



Delft University of Technology

Actris: Session-type based reasoning in separation logic

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

DOI

[10.1145/3371074](https://doi.org/10.1145/3371074)

Publication date

2020

Document Version

Final published version

Published in

Proceedings of the ACM on Programming Languages

Citation (APA)

Hinrichsen, J. K., Bengtson, J., & Krebbers, R. (2020). Actris: Session-type based reasoning in separation logic. *Proceedings of the ACM on Programming Languages*, 4(POPL), 1-30. Article 6.
<https://doi.org/10.1145/3371074>

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: Session-Type Based Reasoning in Separation Logic

JONAS KASTBERG HINRICHSSEN, IT University of Copenhagen, Denmark

JESPER BENGTON, IT University of Copenhagen, Denmark

ROBBERT KREBBERS, Delft University of Technology, The Netherlands

Message passing is a useful abstraction to implement 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 distributed merge sort, a distributed load-balancing mapper, and a variant of the map-reduce model, using relatively simple specifications. Soundness of Actris is proved using a model of its protocol mechanism in the Iris framework. We mechanised the theory of Actris, together with tactics for symbolic execution of programs, as well as all examples in the paper, in the Coq proof assistant.

CCS Concepts: • **Theory of computation** → **Separation logic**; *Program verification*; Programming logic.

Additional Key Words and Phrases: Message passing, actor model, concurrency, session types, Iris

ACM Reference Format:

Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. 2020. Actris: Session-Type Based Reasoning in Separation Logic. *Proc. ACM Program. Lang.* 4, POPL, Article 6 (January 2020), 30 pages. <https://doi.org/10.1145/3371074>

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 [Hewitt et al. 1973] 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.

While message passing is a useful abstraction, it is by no means a silver bullet of concurrent programming. In a qualitative study of larger Scala projects Tasharofi et al. [2013] 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.

Authors' addresses: Jonas Kastberg Hinrichsen, IT University of Copenhagen, Denmark, jkas@itu.dk; Jesper Bengtson, IT University of Copenhagen, Denmark, jebe@itu.dk; Robbert Krebbers, Delft University of Technology, The Netherlands, mail@robbertkrebbers.nl.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2020 Copyright held by the owner/author(s).

2475-1421/2020/1-ART6

<https://doi.org/10.1145/3371074>

For 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. \[2019\]](#) 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 [[Jung et al. 2016, 2018b, 2015](#); [Krebbers et al. 2017a](#)]. 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 the affine variant [[Mostrous and Vasconcelos 2014](#)] of binary session types [[Honda et al. 1998](#)]. We show that dependent separation protocols integrate seamlessly with other concurrency paradigms, allow delegation of linear resources, support channel sharing over multiple concurrent threads using locks, and more.

1.1 Message Passing in Concurrent Separation Logic

Over the last few years, there has been a variety of work on extensions of concurrent separation logic with high-level reasoning principles for message passing [[Craciun et al. 2015](#); [Francalanza et al. 2011](#); [Lozes and Villard 2012](#); [Oortwijn et al. 2016](#)]. These logics typically include some form of mechanism for writing protocol specifications in a high-level manner. Unfortunately, these logics have shortcomings in terms of expressivity. Most importantly, they cannot be used to reason about programs that combine message-passing with other programming and concurrency paradigms, such as higher-order functions, fine-grained shared-memory concurrency, and locks.

In a different line of work, researchers have developed expressive extensions of concurrent separation logic that do support proving strong specifications of programs involving some or all combinations of the aforementioned programming and concurrency paradigms. Examples of such logics are TaDA [[da Rocha Pinto et al. 2014](#)], iCAP [[Svendsen and Birkedal 2014](#)], Iris [[Jung et al. 2015](#)], FCSL [[Nanevski et al. 2014](#)], and VST [[Appel 2014](#)]. However, only a few variants and extensions of these logics address message-passing concurrency.

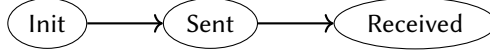
First, there has been work on the use of separation logic to reason about programs that communicate via message passing over a network. The reasoning principles in such logics are geared towards different programming patterns than the ones used in high-level languages like Erlang, Elixir, Go, and Scala. 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 the Disel logic by [Sergey et al. \[2018\]](#) and the Aneris logic by [Krogh-Jespersen et al. \[2019\]](#). Second, there has been work on the use of separation logic to prove compiler correctness of high-level message-passing languages. [Tassarotti et al. \[2017\]](#) 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 expressive 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. \[2017\]](#):

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

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 being 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 Dependent Separation Protocols

Instead of using STSs, we extend separation logic with a new notion called *dependent separation protocols*. This notion is inspired by the session type community, pioneered by [Honda et al. \[1998\]](#), where channel endpoints are given types that describe the expected exchanges. Using binary session types, the channels c and c' in the example above 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 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. \[2010\]](#) and [Craciun et al. \[2015\]](#) 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 $c \multimap \text{prot}$, 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. The 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:

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

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 show that $prog_1$ has the specification $\{ \text{True} \} prog_1 \{ v. v = 42 \}$, where v is the result of the evaluation.

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  $\triangleq$  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 2 to it, and sends it back. This program has the same specification as *prog₁*, while we change the dependent separation protocol as follows (we omit the protocol for the dual endpoint):

$$c \mapsto ! x \langle x \rangle \{True\}. ? \langle x + 2 \rangle \{True\}. end$$

This protocol states that the second exchanged value is exactly the first with 2 added to it. To do so, it makes use of dependency—the variable x , which is used to describe the contents of the first message, is also used to describe the contents of the second message. 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. [2015], which does not support dependent protocols. Their logic is limited to protocols analogous to $! \langle x \rangle \{True\}. ? \langle x + 2 \rangle \{True\}. end$ where x is free, which means the value of x must be known at the point the protocol is created.

While the previous examples could also have been type-checked and verified using the formalisms of Bocchi et al. [2010] and Craciun et al. [2015] the following stateful example cannot:

```
prog3  $\triangleq$  let (c, c') = new_chan () in
  fork {let  $\ell$  = recv c' in  $\ell \leftarrow ! \ell + 2$ ; send c' ()};
  let  $\ell$  = ref 40 in send c  $\ell$ ; recv c; !  $\ell$ 
```

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 *prog₁* and *prog₂*, but the dependent separation protocols is updated as follows:

$$c \mapsto ! \ell x \langle \ell \rangle \{ \ell \mapsto x \}. ? \langle () \rangle \{ \ell \mapsto x + 2 \}. end$$

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. [2010] 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. [2015] cannot verify this program because it does not support dependent protocols. Dependent protocols are crucial here—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 *prog₁*:

```
prog4  $\triangleq$  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 *prog₄* has the same specification as before, but now we use the

following recursive dependent separation protocol:

$$c \mapsto \mu \text{prot}. !x \langle x \rangle \{ \text{True} \}. ? \langle x + 2 \rangle \{ \text{True} \}. \text{prot}$$

This protocol expresses that it is possible to make repeated exchanges with the forked-off thread to increment a number by 2. 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 captures some features of conventional session types. One such example is the *delegation* of channels as seen in the following program:

```
prog5  $\triangleq$  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. However, 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 rather intricate, the following dependent separation protocols describe the communication concisely:

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

The first protocol states that the initial value sent must be a channel endpoint with the protocol of prog_1 , meaning that the main thread must give up the ownership of the channel endpoint c_2 , thereby delegating it. It then expects a value y to be sent, and finally to receive a notification $()$, and also to reacquire ownership of the channel c_2 , which has since taken one step by sending y .

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  $\triangleq$  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:

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

The $!$ 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 2 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.

While it has not been captured in the above examples, protocols are closed under composition and branching. 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 Sections 3, is especially useful when combining channels with locks.

1.3 Contributions and Outline

This paper introduces **Actris**: a higher-order impredicative concurrent separation logic build on top of the Iris framework for reasoning about functional correctness of message-passing programs 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 branching, recursion, and delegation (Section 2).
- 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 (Section 3).
- 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 [2004] (Section 4).
- We give a model of dependent separation protocols in the Iris framework. Using this model we obtain safety (*i.e.*, session fidelity) and postcondition validity of our Hoare triples (Section 5).
- We provide a full mechanisation of Actris [Hinrichsen et al. 2019] using the interactive theorem prover Coq. On top of that, we provide tactics for symbolic execution of dependent separation protocols and mechanise all the examples in the paper (Section 6).

2 A TOUR OF ACTRIS

This section demonstrates the core features of Actris. We first introduce the language (Section 2.1) and the logic (Section 2.2). We then introduce and iteratively extend a simple distributed merge sort algorithm to demonstrate the main features of Actris (Section 2.3–2.8). 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).

2.1 The Actris Language

The language used throughout the paper is an untyped functional language with higher-order functions, higher-order mutable references, fork-based concurrency, and primitives for message-passing over bidirectional asynchronous channels. The syntax is as follows:

$$\begin{aligned}
 v \in \text{Val} &::= () \mid i \mid b \mid \ell \mid c \mid \text{rec } f(x) = e \mid \dots & (i \in \mathbb{Z}, b \in \mathbb{B}, \ell \in \text{Loc}, c \in \text{Chan}) \\
 e \in \text{Expr} &::= v \mid x \mid \text{rec } f(x) = e \mid e_1(e_2) \mid \text{ref } e \mid !e \mid e_1 \leftarrow e_2 \mid \\
 &\quad \text{fork } \{e\} \mid \text{new_chan } () \mid \text{send } e_1 e_2 \mid \text{recv } e \mid \dots
 \end{aligned}$$

We omit the usual operations on pairs, sums, lists, and integers, which are standard. We introduce the following syntactic sugar: lambda abstractions $\lambda x. e$ are defined as $\text{rec } _ (x) = e$, let-bindings $\text{let } x = e_1 \text{ in } e_2$ are defined as $(\lambda x. e_2) e_1$, and sequencing $e_1; e_2$ is defined as $\text{let } _ = e_1 \text{ in } e_2$.

The language features the usual operations for heap manipulation. New references can be created using $\text{ref } e$, dereferenced using $!e$, and assigned to using $e_1 \leftarrow e_2$. Concurrency is supported via $\text{fork } \{e\}$, which spawns a new thread e that is executed in the background. The language also supports atomic operations like compare-and-set (CAS), which can be used to implement lock-free data structures and synchronisation primitives, but these are omitted from the syntax.

The language supports message passing through bidirectional channels, which are represented using pairs of buffers $(\vec{v}_1, \vec{v}_2)$ of unbounded size. The $\text{new_chan } ()$ 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 other's send buffer and vice versa. That means, $\text{send } c_i v$ enqueues the value v in its own buffer, *i.e.*, $\vec{v}_i$, and $\text{recv } c_i$ dequeues a value from the other buffer,

i.e., from $\vec{v}_2$ if $i = 1$ and from $\vec{v}_1$ if $i = 2$. Message passing is asynchronous, meaning that `send` c v will always reduce, while `recv` c will block as long as the receiving buffer is empty.

Throughout the paper, we often use the following syntactic sugar to encapsulate the common behaviour of starting a new process:

$$\text{start } e \triangleq \text{let } f = e \text{ in let } (c, c') = \text{new_chan } () \text{ in fork } \{f \ c'\}; c$$

Here, e should evaluate to a function that takes a channel endpoint.

2.2 The Actris Logic

Actris is a higher-order impredicative concurrent separation logic with a new notion called *dependent separation protocols* to reason about message-passing concurrency. As we will show in Section 5, Actris is built as a library on top of the Iris framework [Jung et al. 2016, 2018b, 2015; Krebbers et al. 2017a] and thus inherits all features of Iris. For the purpose of this section, no prior knowledge of Iris is expected as the majority of Iris's features are orthogonal to Actris's. At this point, we are primarily concerned with Iris's support for nested Hoare triples and guarded recursion, which we need to transfer functions over channels (Section 2.4) and to define recursive protocols (Section 2.6). An extensive overview of Iris can be found in [Jung et al. 2018b].

The grammar of Actris and a selection of its rules are displayed in Figure 1. The Actris grammar includes the polymorphic lambda-calculus¹ with a number of primitive types and terms operating on these types. Most important is the type `iProp` of propositions and the type `iProto` of dependent separation protocols. 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, the meta variable $prot$ for protocols, the meta variable v for values, and the meta variables t and u for general terms of any type. Apart from that, there is the implicit side-condition that recursive predicates defined using the recursion operator $\mu x : \tau. t$ should be *guarded*. That means, the variable x should appear under a *contractive* term construct. As is usual in logics with guarded recursion [Nakano 2000], the later $\triangleright$ modality is contractive so one can define recursive predicates. But moreover, as we will demonstrate in Section 2.6, the constructors $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$ of dependent separation protocols are contractive in the arguments P and $prot$ to enable the construction of recursive protocols. The rule $\mu\text{-UNFOLD}$ says that $\mu x : \tau. t$ is in fact a fixpoint of t .

In order to express program specifications, Actris features Hoare triples $\{P\} e \{v. Q\}$, where P is the precondition and Q the postcondition. The binder v can be used to talk about the return value of e in the postcondition Q , but may be omitted whenever the result is $()$. Note that Hoare triples are propositions of the logic themselves (i.e., they are of type `iProp`), so they can be nested to express specifications of higher-order functions. The rules for Hoare triples are mostly standard, but it is worth pointing out the rule `HT-REC` for recursive functions. This rule has a later $\triangleright$ modality in the precondition, which when combined with the Löb rule allows one to reason about general recursive functions. As usual, the *points-to* connective $\ell \mapsto v$ expresses unique ownership of a location ℓ with value v . Since we consider a garbage collected language, one can discard arbitrary separation logic resources via the rule `AFFINE`.

The novel feature of Actris is its support for dependent separation protocols to reason about message-passing programs. This is done using the $c \multimap prot$ connective, which expresses unique ownership of a channel endpoint c and states that the endpoint follows the protocol $prot$. 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 proposition P denotes the

¹ Actris and Iris, which are both formalised as a shallow embedding in Coq, have in fact a predicative Type hierarchy, while propositions `iProp` are impredicative. For brevity's sake, we omit details about predicativity of Type, as they are standard.

Grammar:

$$\begin{aligned}
\tau, \sigma &::= x \mid 0 \mid 1 \mid \mathbb{B} \mid \mathbb{N} \mid \mathbb{Z} \mid \text{Type} \mid \forall x : \tau. \sigma \mid \\
&\quad \text{Loc} \mid \text{Chan} \mid \text{Val} \mid \text{Expr} \mid \text{iProp} \mid \text{iProto} \mid \dots \\
t, u, P, Q, \text{prot} &::= x \mid \lambda x : \tau. t \mid t(u) \mid t(\tau) \mid \quad (\text{Polymorphic lambda-calculus}) \\
&\quad \text{True} \mid \text{False} \mid P \wedge Q \mid P \vee Q \mid P \Rightarrow Q \mid \quad (\text{Propositional logic}) \\
&\quad \forall x : \tau. P \mid \exists x : \tau. P \mid t = u \mid \quad (\text{Higher-order logic with equality}) \\
&\quad \mu x : \tau. t \mid \triangleright P \mid \quad (\text{Guarded recursion}) \\
&\quad P * Q \mid P \multimap Q \mid \ell \mapsto v \mid \{P\} e \{v. Q\} \mid \quad (\text{Separation logic}) \\
&\quad c \multimap \text{prot} \mid \overline{\text{prot}} \mid \text{prot}_1 \cdot \text{prot}_2 \mid \text{end} \mid \\
&\quad !\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot} \mid ?\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot} \mid \dots \quad (\text{Dependent separation protocols})
\end{aligned}$$
Ordinary affine separation logic:

$ \begin{array}{c} \text{HT-FRAME} \\ \frac{\{P\} e \{w. Q\}}{\{P * R\} e \{w. Q * R\}} \end{array} $	$ \begin{array}{c} \text{HT-VAL} \\ \{\text{True}\} v \{w. w = v\} \end{array} $	$ \begin{array}{c} \text{HT-FORK} \\ \frac{\{P\} e \{\text{True}\}}{\{P\} \text{fork } \{e\} \{w. w = ()\}} \end{array} $
$ \begin{array}{c} \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} $		

Recursion:

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

Heap manipulation:

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

Message passing:

$\{\text{True}\} \text{new_chan } () \{(c, c'). c \multimap \text{prot} * c' \multimap \overline{\text{prot}}\}$	(HT-NEWCHAN)
$\{c \multimap !\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot} * P[\vec{t}/\vec{x}]\} \text{send } c (v[\vec{t}/\vec{x}]) \{c \multimap \text{prot}[\vec{t}/\vec{x}]\}$	(HT-SEND)
$\{c \multimap ?\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot}\} \text{recv } c \{w. \exists \vec{y}. (w = v[\vec{y}/\vec{x}]) * c \multimap \text{prot}[\vec{y}/\vec{x}] * P[\vec{y}/\vec{x}]\}$	(HT-RECV)

Dependent separation protocols:

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

Fig. 1. The grammar and a selection of rules of Actris.

```

sort_service cmp c  $\triangleq$ 
  let l = recv c in
  if |l| ≤ 1 then send c () else
  let l' = split 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;
  merge cmp l l'; send c ()

sort_client cmp l  $\triangleq$ 
  let c = start sort_service cmp in
  send c l;
  recv c

```

Fig. 2. A distributed merge sort algorithm (the code for merge and split is standard and thus elided).

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 *prot*. For example, $!(b : \mathbb{B}) (\ell : \text{Loc}) (i : \mathbb{N}) \langle (b, \ell) \rangle \{ \ell \mapsto i * 10 < i \}$. *prot* expresses that a pair of a Boolean and an integer reference whose value is at least 10 is sent. 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. The $\overline{\text{prot}}$ connective denotes the *dual* of a protocol. As with 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 connective $\text{prot}_1 \cdot \text{prot}_2$ *composes* the protocols prot_1 and prot_2 , which is achieved by substituting any **end** in prot_1 with prot_2 .

The rule **HT-NEWCHAN** allow ascribing any protocol to newly created channels using `new_chan ()`, obtaining ownership of $c \mapsto \text{prot}$ and $c' \mapsto \overline{\text{prot}}$ 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 (*i.e.*, session fidelity, see Section 5.4).

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 send to match up with the ascribed value. Concretely, to send a message w , one need to give up ownership of $c \mapsto !\vec{x}:\vec{\tau} \langle v \rangle \{P\}$. *prot*, pick an appropriate instantiation $\vec{t}$ for the variables $\vec{x}:\vec{\tau}$ so that $w = v[\vec{t}/\vec{x}]$, and give up ownership of the associated resources $P[\vec{t}/\vec{x}]$. Subsequently, one gets back ownership of the protocol tail $\text{prot}[\vec{t}/\vec{x}]$.

The rule **HT-RECV** for `recv c` is essentially dual to the rule **HT-SEND**. One needs to give up ownership of $c \mapsto ?\vec{x}:\vec{\tau} \langle v \rangle \{P\}$. *prot*, and in return gets acquires 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 $\text{prot}[\vec{y}/\vec{x}]$, where $\vec{y}$ are instances of the variables of the protocol.

2.3 Basic Protocols

In order to show the basic features of dependent separation protocols, we will prove the functional correctness of a simple distributed merge sort algorithm, whose code is shown in Figure 2.

The function `sort_client` takes a comparison function *cmp* and a reference to a linked list l that will be sorted using merge sort. The bulk of the work is done by the `sort_service` function that is parameterised by a channel c over which it receives a reference to the linked list to be sorted. If the list is an empty or singleton list, which is trivially sorted, the function immediately sends back a unit value () to inform the caller that the work has been completed, and terminates. Otherwise, the list is split into two partitions using the `split` function, which updates the list in-place so that ℓ points to the first partition, and returns a reference l' to the second partition. These partitions are recursively sorted using two newly started instances of `sort_service`. The results of the processes

$$\begin{array}{ll}
\text{sort_service}_{\text{func}} c \triangleq & \text{sort_client}_{\text{func}} \text{ cml } \triangleq \\
\text{let } \text{cmp} = \text{recv } c \text{ in} & \text{let } c = \text{start sort_service}_{\text{func}} \text{ in} \\
\text{sort_service } \text{cmp } c & \text{send } c \text{ cmp; send } c \text{ l; recv } c
\end{array}$$

Fig. 3. A version of the sort service that receives the comparison function over the channel.

are then requested and merged using the merge function, which updates the list in-place so that l points to the merged list. Finally, the unit value $()$ is sent back along the original channel 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:

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

Here, R is a decidable total relation on an implicit polymorphic 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:

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

The channel c adheres to the following dependent separation protocol:

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

The protocol describes the interaction of sending a list reference, and then receiving a unit value $()$ once the list is sorted and the reference is updated to point to the sorted list. The predicate $\text{sorted_of}_R \vec{x} \vec{y}$ is true iff $\vec{x}$ is a sorted version of $\vec{y}$ with respect to the relation R . The specification of the service and the client is as follows:

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

There are two important things to note about these specifications. First, the protocol sort_prot is written from the point of view of the client. As such, the precondition for sort_service requires that c follows the dual. Second, the pre- and postcondition of sort_service are generalised to have an arbitrary protocol prot composed at the end. It is important to write specifications this way, so they can be embedded in other protocols. We will see examples of that in Section 2.6 and 2.7.

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

2.4 Transferring Functions

The distributed sort_service from the previous section (Figure 2) is parametric on a comparison function. To demonstrate Actris's support for reasoning about functions transferred over channels, we verify the correctness of the program $\text{sort_service}_{\text{func}}$ in Figure 3, which receives the comparison function over the channel instead of via a lambda abstraction. To verify this program, we

extend the protocol `sort_prot` as follows:

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

The new protocol captures that we first send the comparison function `cmp`. In addition, 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:

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

2.5 Branching

Branching behaviour is common in message-passing communication protocols and is readily available in Actris using dependent separation protocols. Branching is encoded in terms of sending and receiving a Boolean value that is matched using an if-then-else construct:

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

We let `left` $\triangleq$ `true` and `right` $\triangleq$ `false` to be used together with `select` for the sake of readability. Due to the higher-order nature of Actris, the usual protocol specifications for branching from session types can be encoded as regular logical branching within the protocols:

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

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`:

HT-SELECT

$$\{c \mapsto \text{prot}_1 \{Q_1\} \oplus \{Q_2\} \text{ prot}_2 * \text{if } b \text{ then } Q_1 \text{ else } Q_2\} \text{ select } c \ b \{c \mapsto \text{if } b \text{ then } \text{prot}_1 \text{ else } \text{prot}_2\}$$

HT-BRANCH

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

Apart from branching on Boolean values, one can use dependent separation protocols to encode branching on any enumeration type (e.g., lists, natural numbers, days of the week, etc.).

2.6 Recursive Protocols

We will now use branching 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 4. The service `sort_servicerec` contains a loop in which branching is used to either terminate the service, or to sort an individual list using the distributed merge sort algorithm `sort_service` from Section 2.3. The client `sort_clientrec` uses the service to sort a nested linked list l of linked lists. It performs this job by starting a single instance of the service at 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.

```

sort_servicerec cmp c  $\triangleq$ 
  branch c with
    left  $\Rightarrow$  sort_service cmp c;
           sort_servicerec cmp c
    | right  $\Rightarrow$  ()
  end

sort_clientrec cmp l  $\triangleq$ 
  let c = start sort_servicerec cmp in
    iter ( $\lambda l'$ . select c left; send c l'; recv c) l;
    select c right

```

Fig. 4. A recursive version of the sort service that can perform multiple jobs in sequence.

A protocol for interacting with the sorting service can be defined as follows:

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

The protocol uses the branching 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 dictates that one can proceed recursively.

It is important to point out that—as is usual in logics with guarded recursion [Nakano 2000]—the variable x should appear under a *contractive* term construct in $\mu x : \tau. t$. In our protocol, the recursive variable rec appears under the argument of $\oplus$, which is defined in terms of $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$, which, similar to $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$, is contractive in the arguments P and $prot$.

The specifications of the service and the client are as follows:

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

We let $J \triangleq \lambda \ell' \vec{y}. \ell' \mapsto_J \vec{y}$ 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 post-composition.

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 $\text{sort_client}_{\text{func}}$, or in a way such that a different comparison function can be used for each sorting job.

2.7 Delegation

Delegation 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. Due to the impredicativity of protocols in Actris, reasoning about programs that make use of delegation is readily available. The protocols $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$ and $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}. prot$ can simply refer to the ownership of protocols $c \mapsto prot'$ in the proposition P .

An example of a program that uses delegation is the $\text{sort_service}_{\text{del}}$ variant of the recursive sorting service in Figure 5, which allows one to perform multiple sorting jobs in parallel. To enable parallelism, it delegates a new channel c' to an inner sorting service for each sorting job.

The client $\text{sort_client}_{\text{del}}$ once again uses the sorting service to sort a nested linked list l of linked lists. The client starts a connection c to the new service, and for each inner list l' , it acquires a delegated channel c' , over which it sends a pointer l' to the inner list that should be sorted.

```

sort_servicedel cmp c  $\triangleq$ 
  branch c with
    left  $\Rightarrow$ 
      let c' = start sort_service cmp in
      send c c';
      sort_servicedel cmp c
    | right  $\Rightarrow$  ()
  end

sort_clientdel cmp l =
  let c = start sort_servicedel cmp in
  let k = new_list () in
  iter ( $\lambda l'. \text{select } c \text{ left};$ 
        let c' = recv c in
        push c' k; send c' l') l
  send c right;
  iter recv k

```

Fig. 5. A recursive version of the sort service that uses delegation to perform multiple jobs in parallel. The code for the function push, which pushes an element to the head of a list, has been elided.

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 `iter 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:

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

The specifications of the service and the client are as follows:

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

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

2.8 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 distributed merge sort service. This version $\text{sort_service}_{\text{fg}}$, as shown in Figure 6, requires the input list to be sent element by element, after which the service sends the sorted list back in the same fashion. We use branching to indicate whether the whole list has been sent (right) or another element remains to be sent (left).

The structure of $\text{sort_service}_{\text{fg}}$ 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, as they are trivially sorted. The inductive case is handled by starting two sub-services at c_1 and c_2 that are sent the two initially received elements, respectively, after which the split_{fg} function is used to receive and forward the remaining values to the sub-services alternatingly. Once the right flag is received, the split_{fg} function terminates, and the algorithm moves to the second phase in which the merge_{fg} function merges the stream of values returned by the sub-services and forwards them to the parent service.

```

sort_servicefg cmp c  $\triangleq$ 
  branch c with
    right  $\Rightarrow$  select c right
  | left  $\Rightarrow$ 
    let x1 = recv c in
    branch c with
      right  $\Rightarrow$  select c left; send c x1;
        select c right
    | left  $\Rightarrow$ 
      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  $\triangleq$ 
  branch c with
    right  $\Rightarrow$  select c1 right; select c2 right
  | left  $\Rightarrow$ 
    let x = recv c in
    select c1 left; send c1 x; splitfg c c2 c1
  end
end

mergefg cmp c c1 c2  $\triangleq$ 
  branch c with
    right  $\Rightarrow$  assert false
  | left  $\Rightarrow$ 
    let x = recv c1 in
    mergefgaux cmp c x c1 c2
  end
end

mergefgaux cmp c x c1 c2  $\triangleq$ 
  branch c2 with
    right  $\Rightarrow$  select c left; send c x1;
      transfer c1 c
  | left  $\Rightarrow$ 
    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
end

sort_clientfg cmp l  $\triangleq$ 
  let c = start sort_servicefg cmp in
  send_all c l; recv_all c l

```

Fig. 6. 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).

The merge_{fg} function initially acquires the first value x from the first sub-service, which it uses in the recursive call as the current largest value. The recursive function $\text{merge}_{fg}^{\text{aux}}$ recursively requests a value y from the sub-service of which the current largest value was not acquired from. It then compares x and y using the comparison function cmp , and forwards the smallest element. This is repeated until the right flag is received from either sub-service, after which the remaining values of the other are forwarded to the parent service using the transfer function.

The interface of the client sort_client_{fg} for this service is similar to the previous ones. It takes a reference to a linked list, which is then sorted. It performs this task by sending the elements of the linked list to the sort service using the send_all function, and puts the received values back into the linked list using the recv_all function.

A suitable protocol for proving functional correctness of the fine-grained sorting service is:

```

sort_protfg (I : T → Val → iProp) (R : T → T →  $\mathbb{B}$ )  $\triangleq$  sort_protfghead I R  $\epsilon$ 
sort_protfghead (I : T → Val → iProp) (R : T → T →  $\mathbb{B}$ )  $\triangleq$   $\mu$  (rec : List T → iProto).
   $\lambda \vec{x}. (! (x : T) (v : Val) \langle v \rangle \{I x v\}. \text{rec } (\vec{x} \cdot [x]) \}) \oplus \text{sort\_prot}_{fg}^{\text{tail}} I R \vec{x} \epsilon$ 
sort_protfgtail (I : T → Val → iProp) (R : T → T →  $\mathbb{B}$ )  $\triangleq$   $\mu$  (rec : List T → List T → iProto).
   $\lambda \vec{x} \vec{y}. (? (y : T) (v : Val) \langle v \rangle \{(\forall i < |\vec{y}|. R \vec{y}_i y) * I y v\}. \text{rec } \vec{x} (\vec{y} \cdot [y]) \}) \{ \text{True} \} \& \{ \vec{x} \equiv_p \vec{y} \} \text{ end}$ 

```

The protocol is split into two phases $\text{sort_prot}_{fg}^{\text{head}}$ and $\text{sort_prot}_{fg}^{\text{tail}}$, mimicking the behaviour of the program. The $\text{sort_prot}_{fg}^{\text{head}}$ 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 $\text{sort_prot}_{fg}^{\text{tail}}$ 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.

The top-level specification of the service and client are similar to the specifications of the coarse grained version of distributed merge sort:

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

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.

3 MANIFEST SHARING VIA LOCKS

Since dependent separation protocols and the connective $c \mapsto \text{prot}$ 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* [Balzer and Pfenning 2017; Balzer et al. 2019], where channel endpoints are shared between multiple parties. Instead of having to extend Actris, we make use of locks and ghost state that Actris readily inherits from Iris. We present the basic idea with a simple introductory example of sharing a channel endpoint between two parties (Section 3.1). We then consider a more challenging example of a distributed load-balancing mapper (Section 3.2).

3.1 Locks and Ghost State

Using the language from Section 2.1 one can implement locks using a spin lock, ticket lock, or a more sophisticated implementation. For the purpose of this paper, we abstract over the concrete implementation and assume that we have operations `new_lock`, `acquire` and `release` that satisfy the common separation logic specifications for locks as shown in Figure 7.

The `new_lock ()` operation creates a new lock, which can be thought of as a mutex. The operation `acquire lk` will atomically take the lock or block in the case the lock is already taken, and `release lk` releases the lock so that it may be acquired by other threads. The specifications in Figure 7 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

$\{R\} \text{ new_lock } () \{lk. \text{is_lock } lk \ R\}$	(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)
$\text{is_lock } lk \ R * \text{is_lock } lk \ R * \text{is_lock } lk \ R$	(LOCK-DUP)

Fig. 7. The rules Actris inherits from Iris for locks.

```

prog_lock  $\triangleq$  let  $c = \text{start}$  ( $\lambda c. \text{let } lk = \text{new\_lock } () \text{ in}$ 
                        fork {acquire  $lk$ ; send  $c$  21; release  $lk$ };
                        acquire  $lk$ ; send  $c$  21; release  $lk$ ) in
recv  $c + \text{recv } c$ 

```

Fig. 8. A sample program that combines locks and channels to achieve manifest sharing.

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

Fig. 9. The authoritative contribution ghost theory.

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.

To show how locks can be used, consider the program in Figure 8, which uses a lock to share a channel endpoint between two threads that each send the integer 21 to the main thread. The following dependent protocol, where n denotes the number of messages that should be exchanged, captures the expected interaction from the point of view of the main thread:

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

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 over time depending on the values that are sent. To tie the number n to the number of contributions made by the threads that share the channel endpoint, we make use of the connectives $\text{auth}_\gamma n$ and contrib_γ , which are defined using Iris’s “ghost theory” mechanism for “user-defined” ghost state [Jung et al. 2018b, 2015].

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. These concepts are made precise by the rules in Figure 9. The rule AUTH-INIT expresses that an authority $\text{auth}_\gamma 0$ can always be created, which is given some fresh ghost identifier γ . Using the

rules **AUTH-ALLOC** and **AUTH-DEALLOC**, one can allocate and deallocate tokens contrib_γ as long as the count 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.

Most of the rules in Figure 9 involve the logical connective $\Rightarrow$ of a so-called *view shift*. The view shift connective, which Actris inherits from Iris, can be thought of as a “ghost update”, which is made precise by the structural rules **Vs-CSQ** and **Vs-FRAME** rules, that establish the connection between $\Rightarrow$ and the Hoare triples of the logic:

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

With the ghost state in place, we can now state 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:

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

To establish the initial lock invariant, we use the rules **AUTH-INIT** and **AUTH-ALLOC** to create the authority $\text{auth}_\gamma 2$ and two contrib_γ tokens. The contrib_γ 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 $\overline{\text{lock_prot } n}$ 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 $\overline{\text{lock_prot } n}$ to $\overline{\text{lock_prot } (n-1)}$. Subsequently, the sending threads can deallocate the token contrib_γ (using **AUTH-DEALLOC**) to decrement the n of $\text{auth}_\gamma n$ accordingly to restore the lock invariant.

3.2 A Distributed 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 distributed load-balancing mapper that maps a function f over a list. Our distributed mapper consists of one client that distributes the work, and a number of workers that perform the function f on individual elements x 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, which can be found in Figure 10, 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 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 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.²

²The interested reader can find the entire code in the accompanied Coq development [Hinrichsen et al. 2019].

```

mapper_worker  $f_v$   $lk$   $c \triangleq$ 
  acquire  $lk$ ; select  $c$  left;
  branch  $c$  with
    right  $\Rightarrow$  release  $lk$ 
  | left  $\Rightarrow$  let  $x = \text{recv } c$  in release  $lk$ ;      (* acquire work *)
    let  $y = f_v x$  in                             (* map it *)
    acquire  $lk$ ;
    select  $c$  right; send  $c$   $y$ ;                 (* send it back *)
    release  $lk$ ;
  mapper_worker  $f_v$   $lk$   $c$ 
end

```

Fig. 10. A worker of the distributed mapper service.

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

Fig. 11. The authoritative contribution ghost theory extended with multisets.

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

```

mapper_prot ( $I_T : T \rightarrow \text{Val} \rightarrow \text{iProp}$ ) ( $I_U : U \rightarrow \text{Val} \rightarrow \text{iProp}$ ) ( $f : T \rightarrow \text{List } U$ )  $\triangleq$ 
   $\mu$  ( $\text{rec} : \mathbb{N} \rightarrow \text{MultiSet } T \rightarrow \text{iProto}$ ).  $\lambda n X.$ 
    if  $n = 0$  then end else
      (! ( $x : T$ ) ( $v : \text{Val}$ )  $\langle v \rangle \{I_T x v\}. \text{rec } n (X \uplus \{x\})$ )  $\oplus \text{rec } (n - 1) X$ 
      {( $n=1 \Rightarrow (X=\emptyset)$ ) & {True}}
      ? ( $x : T$ ) ( $\ell : \text{Loc}$ )  $\langle \ell \rangle \{x \in X * \ell \mapsto_{I_U} (f x)\}. \text{rec } n (X \setminus \{x\})$ 

```

The protocol is parameterised by representation predicates I_T and I_U that relate language values to elements of type T and U in the logic, as well as a function $f : T \rightarrow \text{List } U$ that specifies the behaviour of the language-level function f_v . The connection between f and f_v is formalised as:

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

Similar to `lock_prot` from Section 3.1, `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 $(n = 1) \Rightarrow (X = \emptyset)$ on the branch (&) 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 Section 3.1. 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 11. 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 $\text{contrib}_\gamma Y$ can be updated as long as it is done in accordance with the multiset X of $\text{auth}_\gamma n X$. Finally, the **AUTHM-CONTRIB-AGREE** rule expresses that the multiset Y of $\text{contrib}_\gamma Y$ must be a subset of the multiset X of $\text{auth}_\gamma n X$, while the stricter rule **AUTHM-CONTRIB-AGREE1** asserts equality between X and Y when only one contribution remains.

The specifications of `mapper_worker` and a possible top-level client `mapper_client` that uses n workers to map f_v over the linked list ℓ are as follows:

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

$$\text{mapper_worker } f_v \ lk \ c \quad \text{mapper_client } n \ f_v \ \ell$$

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

The lock invariant and specification of `mapper_worker` are similar to those used in the simple example in Section 3.1. The specification of `mapper_client` $n \ f_v \ \ell$ simply states that the resulting linked list points to 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 `mapper_prot` $I_T \ I_U \ f \ n$, where n is the initial number of workers. Subsequently, we use the rules **AUTHM-INIT** and **AUTHM-ALLOC** to create the authority $\text{auth}_\gamma n \emptyset$ and n tokens $\text{contrib}_\gamma \emptyset$, 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 $\text{contrib}_\gamma \emptyset$, 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.

4 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 distributed load-balancing implementation of the map-reduce model by [Dean and Ghemawat \[2004\]](#).

Since Actris is not concerned with distributed systems over networks, we consider a version of map-reduce that distributes 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 distributed merge sort algorithm (Section 2.8) and the distributed load-balancing mapper (Section 3.2). 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 [[Hinrichsen et al. 2019](#)].

4.1 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 again 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:

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

4.2 Implementation of Map-Reduce

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

Our variant of the map-reduce model is defined as a function $\text{map_reduce}_v n m f_v g_v \ell$, which coordinates the work for performing map-reduce on a linked list ℓ between n mappers performing the “map” function f_v and m workers performing the “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 `mapper_worker` from Section 3, parameterised with the f_v function. Additionally start an instance of `sort_servicefg` from Section 2, 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 $\mathbb{Z}$ 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 `mapper_worker`, parameterised by g_v .
- (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.

4.3 Functional Correctness of Map-Reduce

The specification of the program is as follows:

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

The f_spec predicates (as introduced in Section 3.2) establish a connection between the functions f and g on the logical level and the functions f_v and g_v in the language. 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 $\ell \mapsto_{I_T} \vec{x}$ 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 Section 4.1.

5 THE MODEL OF ACTRIS

Actris is defined as an internal logic embedded in the Iris framework [Jung et al. 2016, 2018b, 2015; Krebbers et al. 2017a]. This means that the type `iProto` of dependent separation protocols and the connective $c \mapsto \text{prot}$ for the ownership of a channel endpoint are definitions in the Iris logic, and that the Actris proof rules are lemmas in the Iris logic. In this section we describe the relevant aspects of this embedding. First, we present our definitional semantics of bidirectional channels (Section 5.1). We then show how the type `iProto` is modelled (Section 5.2) and how the connective $c \mapsto \text{prot}$ for channel ownership is defined (Section 5.3). Finally, we show how adequacy (safety and postcondition validity of Hoare triples) follows from the embedding in Iris (Section 5.4).

5.1 Semantics of Channels

Since the Iris framework is parametric in the programming language that is being used, there are various approaches to extend it 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 [Jung et al. 2015] and by Tassarotti et al. [2017].
- 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. [2019] for a lock-free implementation of channels.

In this paper we went for the second approach. We used `HeapLang`, the default concurrent language shipped with Iris, and implemented bidirectional channels using a pair of linked-lists protected by a

lock. Although this implementation is not efficient, it has the benefit that it gives a clear declarative semantics that corresponds exactly to the intuitive semantics given in Section 2.1.

At the level of the logic, the state of both buffers of a bidirectional channel is described using the representation predicate $(c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2)$. We prove [Jacobs and Piessens \[2011\]](#)-style logically atomic specifications that roughly correspond to the following Hoare triples:

$$\begin{array}{ll} \{\text{True}\} \text{new_chan } () \{ (c_1, c_2). (c_1, c_2) \mapsto (\epsilon, \epsilon) \} \\ \{(c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2)\} \text{send } c_1 \ w \{ (c_1, c_2) \mapsto (\vec{v}_1 \cdot [w], \vec{v}_2) \} \\ \{(c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2)\} \text{send } c_2 \ w \{ (c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2 \cdot [w]) \} \\ \{(c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2)\} \text{recv } c_1 \{ w. (\vec{v}_2 = [w] \cdot \vec{w}) * (c_1, c_2) \mapsto (\vec{v}_1, \vec{w}) \} \\ \{(c_1, c_2) \mapsto (\vec{v}_1, \vec{v}_2)\} \text{recv } c_2 \{ w. (\vec{v}_1 = [w] \cdot \vec{w}) * (c_1, c_2) \mapsto (\vec{w}, \vec{v}_2) \} \end{array}$$

In Section 5.4 we show how Actris's rules for channels (**HT-NEWCHAN**, **HT-SEND** and **HT-RECV**) are derived from these logically atomic specifications. It is worth pointing out that the embedding of Actris only makes use of these specifications, and thus abstracts from the implementation details of the channel. This means that the channel implementation could be replaced with a more efficient version, or with a version that is natively integrated into the language.

5.2 The Model of Dependent Separation Protocols

Dependent separation protocols are streams of the form $!\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$, which describe the expected values v and associated resources P that should be transferred over a bidirectional channel, after which it continues as the tail protocol $prot$. These streams are either infinite or finite, where the finite streams are terminated by an **end** constructor. What makes dependent separation protocols different from ordinary streams (that are typically defined as a coinductive data type), is the sequence of logical variables $\vec{x}:\vec{\tau}$ that bind into the expected value v , the resources P , and the tail protocol $prot$. The variables $\vec{x}:\vec{\tau}$ are higher-order and impredicative, meaning that the types $\vec{\tau}$ can be any type of Iris, including functions, propositions, predicates, and protocols themselves.

In order to give a model of dependent separation protocols, we first need to consider what $!\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$ mean semantically. As the variables $\vec{x}:\vec{\tau}$ bind into v , P , and $prot$, these protocols can be thought of as predicates over the transferred physical value and actual tail of the protocol. Since dependent separation protocols describe the logical resources that are being transferred along with the message, they should be modelled using Iris predicates (**iProp**) instead of meta-level predicates (**Prop**). This leads to the following formal intermediate definition:

$$\begin{aligned} \widehat{\text{iProto}} &\triangleq 1 + (\text{action} \times (\text{Val} \rightarrow \blacktriangleright \widehat{\text{iProto}} \rightarrow \text{iProp})) \quad (\text{where } \text{action} ::= \text{send} \mid \text{recv}) \\ \text{end} &\triangleq \text{inj}_1 () \\ !\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot &\triangleq \text{inj}_2 (\text{send}, \lambda w \text{ prot}'. \exists \vec{x}:\vec{\tau}. (v = w) * (\blacktriangleright P) * (\text{prot}' = \text{next } prot)) \\ ?\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot &\triangleq \text{inj}_2 (\text{recv}, \lambda w \text{ prot}'. \exists \vec{x}:\vec{\tau}. (v = w) * (\blacktriangleright P) * (\text{prot}' = \text{next } prot)) \end{aligned}$$

We define $\widehat{\text{iProto}}$ using Iris's support for solving *guarded recursive domain equations* [[America and Rutten 1989](#); [Birkedal et al. 2012](#)]. This means that the recursive occurrence of $\widehat{\text{iProto}}$ must appear below a *later* ($\blacktriangleright$) operator (whose only constructor is $\text{next} : T \rightarrow \blacktriangleright T$). The left part of the sum type indicates that the protocol has terminated, while the right part is a predicate that describes the exchange. The existential quantifiers in the definitions make sure that the variables $\vec{x}:\vec{\tau}$ bind into v , P , and $prot$. Moreover, to make it possible to construct recursive protocols, $!\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$ and $?\vec{x}:\vec{\tau}\langle v \rangle\{P\}.prot$ are contractive in P and $prot$ due to the uses of $\blacktriangleright P$ and $\text{next } prot$.

Since Iris's ghost theory will be instantiated with $\widehat{\text{iProto}}$ in Section 5.3, $\widehat{\text{iProto}}$ and iProp need to be defined in a mutually recursive fashion. Unfortunately, as a consequence of this, along with the

negative occurrence of iProto , the above definition of $\widehat{\text{iProto}}$ cannot be constructed in the current Coq development of Iris without significant manual effort. Although there is no fundamental restriction preventing this definition [Birkedal et al. 2012, § 7], the necessary changes to the Coq development needed to construct such types are orthogonal to the purpose of this paper. As such, the actual—albeit less intuitive—definition of iProto is as follows:

$$\begin{aligned} \text{iProto} &= 1 + (\text{action} \times (\text{Val} \rightarrow (\blacktriangleright \text{iProto} \rightarrow \text{iProp}) \rightarrow \text{iProp})) \\ \text{end} &\triangleq \text{inj}_1 () \\ !\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot} &\triangleq \text{inj}_2 (\text{send}, \lambda w (f : \blacktriangleright \text{iProto} \rightarrow \text{iProp}). \exists \vec{x} : \vec{\tau}. (v = w) * \blacktriangleright P * f(\text{next prot})) \\ ?\vec{x} : \vec{\tau} \langle v \rangle \{P\}. \text{prot} &\triangleq \text{inj}_2 (\text{recv}, \lambda w (f : \blacktriangleright \text{iProto} \rightarrow \text{iProp}). \exists \vec{x} : \vec{\tau}. (v = w) * \blacktriangleright P * f(\text{next prot})) \end{aligned}$$

In this definition, we use a predicate $f : \blacktriangleright \text{iProto} \rightarrow \text{iProp}$ to restrict the tail protocol, causing iProto to appear in a positive position. When using these protocols in Section 5.3, we always pick the predicate f to be $(\lambda \text{prot}'' . \text{prot}' = \text{prot}'')$ where prot' is the actual tail, which would otherwise have been supplied as the tail argument to the ideal definition. Hence, the definition iProto is functionally equivalent to the ideal definition $\widehat{\text{iProto}}$.

With the actual definition at hand, the dual $(\overline{_})$ and composition $(_ \cdot _)$ operations are defined using Iris's guarded recursion operator $(\mu x : \tau. t)$ as follows:

$$\begin{aligned} \overline{(_)} &\triangleq \mu \text{rec}. \lambda \text{prot}. \begin{cases} \text{inj}_1 () & \text{if } \text{prot} = \text{inj}_1 () \\ \text{inj}_2 (\overline{a}, \lambda v f. \Phi v (f \circ \text{map rec})) & \text{if } \text{prot} = \text{inj}_2 (a, \Phi) \end{cases} \\ (_ \cdot \text{prot}_2) &\triangleq \mu \text{rec}. \lambda \text{prot}_1. \begin{cases} \text{prot}_2 & \text{if } \text{prot}_1 = \text{inj}_1 () \\ \text{inj}_2 (a, \lambda v f. \Phi v (f \circ \text{map rec})) & \text{if } \text{prot}_1 = \text{inj}_2 (a, \Phi) \end{cases} \end{aligned}$$

Here, we let $\overline{\text{send}} \triangleq \text{recv}$, $\overline{\text{recv}} \triangleq \text{send}$, and $\text{map} : (T \rightarrow U) \rightarrow (\blacktriangleright T \rightarrow \blacktriangleright U)$, which is the mapping function on the $\blacktriangleright$ -type.

5.3 The Model of Channel Ownership

Now that we have given a semantics of channels (Section 5.1) and a model of dependent separation protocols (Section 5.2), it remains to use these concepts to define the $c \multimap \text{prot}$ connective.

In order to define Actris's connectives $c_1 \multimap \text{prot}_1$ and $c_2 \multimap \text{prot}_2$ for ownership of the endpoints c_1 and c_2 of a bidirectional channel, we need to tie prot_1 and prot_2 to the physical contents of the channel buffers as captured by the representation predicate $(c_1, c_2) \multimap (\vec{v}_1, \vec{v}_2)$. To obtain duality of the endpoint protocols, they must always be in one of two configurations:

- There are zero or more messages $\vec{v}_1$ in transit from the first endpoint to the second. That is, the buffer contents is $(c_1, c_2) \multimap (\vec{v}_1, \epsilon)$, and prot_2 starts with a series of receive (?) nodes that match up with the contents $\vec{v}_1$ of the first buffer.
- There are zero or more messages $\vec{v}_2$ in transit from the second endpoint to the first. That is, the buffer contents is $(c_1, c_2) \multimap (\epsilon, \vec{v}_2)$, and prot_1 starts with a series of receive (?) nodes that match up with the contents $\vec{v}_2$ of the second buffer.

Note that dependent separation protocols (and binary session types) enforce that either of the two channel buffers is always empty as a channel endpoint can only start sending messages (!) once it has received (?) all messages sent by the other endpoint. This property is not generally true for bidirectional channels.

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

Fig. 12. Higher-order ghost variables in Iris.

This informal invariant can be formalised in Iris using the definition I , which, in turn, is used to define the connective $c \rightsquigarrow \text{prot}$ for channel ownership:

$$\begin{aligned}
& \text{interp} \in \text{prot}_1 \text{ prot}_2 \triangleq (\text{prot}_1 = \overline{\text{prot}_2}) \\
& \text{interp} ([v] \cdot \vec{v}) \text{ prot}_1 \text{ prot}_2 \triangleq \exists \Phi \text{ prot}_2'. (\text{prot}_2 = \text{inj}_2(\text{recv}, \Phi)) * \\
& \quad \Phi v (\lambda \text{prot}'''. \text{next } \text{prot}_2' = \text{prot}''') * \\
& \quad \triangleright \text{interp } \vec{v} \text{ prot}_1 \text{ prot}_2' \\
& I \gamma_1 \gamma_2 c_1 c_2 \triangleq \exists \vec{v}_1 \vec{v}_2 \text{ prot}_1 \text{ prot}_2. (c_1, c_2) \rightsquigarrow (\vec{v}_1, \vec{v}_2) * \\
& \quad \gamma_1 \mapsto_{\bullet} \text{prot}_1 * \gamma_2 \mapsto_{\bullet} \text{prot}_2 * \\
& \quad \triangleright \left(\begin{aligned} & (\vec{v}_2 = \epsilon * \text{interp } \vec{v}_1 \text{ prot}_1 \text{ prot}_2) \vee \\ & (\vec{v}_1 = \epsilon * \text{interp } \vec{v}_2 \text{ prot}_2 \text{ prot}_1) \end{aligned} \right) \\
& c \rightsquigarrow \text{prot} \triangleq \exists \gamma_1 \gamma_2 c_1 c_2. \boxed{I \gamma_1 \gamma_2 c_1 c_2} * \\
& \quad ((c = c_1 * \gamma_1 \mapsto_{\circ} \text{prot}) \vee (c = c_2 * \gamma_2 \mapsto_{\circ} \text{prot}))
\end{aligned}$$

Setting the deeper technicalities related to the encoding in Iris aside, the most important parts are as follows. The two configurations informally described in the bullet list above are captured by the disjunction in the invariant I . The predicate $\text{interp } \vec{v} \text{ prot}_1 \text{ prot}_2$ captures that prot_2 contains zero or more receive (?) nodes that match up with the values $\vec{v}$ currently in the channel buffer. The clause $\Phi v (\lambda \text{prot}'''. \text{next } \text{prot}_2' = \text{prot}''')$ in the definition of interp describes the resources held by the messages in the channel buffer, while unifying the physical value v and the protocol tail prot_2' that is then used in the recursive call to interp .

In order to tie the proposition $c \rightsquigarrow \text{prot}$ to the contents of the bidirectional channel c we follow the usual Iris methodology: we make use of Iris's invariants connective $\boxed{R}$ and a suitable “ghost theory”. The invariant connective $\boxed{R}$ expresses that the proposition R holds at any time, *i.e.*, is invariant. As for the ghost theory we use higher-order ghost variables, whose rules are shown in Figure 12. Higher-order ghost variables come in pairs $\gamma \mapsto_{\bullet} \text{prot}$ and $\gamma \mapsto_{\circ} \text{prot}$, which should always hold the same protocol prot . They can be allocated together (HO-GHOST-ALLOC), are always required to hold the same protocol (HO-GHOST-AGREE), and can only be updated together (HO-GHOST-UPDATE). The subtle part of higher-order ghost variables is that they involve ownership of protocols of type iProto , which are defined in terms of Iris propositions iProp , which themselves are defined in terms of iProto . Due to this mutual recursion, the rule HO-GHOST-AGREE only gives the equality between dependent separation protocols under a later modality.

5.4 Adequacy of Dependent Separation Protocols

With the above definitions at hand, we can now prove the rules HT-NEWCHAN, HT-SEND and HT-RCV from the combination of the logically atomic specifications in Section 5.1 and Iris's rules for invariants and ghost state. The proof of rule HT-NEWCHAN is straightforward from HO-GHOST-ALLOC, which is used to allocate two pairs of higher-order ghost variables for the chosen protocol and its dual. The crux of the proofs of HT-SEND and HT-RCV is to update the higher-order ghost variables $\gamma \mapsto_{\bullet} \text{prot}$ and $\gamma \mapsto_{\circ} \text{prot}$ using HO-GHOST-UPDATE to the tail of the protocol, and to transfer the

resources described by the head of the protocol in and out of the invariant, when sending and receiving, respectively. This gives rise to the main result:

THEOREM 5.1 (ADEQUACY OF ACTRIS). *Let ϕ be a first-order predicate over values and suppose the Hoare triple $\{\text{True}\} e \{\phi\}$ is derivable in Actris, then:*

- **(Safety):** *The program e will not get stuck.*
- **(Postcondition validity):** *If the main thread of e terminates with a value v , then the postcondition ϕv holds at the meta-level.*

Since Actris is an internal logic embedded in Iris, the proof is an immediate consequence of Iris's adequacy theorem [Jung et al. 2018b; Krebbers et al. 2017a]. Finally, note that safety implies *session fidelity*—any message that is received has in fact been sent.

6 COQ MECHANISATION

The definition of the Actris logic, its model, and the proofs of the examples in this paper have been fully mechanised using the Coq proof assistant [Coq Development Team 2019]. The mechanisation is built on top of the mechanisation of Iris [Jung et al. 2016, 2018b; Krebbers et al. 2017a] and the MoSeL Proof Mode (formerly Iris Proof Mode) [Krebbers et al. 2018, 2017b], which essentially provides an embedded proof assistant for separation logic in Coq. Building Actris on top of the Iris framework in Coq has a number of tangible advantages:

- By defining channels on top of HeapLang, the default concurrent language shipped with Iris, we do not have to define a full programming language semantics and can reuse all of the Coq machinery, including the tactics for symbolic execution of non message-passing programs.
- Since Actris is essentially mechanised as an Iris library that provides support for the iProto type, the $c \multimap \text{prot}$ connective, the various operations on protocols, and the proof rules as lemmas, we get all present features of Iris for free.
- 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 give rise to a very small Coq development. The total size is about 3500 lines of code (comments and whitespace are included), of which 200 lines are used for the definitional semantics of channels, 1000 lines for the model, 450 lines for tactics, 1250 lines for the examples, and 600 lines for utilities (linked lists, permutations, *etc.*).

In order to implement Actris's tactics for symbolic execution, we followed the methodology described in the original Iris Proof Mode paper [Krebbers et al. 2017b], which means that the logic in Coq is presented in weakest precondition style rather than using Hoare triples. For handling `send` or `recv` we defined 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.
- Find a corresponding $c \multimap \text{prot}$ hypothesis in the separation logic context.
- Normalise the protocol prot using the rules for duals, composition, and fixpoints so it can be written with a $! \vec{x} : \vec{\tau} \langle v \rangle \{P\}$, prot or $? \vec{x} : \vec{\tau} \langle v \rangle \{P\}$, prot construct in head position.
- In case of `wp_send`, instantiate the variables $\vec{x} : \vec{\tau}$ using the terms $(t1 \dots tn)$, and create a goal for the proposition P with the hypotheses $[H1 \dots Hn]$. Hypotheses prefixed with \$ are framed. In case the terms $(t1 \dots tn)$ are omitted, these are determined using unification.

- In case of `wp_recv`, introduce the variables $\vec{x}:\vec{\tau}$ into the context by naming them ($y_1 \dots y_n$), and create a hypothesis H for P .

7 RELATED WORK

As Actris combines results from both the separation logic and session types community, there is an abundance of related work. This section briefly elaborates on the relation to message passing in separation logic (Section 7.1) and process calculi (Section 7.2), session types (Section 7.3), endpoint sharing (Section 7.4), and verification efforts of map-reduce (Section 7.5).

7.1 Message Passing and Separation Logic

Lozes and Villard [2012] proposed a logic, based on previous work by Villard et al. [2009], to reason about programs written in a small imperative language with message passing using channels similar to ours. Messages are labelled, and protocols are handled with a combination of finite-state automata (FSA) with correspondingly labelled transitions and predicates associated with each state of the automata. This combination is similar to, but less general than, STSs in Iris. Their language does not support higher-order functions or delegation, but since their language is restricted to structured concurrency (*i.e.*, not fork-based) and their logic is linear (*i.e.*, not affine), they ensure that all resources like channels and memory are properly deallocated.

Craciun et al. [2015] 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 via message-passing, delegation through higher-order channels, branching 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, however, is not higher-order, which means that sending functions over channels is not possible. Moreover, their logic does not support protocol-level binders that can connect the transferred message with the tail protocol. It is therefore not possible to model dependent protocols like we do in Actris. 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 [Jung et al. 2015] 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 only reason about the entire channel buffer, protocol reasoning must be done via STSs or other forms of ghost state.

Hamin and Jacobs [2019] 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 did not provide abstractions akin to our session-type based protocols. Instead one has to reason using invariants and ghost state explicitly.

Mansky et al. [2017] take yet another direction and verify the functional correctness of a message-passing system written in C using VST [Appel 2014]. 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.

The work most closely related to ours, and the only work we know of that combines session types with separation logic, is by Tassarotti et al. [2017] whose main contribution is to prove correctness and termination preservation of a compiler from a simple language with session types to a functional language with mutable state, where the channels are implemented using references

on the heap. This work is also done in Iris. 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. \[2018\]](#) and the Aneris logic by [Krogh-Jespersen et al. \[2019\]](#) 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 address a different use case than we do.

7.2 Separation Logic and Process Calculi

Another approach is to verify message-passing programs written in some dialect of process calculus. We focus on related work that combines process calculus with separation logic.

[Francalanza et al. \[2011\]](#) 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. \[2016\]](#) 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 source code, but there is no formal connection between the specification logic and the underlying language.

Neither approach supports delegation or any concurrency paradigms other than message passing.

7.3 Session Types

Seminal work on linear type systems for the pi calculus by [Kobayashi et al. \[1996\]](#) led to the creation of binary session types by [Honda et al. \[1998\]](#).

[Bocchi et al. \[2010\]](#) pushed the boundaries of what can be verified with multi-party 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.

Later work by [Dardha et al. \[2012\]](#) 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 [[Lange et al. 2018](#)], OCaml [[Imai et al. 2019](#); [Padovani 2017](#)], and Java [[Hu et al. 2010](#)]. These works focus on adding session-typed support for the Actor model in existing languages, but do not target proving functional correctness of programs.

7.4 Endpoint sharing

One of the key features of session types is that endpoints are owned by a single process. While these endpoints can be delegated (*i.e.*, transferred from one process to another), they typically cannot be shared (*i.e.*, be accessed by multiple processes concurrently). However, as we demonstrate in Section 3, sharing channels endpoints is often desirable.

In the pi calculus community there has been prior work on endpoint sharing, *e.g.*, by [Atkey et al. \[2016\]](#); [Kobayashi \[2006\]](#); [Padovani \[2014\]](#). The latest contribution in this line of work is by [Balzer and Pfenning \[2017\]](#); [Balzer et al. \[2019\]](#), 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.

7.5 Verification of Map-Reduce

To our knowledge the only verification related to the map-reduce model [Dean and Ghemawat 2004] is by Ono et al. [2011], 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 [Marché et al. 2004], 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.

8 FUTURE WORK

One of the most prominent extensions of binary session types is multi-party session types [Honda et al. 2008], often called choreographies, which allow concise specifications of message transfers between more than two parties. It would be interesting to explore a multi-party version of dependent separation protocols, similar to the multi-party version of session logic by Costea et al. [2018], to allow Actris to verify programs that make use of multi-party communication.

In addition to safety (*i.e.*, session fidelity), conventional session type systems guarantee properties like deadlock and resource-leak freedom. Since Actris is an extension of concurrent separation logic that supports reasoning about several concurrency primitives including message passing, ensuring deadlock freedom is hard. The only prior work in this direction that we are aware of is by Hamin and Jacobs [2019] and Craciun et al. [2015], but it is not immediately obvious how to integrate that with Iris or Actris. Leak freedom has been studied in Iron, an extension of Iris by Bizjak et al. [2019], which makes it possible to prove resource-leak freedom of non-structured fork-based concurrent programs. It would be interesting to build dependent separation protocols on top of Iron instead of Iris.

Another direction for future work is to use dependent separation protocols for a logical relations model of session types. Similar to the RustBelt project [Jung et al. 2018a], this would give rise to an extensible approach for proving type soundness of type systems with session types, which can be used to establish that unsafe code in libraries has been safely encapsulated by its ADT.

ACKNOWLEDGMENTS

We thank the anonymous reviewers for their helpful feedback. We are grateful to Fabrizio Montesi, Daniel Gratzer, Andreea Costea, and the participants of the Iris workshop 2019 for discussions. Robbert Krebbers was supported by the Netherlands Organisation for Scientific Research (NWO), project number 016.Veni.192.259.

REFERENCES

- Pierre America and Jan J. M. M. Rutten. 1989. Solving Reflexive Domain Equations in a Category of Complete Metric Spaces. *JCSS* 39, 3 (1989), 343–375.
- Andrew W Appel. 2014. *Program logics for certified compilers*. Cambridge University Press.
- Robert Atkey, Sam Lindley, and J. Garrett Morris. 2016. Conflation Confers Concurrency. In *Essays Dedicated to Philip Wadler on the Occasion of His 60th Birthday (LNCS)*, Vol. 9600. 32–55.
- Stephanie Balzer and Frank Pfenning. 2017. Manifest Sharing with Session Types. *PACMPL* 1, ICFP (2017), 37:1–37:29.

- Stephanie Balzer, Bernardo Toninho, and Frank Pfenning. 2019. Manifest Deadlock-Freedom for Shared Session Types. In *ESOP*, Vol. 11423 LNCS. 611–639.
- Lars Birkedal, Rasmus Ejlers Møgelberg, Jan Schwinghammer, and Kristian Støvring. 2012. First Steps in Synthetic Guarded Domain Theory: Step-Indexing in the Topos of Trees. *LMCS* 8, 4 (2012).
- Ales Bizjak, Daniel Gratzer, Robbert Krebbers, and Lars Birkedal. 2019. Iron: Managing Obligations in Higher-Order Concurrent Separation Logic. *PACMPL* 3, POPL (2019), 65:1–65:30.
- Laura Bocchi, Kohei Honda, Emilio Tuosto, and Nobuko Yoshida. 2010. A Theory of Design-by-Contract for Distributed Multiparty Interactions. In *CONCUR*. 162–176.
- Coq Development Team. 2019. The Coq Proof Assistant Reference Manual, Version 8.9. (2019). <https://coq.inria.fr/distrib/current/refman/>
- Andreea Costea, Wei-Ngan Chin, Shengchao Qin, and Florin Craciun. 2018. Automated Modular Verification for Relaxed Communication Protocols. In *APLAS (LNCS)*, Vol. 11275. 284–305.
- Florin Craciun, Tibor Kiss, and Andreea Costea. 2015. Towards a Session Logic for Communication Protocols. In *ICECCS*. 140–149.
- Pedro da Rocha Pinto, Thomas Dinsdale-Young, and Philippa Gardner. 2014. TaDA: A Logic for Time and Data Abstraction. In *ECOOP*. 207–231.
- Ornela Dardha, Elena Giachino, and Davide Sangiorgi. 2012. Session Types Revisited. In *PPDP*. 139–150.
- Jeffrey Dean and Sanjay Ghemawat. 2004. MapReduce: Simplified Data Processing on Large Clusters. In *OSDI*. 137–150.
- Adrian Francalanza, Julian Rathke, and Vladimiro Sassone. 2011. Permission-Based Separation Logic for Message-Passing Concurrency. *LMCS* 7, 3 (2011).
- Jafar Hamin and Bart Jacobs. 2019. Transferring Obligations Through Synchronizations. In *ECOOP*.
- Carl Hewitt, Peter Bishop, and Richard Steiger. 1973. A Universal Modular ACTOR Formalism for Artificial Intelligence. In *IJCAI*. 235–245.
- Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. 2019. Coq Mechanization of Actris. Available online at <https://gitlab.mpi-sws.org/iris/actris>.
- Kohei Honda, Vasco Thudichum Vasconcelos, and Makoto Kubo. 1998. Language Primitives and Type Discipline for Structured Communication-Based Programming. In *proceedings of ESOP*. 122–138.
- Kohei Honda, Nobuko Yoshida, and Marco Carbone. 2008. Multiparty Asynchronous Session Types. In *POPL*. 273–284.
- Raymond Hu, Dimitrios Kouzapas, Olivier Pernet, Nobuko Yoshida, and Kohei Honda. 2010. Type-Safe Eventful Sessions in Java. In *ECOOP*. 21–25.
- Keigo Imai, Nobuko Yoshida, and Shoji Yuen. 2019. Session-OCaml: A Session-Based Library with Polarities and Lenses. *Science of Computer Programming* 172 (2019), 135–159.
- Bart Jacobs and Frank Piessens. 2011. Expressive Modular Fine-Grained Concurrency Specification. In *POPL*. 271–282.
- Ralf Jung, Jacques-Henri Jourdan, Robbert Krebbers, and Derek Dreyer. 2018a. RustBelt: Securing the Foundations of the Rust Programming Language. *PACMPL* 2, POPL (2018), 66:1–66:34.
- Ralf Jung, Robbert Krebbers, Lars Birkedal, and Derek Dreyer. 2016. Higher-Order Ghost State. In *ICFP*. 256–269.
- Ralf Jung, Robbert Krebbers, Jacques-Henri Jourdan, Ales Bizjak, Lars Birkedal, and Derek Dreyer. 2018b. Iris From the Ground Up: A Modular Foundation for Higher-Order Concurrent Separation Logic. *JFP* 28 (2018), e20.
- Ralf Jung, David Swasey, Filip Sieczkowski, Kasper Svendsen, Aaron Turon, Lars Birkedal, and Derek Dreyer. 2015. Iris: Monoids and Invariants as an Orthogonal Basis for Concurrent Reasoning. In *POPL*. 637–650.
- Naoki Kobayashi. 2006. A New Type System for Deadlock-Free Processes. In *CONCUR (LNCS)*, Vol. 4137. 233–247.
- Naoki Kobayashi, Benjamin C. Pierce, and David N. Turner. 1996. Linearity and the pi-Calculus. In *POPL*. 358–371.
- Robbert Krebbers, Jacques-Henri Jourdan, Ralf Jung, Joseph Tassarotti, Jan-Oliver Kaiser, Amin Timany, Arthur Charguéraud, and Derek Dreyer. 2018. MoSeL: A General, Extensible Modal Framework for Interactive Proofs in Separation Logic. *PACMPL* 2, ICFP (2018), 77:1–77:30.
- Robbert Krebbers, Ralf Jung, Ales Bizjak, Jacques-Henri Jourdan, Derek Dreyer, and Lars Birkedal. 2017a. The Essence of Higher-Order Concurrent Separation Logic. In *ESOP*. 696–723.
- Robbert Krebbers, Amin Timany, and Lars Birkedal. 2017b. Interactive Proofs in Higher-Order Concurrent Separation Logic. In *POPL*. 205–217.
- Morten Krogh-Jespersen, Amin Timany, Marit Edna Ohlenbusch, Simon Oddershede Gregersen, and Lars Birkedal. 2019. Aneris: A Mechanised Logic for Modular Reasoning about Distributed Systems. Submitted for publication.
- Julien Lange, Nicholas Ng, Bernardo Toninho, and Nobuko Yoshida. 2018. A Static Verification Framework for Message Passing in Go Using Behavioural Types. *ICSE* (2018), 1137–1148.
- Étienne Lozes and Jules Villard. 2012. Shared Contract-Obedient Endpoints. In *ICE*. 17–31.
- William Mansky, Andrew W. Appel, and Aleksey Nogin. 2017. A Verified Messaging System. *PACMPL* 1, OOPSLA (2017), 87:1–87:28.

- Claude Marché, Christine Paulin-Mohring, and Xavier Urbain. 2004. The KRAKATOA Tool for Certification of JAVA/JAVAC-ARD Programs Annotated in JML. *JLP* 58, 1-2 (2004), 89–106.
- Dimitris Mostrous and Vasco Thudichum Vasconcelos. 2014. Affine Sessions. In *COORDINATION*. 115–130.
- Hiroshi Nakano. 2000. A Modality for Recursion. In *LICS*. 255–266.
- Aleksandar Nanevski, Ruy Ley-Wild, Ilya Sergey, and Germán Andrés Delbianco. 2014. Communicating State Transition Systems for Fine-Grained Concurrent Resources. In *ESOP*. 290–310.
- Kosuke Ono, Yoichi Hirai, Yoshinori Tanabe, Natsuko Noda, and Masami Hagiya. 2011. Using Coq in Specification and Program Extraction of Hadoop MapReduce Applications. In *SEFM*. 350–365.
- Wytse Oortwijn, Stefan Blom, and Marieke Huisman. 2016. Future-based Static Analysis of Message Passing Programs. In *PLACES*. 65–72.
- Luca Padovani. 2014. Deadlock and Lock Freedom in the Linear π -Calculus. In *CSL*. 72:1–72:10.
- Luca Padovani. 2017. A Simple Library Implementation of Binary Sessions. *JFP* 27, 2010 (2017), e4.
- Ilya Sergey, James R. Wilcox, and Zachary Tatlock. 2018. Programming and Proving with Distributed Protocols. *PACMPL* 2, POPL (2018), 28:1–28:30.
- Kasper Svendsen and Lars Birkedal. 2014. Impredicative Concurrent Abstract Predicates. In *ESOP*. 149–168.
- Samira Tasharofi, Peter Dinges, and Ralph E. Johnson. 2013. Why Do Scala Developers Mix the Actor Model with Other Concurrency Models?. In *ECOOP*. 302–326.
- Joseph Tassarotti, Ralf Jung, and Robert Harper. 2017. A Higher-Order Logic for Concurrent Termination-Preserving Refinement. In *ESOP*. 909–936.
- Tengfei Tu, Xiaoyu Liu, Linhai Song, and Yiyang Zhang. 2019. Understanding Real-World Concurrency Bugs in Go. In *ASPLOS*. 865–878.
- Jules Villard, Étienne Lozes, and Cristiano Calcagno. 2009. Proving Copyless Message Passing. In *APLAS*. 194–209.