underlying language.

11.3. Session types. Seminal work on linear type systems for the π -calculus by Kobayashi *et al.* [KPT96] led to the creation of binary session types by Honda *et al.* [HVK98], and consequentially multiparty session types by Honda *et al.* [HYC08].

Later work by Dardha *et al.* [DGS12] helped merge the linear type systems of Kobayashi with Honda's session types, which facilitated the incorporation of session types in mainstream programming languages like Go [LNTY18], OCaml [Pad17; IYY19], and Java [HKP⁺10]. These works focus on adding session-typed support for message passing in existing languages, but do not target functional correctness.

Bocchi *et al.* [BHTY10] pushed the boundaries of what can be verified with (multiparty) session types while staying within a decidable fragment of first-order logic. They use first-order predicates to describe properties of values being sent and received. Decidability is maintained by imposing restrictions on these predicates, such as ensuring that nothing is sent that will be invalidated down the line. The constraints on the logic do, however, limit what programs can be verified. The work includes standard subtyping on communicated values and on choices, but no notion of swapping sends ahead of receives.

Caires and Pfenning [CP10] discovered a correspondence between intuitionistic linear logic and π -calculus with session types, which was extended with quantifiers and dependent types by Toninho *et al.* [TCP11]. These quantifiers range over both terms and propositions of an LF-based logic [CP96], and can be used to specify basic properties of the exchanged values. Toninho and Yoshida [TY18] extended this work by allowing the structure of the protocol to depend on the quantifiers. This notion of dependency allows for protocols where the length of the (tail) protocol depends on the values that were previously exchanged, similar to what we do in § 5.6. Finally, Das and Pfenning [DP20a; DP20b] developed a dependent session-type system with domain-specific logic for verifying arithmetic properties of programs with message passing.

Another approach to dependent session types was carried out by Thiemann and Vasconcelos [TV20] who introduced label-dependent session types. They unify universal and existential quantifiers with the send and receive primitives of conventional session types. Hence, similar to Actris, the choice connectives ($\&$ and $\oplus$) can be derived.

Toninho *et al.* [TCP14] and Lindley and Morris [LM16] developed session-type systems with termination guarantees in the presence of recursive (session) types. This is achieved by imposing a discipline similar to (co)inductive definitions in Coq and Agda. In contrast, Actris poses no usage discipline on recursive dependent separation protocols, and hence guarantees partial correctness.

11.4. Session subtyping. Actris’s subprotocol relation is inspired by the notion of session subtyping, for which seminal work was carried out by Gay and Hole [GH05]. Mostrous *et al.* [MYH09] extended session subtyping to multiparty asynchronous session types, and as part of that, introduced the notion of swapping sends ahead of receives for independent channels. Mostrous *et al.* [MY15] later considered swapping over the same channel in the context of binary session types. Our subprotocol relation is most closely related to the work of Mostrous *et al.* [MY15], although they define subtyping as a simulation on infinite trees, using so-called asynchronous contexts, whereas we define it using Iris’s support for guarded recursion. It should be noted that the work by Gay and Hole [GH05] differs from the work by Mostrous *et al.* [MYH09] and Mostrous *et al.* [MY15] in the orientation of the subtyping relation, as discussed by Gay [Gay16]. Our subprotocol relation uses the orientation of Gay and Hole [GH05].

Session subtyping for recursive type systems is universally carried out as a type simulation on infinite trees [GH05; MYH09; MY15], which complicates subtyping under the recursion operator. Bernardi *et al.* [BDGK14] and Gay *et al.* [GTV20] provide further insights on this problem, although they primarily investigate duality rather than subtyping.

To reason about recursive subtyping, Brand and Henglein [BH98] present a coinductive formulation of subtyping (which they apply to regular type systems, rather than session types). We use a similar coinductive formulation, but instead of ordinary coinduction, we use Iris’s support for guarded recursion, which lets us prove subtyping relations of recursive protocols using Löb induction.

11.5. Endpoint sharing. One of the key features of conventional session types is that endpoints are owned by a single thread. While endpoints can be delegated (*i.e.*, transferred from one thread to another), they typically cannot be shared (*i.e.*, be accessed by multiple threads concurrently). However, as demonstrated in § 7, sharing channels endpoints is often desirable, and possible in Actris.

As a simple way to relax this limitation of sharing in conventional session types, Vasconcelos [Vas12] allows session types of the form $(\mu rec. !T. rec)$ or $(\mu rec. ?T. rec)$ to be shared. Lozes and Villard [LV12] present a similar idea in the context of their contract-based separation logic (see also § 11.1) by equipping the connective for channel endpoint ownership with a fractional permission. If the fraction is smaller than 1, then the endpoint can be shared, but at the cost of only permitting transitions to the same contract state. Using fractional permissions they prove a lock specification à la Gotsman *et al.* [GBC⁺07] of an implementation of locks in terms of channels. This approach to locks is dual to ours in Actris, where we implement channels in terms of locks. Unlike Iris (and Actris), their logic does not support ghost state, so it cannot express complex protocols like the ones from § 7.

In the π -calculus community there has been prior work on endpoint sharing, *e.g.*, by Atkey *et al.* [ALM16], Kobayashi [Kob06], and Padovani [Pad14]. The latest contribution in this line of work is by Balzer *et al.* [BTP19], who developed a type system based on

session types with support for manifest sharing. Manifest sharing is the notion of sharing a channel endpoint between multiple processes using a lock-like structure to ensure mutual exclusion. Their key idea to ensure mutual exclusion using a type system is to use adjoint modalities to connect two classes of types: types that are linear, and thus denote unique channel ownership, and types that are unrestricted, and thus can be shared. The approach to endpoint sharing in Actris is different: dependent separation protocols do not include a built-in notion for endpoint sharing, but can be combined with Iris’s general-purpose mechanisms for sharing, like locks.

11.6. Verification of map-reduce. To our knowledge the only verification related to the map-reduce model [DG04] is by Ono *et al.* [OHT⁺11], who made two mechanisations in Coq. The first took a functional model of map-reduce and verified a few specific mappers and reducers, extracted these to Haskell, and ran them using Hadoop Streaming. The second did the same by annotating Java mappers and reducers using JML and proving them correct using the Krakatoa tool [MPU04], using a combination of SAT-solvers and the Coq proof assistant. While they worked on verifying specific mappers and reducers, our case study focuses on verifying the communication of a map-reduce model that can later be parameterised with concrete mappers and reducers.

12. CONCLUSION AND FUTURE WORK

In this paper, we have given a comprehensive account of the Actris concurrent separation logic for proving functional correctness of programs that combine message-passing with other programming and concurrency paradigms. The core feature of Actris is its mechanism of dependent separation protocols, which is inspired by session types. Considering the rich literature on session types and concurrent separation logic, we expect there to be many promising directions for future work.

Multi-party. The formalism of multi-party session types [HYC08] applies to message-passing communication between more than two parties (threads or processes). The key ingredient of multi-party session types is the notion of a *global protocol*, which specifies the permitted communication for multiple parties of a system. From the global protocol one can then generate *local protocols* for the individual parties. It would be interesting to explore a multi-party version of dependent separation protocols. Prior work by Costea *et al.* [CCQC18] on multi-party session logic and Zhoud *et al.* [ZFH⁺20] on refined multiparty session types could serve as a starting point.

Deadlock freedom. As discussed in §4.3, deadlocks are valid behaviours according to the notion of safety used in Iris (and thus Actris). Many conventional session type systems do not consider deadlocks to be valid behaviours, but achieve that at the expense of prohibiting valid (deadlock free) programs that can be verified in Actris.

A direction for future work is to develop a variant of Actris that incorporates the usual restrictions of session-type systems like linearity and a **start** primitive for combined channel and thread creation. To prove an adequacy theorem that ensures that this variant of Actris indeed prohibits deadlocks, one needs to change the model of Actris to ensure acyclicity of the dependency structure among the threads and channels. This could be achieved by

building upon recent work by Bizjak *et al.* [BGKB19] on linearity in Iris and by Jacobs *et al.* [JBK21] on a separation-logic based proof method for deadlock freedom of session types. Additionally, one could consider a version of Actris without garbage collection but with a `close` instruction for channel deallocation, and prove that it indeed guarantees memory-leak freedom.

Another direction for future work is to develop a separation logic that combines session-type based deadlock freedom with lock-order based deadlock freedom to prove deadlock freedom of programs that combine message passing with other concurrency mechanisms like locks. The work by Hamin and Jacobs [HJ19] on reasoning about lock orders in separation logic, and the work by Balzer *et al.* [BTP19] on deadlock freedom for manifest sharing might provide valuable insights, but figuring out how to combine these two approaches with Iris and Actris is a challenging open problem.

ACKNOWLEDGMENTS

We thank the anonymous reviewers of both this paper and the POPL'20 conference version for their helpful feedback. We are grateful to Andreea Costea, Daniel Gratzer, Daniël Louwink, Fabrizio Montesi, Marco Carbone, and the participants of the Iris workshop 2019 for discussions. The third author (Robbert Krebbers) was supported by the Dutch Research Council (NWO), project 016.Veni.192.259.

REFERENCES

- [Ahm04] Amal Ahmed. *Semantics of Types for Mutable State*. PhD thesis, Princeton University, 2004.
- [ALM16] Robert Atkey, Sam Lindley, and J. Garrett Morris. Conflation confers concurrency. In *Essays Dedicated to Philip Wadler on the Occasion of His 60th Birthday*, pages 32–55, 2016.
- [AM01] Andrew W. Appel and David McAllester. An indexed model of recursive types for foundational proof-carrying code. *TOPLAS*, 23(5):657–683, 2001.
- [AMRV07] Andrew W. Appel, Paul-André Melliès, Christopher D. Richards, and Jérôme Vouillon. A very modal model of a modern, major, general type system. In *POPL*, pages 109–122, 2007.
- [App14] Andrew W. Appel. *Program Logics for Certified Compilers*. Cambridge University Press, 2014.
- [AR89] Pierre America and Jan J. M. M. Rutten. Solving reflexive domain equations in a category of complete metric spaces. *JCSS*, 39(3):343–375, 1989.
- [BB20] Lars Birkedal and Aleš Bizjak. Lecture notes on iris: Higher-order concurrent separation logic. <https://iris-project.org/tutorial-material.html>, 2020.
- [BDGK14] Giovanni Bernardi, Ornela Dardha, Simon J. Gay, and Dimitrios Kouzapas. On duality relations for session types. In *TGC*, volume 8902 of *LNCS*, pages 51–66, 2014.
- [BGKB19] Ales Bizjak, Daniel Gratzer, Robbert Krebbers, and Lars Birkedal. Iron: Managing obligations in higher-order concurrent separation logic. *PACMPL*, 3(POPL):65:1–65:30, 2019.

- [BH98] Michael Brandt and Fritz Henglein. Coinductive axiomatization of recursive type equality and subtyping. *Fundamenta Informaticae*, 33(4):309–338, 1998.
- [BHTY10] Laura Bocchi, Kohei Honda, Emilio Tuosto, and Nobuko Yoshida. A theory of design-by-contract for distributed multiparty interactions. In *CONCUR*, pages 162–176, 2010.
- [BMSS12] Lars Birkedal, Rasmus Ejlers Møgelberg, Jan Schwinghammer, and Kristian Støvring. First steps in synthetic guarded domain theory: Step-indexing in the topos of trees. *LMCS*, 8(4), 2012.
- [BP17] Stephanie Balzer and Frank Pfenning. Manifest sharing with session types. *PACMPL*, 1(ICFP):37:1–37:29, 2017.
- [BST10] Lars Birkedal, Kristian Støvring, and Jacob Thamsborg. The category-theoretic solution of recursive metric-space equations. *TCS*, 411(47):4102–4122, 2010.
- [BTP19] Stephanie Balzer, Bernardo Toninho, and Frank Pfenning. Manifest Deadlock-Freedom for Shared Session Types. In *ESOP*, pages 611–639, 2019.
- [CCQC18] Andreea Costea, Wei-Ngan Chin, Shengchao Qin, and Florin Craciun. Automated modular verification for relaxed communication protocols. In *APLAS*, pages 284–305, 2018.
- [CKC15] Florin Craciun, Tibor Kiss, and Andreea Costea. Towards a session logic for communication protocols. In *ICECCS*, pages 140–149, 2015.
- [Coq20] Coq Development Team. The Coq proof assistant reference manual, version 8.12.0. 2020. URL: <https://coq.inria.fr/distrib/current/refman/>.
- [CP96] Ilario Cervesato and Frank Pfenning. A linear logical framework. In *LICS*, pages 264–275, 1996.
- [CP10] Luís Caires and Frank Pfenning. Session types as intuitionistic linear propositions. In *CONCUR*, pages 222–236, 2010.
- [DAB11] Derek Dreyer, Amal Ahmed, and Lars Birkedal. Logical step-indexed logical relations. *LMCS*, 7(2), 2011.
- [DG04] Jeffrey Dean and Sanjay Ghemawat. Mapreduce: Simplified data processing on large clusters. In *OSDI*, pages 137–150, 2004.
- [DGS12] Ornella Dardha, Elena Giachino, and Davide Sangiorgi. Session types revisited. In *PPDP*, pages 139–150, 2012.
- [DP20a] Ankush Das and Frank Pfenning. Session types with arithmetic refinements. In *CONCUR*, volume 171 of *LIPICs*, pages 13:1–13:18, 2020.
- [DP20b] Ankush Das and Frank Pfenning. Verified linear session-typed concurrent programming. In *PPDP*, pages 7:1–7:15. ACM, 2020.
- [dRPDG14] Pedro da Rocha Pinto, Thomas Dinsdale-Young, and Philippa Gardner. Tada: A logic for time and data abstraction. In *ECOOP*, pages 207–231, 2014.
- [FH92] Matthias Felleisen and Robert Hieb. The revised report on the syntactic theories of sequential control and state. *TCS*, 103(2):235–271, 1992.
- [FRS11] Adrian Francalanza, Julian Rathke, and Vladimiro Sassone. Permission-based separation logic for message-passing concurrency. *LMCS*, 7(3), 2011.
- [Gay16] Simon J. Gay. Subtyping supports safe session substitution. In *A List of Successes That Can Change the World - Essays Dedicated to Philip Wadler on the Occasion of His 60th Birthday*, pages 95–108, 2016.
- [GBC⁺07] Alexey Gotsman, Josh Berdine, Byron Cook, Noam Rinetzkky, and Mooly Sagiv. Local reasoning for storable locks and threads. In Zhong Shao, editor, *APLAS*, volume 4807 of *LNCs*, pages 19–37, 2007.

- [GH05] Simon J. Gay and Malcolm Hole. Subtyping for session types in the pi calculus. *Acta Informatica*, 42(2-3):191–225, 2005.
- [GST⁺20] Paolo G. Giarrusso, Léo Stefanescu, Amin Timany, Lars Birkedal, and Robbert Krebbers. Scala step-by-step: Soundness for DOT with step-indexed logical relations in Iris. *PACMPL*, 4(ICFP):114:1–114:29, 2020.
- [GTV20] Simon J. Gay, Peter Thiemann, and Vasco T. Vasconcelos. Duality of session types: The final cut. In *PLACES*, volume 314 of *EPTCS*, pages 23–33, 2020.
- [HAN08] Aquinas Hobor, Andrew Appel, and Francesco Nardelli. Oracle semantics for concurrent separation logic. In *ESOP*, volume 4960 of *LNCS*, pages 353–367, 2008.
- [HBK20] Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. Actris: Session-type based reasoning in separation logic. *PACMPL*, 4(POPL):6:1–6:30, 2020.
- [HBK21] Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. Coq mechanisation of actris, 2021. Available online at <https://gitlab.mpi-sws.org/iris/actris/-/tree/lmcs>.
- [HBS73] Carl Hewitt, Peter Bishop, and Richard Steiger. A Universal Modular ACTOR Formalism for Artificial Intelligence. In *IJCAI*, pages 235–245, 1973.
- [HJ19] Jafar Hamin and Bart Jacobs. Transferring obligations through synchronizations. In *ECOOP*, pages 19:1–19:58, 2019.
- [HKP⁺10] Raymond Hu, Dimitrios Kouzapas, Olivier Pernet, Nobuko Yoshida, and Kohei Honda. Type-safe eventful sessions in java. In *ECOOP*, pages 21–25, 2010.
- [HLKB21] Jonas Kastberg Hinrichsen, Daniël Louwink, Robbert Krebbers, and Jesper Bengtson. Machine-checked semantic session typing. In *CPP*, pages 178–198, 2021.
- [HVK98] Kohei Honda, Vasco Thudichum Vasconcelos, and Makoto Kubo. Language primitives and type discipline for structured communication-based programming. In *ESOP*, pages 122–138, 1998.
- [HYC08] Kohei Honda, Nobuko Yoshida, and Marco Carbone. Multiparty asynchronous session types. In *POPL*, pages 273–284, 2008.
- [Iri21] Iris Development Team. The mechanisation of Iris. 2021. URL: <https://gitlab.mpi-sws.org/iris/iris/>.
- [IYY19] Keigo Imai, Nobuko Yoshida, and Shoji Yuen. Session-ocaml: A session-based library with polarities and lenses. *Science of Computer Programming*, 172:135–159, 2019.
- [JBK21] Jules Jacobs, Stephanie Balzer, and Robbert Krebbers. Connectivity graphs: A method for proving deadlock freedom based on separation logic, 2021. Manuscript under submission.
- [JJKD18] Ralf Jung, Jacques-Henri Jourdan, Robbert Krebbers, and Derek Dreyer. Rust-belt: Securing the foundations of the rust programming language. *PACMPL*, 2(POPL):66:1–66:34, 2018.
- [JJKD21] Ralf Jung, Jacques-Henri Jourdan, Robbert Krebbers, and Derek Dreyer. Safe systems programming in rust. *CACM*, 64(4):144–152, 2021.
- [JKBD16] Ralf Jung, Robbert Krebbers, Lars Birkedal, and Derek Dreyer. Higher-order ghost state. In *ICFP*, pages 256–269, 2016.

- [JKJ⁺18] Ralf Jung, Robbert Krebbers, Jacques-Henri Jourdan, Ales Bizjak, Lars Birkedal, and Derek Dreyer. Iris from the ground up: A modular foundation for higher-order concurrent separation logic. *JFP*, 28:e20, 2018.
- [JSS⁺15] Ralf Jung, David Swasey, Filip Sieczkowski, Kasper Svendsen, Aaron Turon, Lars Birkedal, and Derek Dreyer. Iris: Monoids and invariants as an orthogonal basis for concurrent reasoning. In *POPL*, pages 637–650, 2015.
- [KJB⁺17] Robbert Krebbers, Ralf Jung, Ales Bizjak, Jacques-Henri Jourdan, Derek Dreyer, and Lars Birkedal. The essence of higher-order concurrent separation logic. In *ESOP*, pages 696–723, 2017.
- [KJJ⁺18] Robbert Krebbers, Jacques-Henri Jourdan, Ralf Jung, Joseph Tassarotti, Jan-Oliver Kaiser, Amin Timany, Arthur Charguéraud, and Derek Dreyer. Mosel: A general, extensible modal framework for interactive proofs in separation logic. *PACMPL*, 2(ICFP):77:1–77:30, 2018.
- [Kob06] Naoki Kobayashi. A new type system for deadlock-free processes. In *CONCUR*, pages 233–247, 2006.
- [KPT96] Naoki Kobayashi, Benjamin C. Pierce, and David N. Turner. Linearity and the pi-calculus. In *POPL*, pages 358–371, 1996.
- [KTB17] Robbert Krebbers, Amin Timany, and Lars Birkedal. Interactive proofs in higher-order concurrent separation logic. In *POPL*, pages 205–217, 2017.
- [KTO⁺20] Morten Krogh-Jespersen, Amin Timany, Marit Edna Ohlenbusch, Simon Oddershede Gregersen, and Lars Birkedal. Aneris: A mechanised logic for modular reasoning about distributed systems. In *ESOP*, pages 336–365, 2020.
- [LM16] Sam Lindley and J. Garrett Morris. Talking bananas: Structural recursion for session types. In *ICFP*, pages 434–447, 2016.
- [LNTY18] Julien Lange, Nicholas Ng, Bernardo Toninho, and Nobuko Yoshida. A static verification framework for message passing in go using behavioural types. *ICSE*, pages 1137–1148, 2018.
- [LV12] Étienne Lozes and Jules Villard. Shared contract-obedient endpoints. In *ICE*, pages 17–31, 2012.
- [MAN17] William Mansky, Andrew W. Appel, and Aleksey Nogin. A verified messaging system. *PACMPL*, 1(OOPSLA):87:1–87:28, 2017.
- [MJ20] Yusuke Matsushita and Jacques-Henri Jourdan. Flexible number of logical steps per physical step, 2020. Iris merge request. URL: https://gitlab.mpi-sws.org/iris/iris/-/merge_requests/595.
- [MPU04] Claude Marché, Christine Paulin-Mohring, and Xavier Urbain. The KRAKATOA tool for certification of JAVA/JAVACARD programs annotated in JML. *JLP*, (1-2):89–106, 2004.
- [MY15] Dimitris Mostrous and Nobuko Yoshida. Session typing and asynchronous subtyping for the higher-order π -calculus. *Information and Computation*, 241:227–263, 2015.
- [MYH09] Dimitris Mostrous, Nobuko Yoshida, and Kohei Honda. Global principal typing in partially commutative asynchronous sessions. In *ESOP*, pages 316–332, 2009.
- [Nak00] Hiroshi Nakano. A modality for recursion. In *LICS*, pages 255–266, 2000.
- [NLSD14] Aleksandar Nanevski, Ruy Ley-Wild, Ilya Sergey, and Germán Andrés Delbianco. Communicating state transition systems for fine-grained concurrent resources. In *ESOP*, pages 290–310, 2014.

- [OBH16] Wytse Oortwijn, Stefan Blom, and Marieke Huisman. Future-based static analysis of message passing programs. In *PLACES*, pages 65–72, 2016.
- [OHT⁺11] Kosuke Ono, Yoichi Hirai, Yoshinori Tanabe, Natsuko Noda, and Masami Hagiya. Using coq in specification and program extraction of hadoop mapreduce applications. In *SEFM*, pages 350–365, 2011.
- [Pad14] Luca Padovani. Deadlock and lock freedom in the linear π -calculus. In *CSL*, pages 72:1–72:10, 2014.
- [Pad17] Luca Padovani. A simple library implementation of binary sessions. *JFP*, 27(2010):e4, 2017.
- [SB14] Kasper Svendsen and Lars Birkedal. Impredicative concurrent abstract predicates. In *ESOP*, pages 149–168, 2014.
- [SGG⁺21] Simon Spies, Lennard Gäher, Daniel Gratzer, Joseph Tassarotti, Robbert Krebbers, Derek Dreyer, and Lars Birkedal. Transfinite Iris: Resolving an existential dilemma of step-indexed separation logic. In *PLDI*, pages 80–95, 2021.
- [SRF⁺20] Nikhil Swamy, Aseem Rastogi, Aymeric Fromherz, Denis Merigoux, Danel Ahman, and Guido Martínez. SteelCore: An extensible concurrent separation logic for effectful dependently typed programs. *PACMPL*, 4(ICFP):121:1–121:30, 2020.
- [SSB16] Kasper Svendsen, Filip Sieczkowski, and Lars Birkedal. Transfinite step-indexing: Decoupling concrete and logical steps. In *ESOP*, pages 727–751, 2016.
- [SWT18] Ilya Sergey, James R. Wilcox, and Zachary Tatlock. Programming and proving with distributed protocols. *PACMPL*, 2(POPL):28:1–28:30, 2018.
- [TCP11] Bernardo Toninho, Luís Caires, and Frank Pfenning. Dependent session types via intuitionistic linear type theory. In *PPDP*, pages 161–172, 2011.
- [TCP14] Bernardo Toninho, Luís Caires, and Frank Pfenning. Corecursion and non-divergence in session-typed processes. In *TGC*, volume 8902 of *LNCS*, pages 159–175, 2014.
- [TDJ13] Samira Tasharofi, Peter Dinges, and Ralph E. Johnson. Why do scala developers mix the actor model with other concurrency models? In *ECOOP*, pages 302–326, 2013.
- [TJH17] Joseph Tassarotti, Ralf Jung, and Robert Harper. A higher-order logic for concurrent termination-preserving refinement. In *ESOP*, pages 909–936, 2017.
- [TLSZ19] Tengfei Tu, Xiaoyu Liu, Linhai Song, and Yiyang Zhang. Understanding real-world concurrency bugs in go. In *ASPLOS*, pages 865–878, 2019.
- [TV20] Peter Thiemann and Vasco T. Vasconcelos. Label-dependent session types. *PACMPL*, 4(POPL):67:1–67:29, 2020.
- [TY18] Bernardo Toninho and Nobuko Yoshida. Depending on session-typed processes. In *FOSSACS*, volume 10803 of *LNCS*, pages 128–145, 2018.
- [Vas12] Vasco Thudichum Vasconcelos. Fundamentals of session types. *Information and Computation*, 217:52–70, 2012.
- [VLC09] Jules Villard, Étienne Lozes, and Cristiano Calcagno. Proving copyless message passing. In *APLAS*, pages 194–209, 2009.
- [ZFH⁺20] Fangyi Zhou, Francisco Ferreira, Raymond Hu, Rumyana Neykova, and Nobuko Yoshida. Statically verified refinements for multiparty protocols. *PACMPL*, 4(OOPSLA), 2020.