

An Operational Approach to Library Abstraction under Relaxed Memory Concurrency (Extended Version)

ABHISHEK KR SINGH, Tel Aviv University, Israel

ORI LAHAV, Tel Aviv University, Israel

Concurrent data structures and synchronization mechanisms implemented by expert developers are indispensable for modular software development. In this paper, we address the fundamental problem of library abstraction under weak memory concurrency, and identify a general library correctness condition allowing clients of the library to reason about program behaviors using the specification code, which is often much simpler than the concrete implementation. We target (a fragment of) the RC11 memory model, and develop an equivalent operational presentation that exposes knowledge propagation between threads, and is sufficiently expressive to capture library behaviors as totally ordered operational execution traces. We further introduce novel access modes to the language that allow intricate specifications accounting for library internal synchronization that is not exposed to the client, as well as the library's demands on external synchronization by the client. We illustrate applications of our approach in several examples of different natures.

CCS Concepts: • **Theory of computation** → **Concurrency**; **Operational semantics**; **Program verification**; • **Software and its engineering** → **Semantics**; **Software verification**; **Abstraction, modeling and modularity**.

Additional Key Words and Phrases: Relaxed memory consistency, Concurrent objects, Linearizability, Library abstraction

ACM Reference Format:

Abhishek Kr Singh and Ori Lahav. 2023. An Operational Approach to Library Abstraction under Relaxed Memory Concurrency (Extended Version). *Proc. ACM Program. Lang.* 7, POPL, Article 53 (January 2023), 45 pages. <https://doi.org/10.1145/3571246>

1 Introduction

Library abstraction constitutes a powerful means to achieve modularity in software development. It allows expert library developers to write optimized implementations of common programming tasks, and “once and for all” establish that these implementations admit their corresponding specifications. In turn, users of these implementations, called *clients* of the library, may reason about program behaviors assuming only the libraries’ specifications, with no understanding, or even no access, to the implementations. From a formal standpoint, a prerequisite for applying library abstraction is to identify a condition that provably allows library developers to abstract away from a particular client program when verifying their implementations, while ensuring the soundness of client reasoning using the specifications in any (valid) client program.

In the case of sequential programs, library abstraction is straightforward: specifications may be given using pre- and post-conditions, and it is easy to establish the soundness of client reasoning that is based on them. Concurrent programs, however, require more attention. Assuming

Authors’ Contact Information: [Abhishek Kr Singh](#), Tel Aviv University, Israel, abhishek.uor@gmail.com; [Ori Lahav](#), Tel Aviv University, Israel, orilahav@tau.ac.il.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

© 2023 Copyright held by the owner/author(s).

ACM 2475-1421/2023/1-ART53

<https://doi.org/10.1145/3571246>

an underlying sequentially consistent (SC) memory system [Lamport 1979], classical linearizability (w.r.t. some sequential specification) [Herlihy and Wing 1990] ensures refinement w.r.t. a sequential object [Filipović et al. 2010], and can be seen as a library abstraction condition. More recent approaches employ “code as specification” (see, e.g., [Gotsman and Yang 2011]), where implementations are specified using (simple) code, and library abstraction amounts to contextual refinement between two pieces of code. Linearizability can be seen as a particular instance of this approach, where specifications are obtained by wrapping a sequential implementation within a global lock [Bouajjani et al. 2015]. Furthermore, to make such specification formalism more useful, one often enriches the given programming language with special specification constructs, that are employed only in the specification code and are relatively easy to reason about (e.g., atomic blocks in concurrent programs).

The current paper studies library abstraction under relaxed (a.k.a. weak) shared-memory concurrency. In particular, we consider a C11-style weak memory model with several kinds of memory-access modes (a.k.a. memory orderings), placed differently on the spectrum between efficient implementation (with an optimizing compiler targeting a modern multicore hardware) and strong consistency guarantees. The infamous complexity of programming under such a model makes library abstraction indispensable. In particular, library abstraction allows clients to use the programming guarantees supplied by the model, e.g., clients should be able to rely on the synchronization induced by a library (such as a lock library) for ensuring data-race freedom in their programs, and then apply the model’s guarantee that data-race free programs have strong semantics.

Remark 1. Like many formal verification frameworks for C11, we are unable to work with the original model, which allows unrestricted cycles in the union of ‘program order’ and ‘reads-from’, and thus exhibits “out-of-thin-air” behaviors and fails to provide the most basic data-race-freedom guarantee. We follow [Boehm and Demsky 2014] and its formalization in the RC11 model in [Lahav et al. 2017] to conservatively disallow *all* such cycles. This entails a certain performance penalty for maintaining the load-store order between relaxed accesses [Ou and Demsky 2018]. We also note that the fragment of the RC11 handled in this paper lacks release/acquire fences and sequentially consistent accesses, which are left to future work.

Our main contribution is a correctness condition for libraries that provably ensures (contextual) refinement between implementations and their respective specifications under the weak memory model, with several distinctive properties:

- Our proposed condition is based on *totally ordered execution traces* of the library in question (a decision motivated below). Accordingly, we present a novel equivalent operational version of the (originally) declarative memory model, and use it in the correctness condition.
- To account for various restrictions that libraries may impose on their clients (e.g., refrain from data races between two specific methods), our result supports rich library calling policies, and refinement is conditioned on the client’s adherence to the calling policy. Importantly, whether a program adheres to the policy or not is checked against the library specification, which allows the application of the abstraction theorem by the client without any knowledge of the implementation.
- Our specification language allows simple *lock-based* specifications of libraries that do not provide synchronization guarantees for their clients, and dually of libraries that rely on the client for performing the synchronization between library calls. For that matter, we introduce a novel access mode to the programming language (and memory model).

We illustrate the application of our approach for an RCU synchronization mechanism which we specify by relying solely on locks (§8.1), as well as for a relaxed concurrent queue object that does not expose its internal synchronization to its clients (§8.2). We also derive a (local) data-race-freedom guarantee for the memory model as an instance of library abstraction (§8.3).

Outline. This paper is organized as follows. In §2 we present an informal overview of the challenges we address and our solutions. In §3 we define concurrent programs and their semantics independent of a memory model. In §4 we present the memory model semantics, which we call dRC11 (d for ‘declarative’), as a fragment of RC11 extended with novel constraints for the new access modes. In §5 we present an operational version of the memory model, which we call pRC11 (p for ‘propagation’), that is needed for defining the library correctness condition. In §6 we introduce the notions required to formulate the library abstraction theorem. In §7 we state and prove the main theorem. In §8 we demonstrate applications of the abstraction theorem (which can assist in understanding the crux of the abstraction theorem even before reading the more technical material in Section 3 to 7). We discuss the relation to other work in §9 and conclude in §10. Additional detailed proofs are given in §A.

2 Key Challenges and Ideas

We outline the main challenges and the key ideas in our solutions. We keep the discussion and examples informal, leaving the formal development to later sections.

2.1 Library Correctness Criterion

The main challenge lies in establishing a library correctness condition that will ensure contextual refinement between a library specification $L^\#$ (given as code) and its implementation L under weak memory semantics. Naturally, such condition should consider each object in isolation and avoid quantification over all possible clients. Moreover, we opt for an *operational* condition that, like standard linearizability, is based on totally ordered histories generated by the object in question. The latter desideratum is in contrast with previous work on library specifications in (R)C11 [Batty et al. 2013; Raad et al. 2019] and adaptations of linearizability to weak memory concurrency (see, e.g., [Dongol et al. 2018]), which employ multiple partial orders in their correctness criteria. We believe that an operational approach, based on *one* total timeline, may be more intuitive for users (also considering informal arguments), and will allow easier adoption of standard techniques and tools that were applied before to verify refinement between transition systems (see e.g., our initial experiment with the FDR refinement checker in §8.1).

While having an operational correctness criterion may seem contradicting to the fact that standard weak memory formulations, e.g., (R)C11, are *declarative* (a.k.a. axiomatic memory models), unlike previous work (e.g., [Batty et al. 2013]), we consider this to be only a superficial matter. Indeed, declarative specifications of models that forbid “reads from the future” (i.e., impose acyclicity of the union of the ‘program order’ and the ‘read-from’ relations; called an ‘in-order’ semantics in [Cho et al. 2021]) can be equally characterized as operational memory models: the states consist of the current execution trace ordered by multiple partial order relations as needed, and transitions are only between *consistent* states as specified by the consistency condition of the declarative model (see, e.g., the RAG transition system for the release/acquire fragment of C11 in [Lahav and Margalit 2019]). More concise (and possibly more intuitive) formulations of such semantics may use timestamps and views, and the equivalence to the “operationalized” declarative model is witnessed by a standard forward simulation relation (see, e.g., [Kaiser et al. 2017; Kang et al. 2017]).

Now, operational semantics (in particular, operationalized versions of declarative models), naturally give rise to a set of *library histories*, and allow the application of standard linearizability. For that matter, one considers a “most general client” program that invokes the library methods in the most general way expected by the library. Then, the set of library histories consists of all sequences of invocations (with the values of the arguments) and responses (with the returned values) of the library methods by different threads that are obtained in operational traces of that client. Under SC,

inclusion of the sets of histories between two libraries may serve as a library abstraction condition (i.e., if every history of a library L is also a history of $L^\#$, then L refines $L^\#$).

Aiming to apply a similar condition, following previous work on abstraction under TSO [Burckhardt et al. 2012] and under non-volatile memory [Khyzha and Lahav 2022], we observe that under weak memory semantics, inclusion of sets of histories is generally *unsound* as a library correctness condition.¹ Next, we provide a simple (contrived) example of this issue.

Example 2.1. Consider a library L with two methods `foo` and `bar` and the following specification (here and henceforth we assume that all variables are initialized to 0):

```
foo() :  store(x, 1, rel); return();
```

```
bar() :  a := load(x, acq); return(a);
```

Suppose that the library enforces a policy on its clients: (1) `foo` and `bar` must be called *exactly once* in *different* threads; and (2) `bar` must be called *after* `foo` in the execution order. For example, the following program adheres to this call policy:

```
foo();
store(y, 1, rlz); || b := load(y, rlz);
                   || if b = 1 then
                   ||   c := bar();
```

The accesses to `y` ensure that in the generated traces the call to `bar` indeed appears only after `foo` returns. Nevertheless, speaking in C11 terminology, since these accesses are relaxed (as annotated by `rlz` mode), they do *not* induce a *happens-before* order between the calls, and thus, following the specification, the read from `x` inside `bar` may return 0 (the initial value) or 1. Accordingly, a “naive” library developer may attempt to efficiently implement L as follows (where $\oplus$ denotes a non-deterministic choice):

```
foo() :  return();
```

```
bar() :  return(0)  $\oplus$  return(1);
```

Is this implementation correct? In histories of the most general client of L that respects the library’s calling policy (i.e., in histories generated by executing the above program) one cannot observe a difference between the implementation and the specification. But, nevertheless, contextual refinement does *not* hold (so the implementation cannot be considered correct) for two different reasons illustrated by the following client programs:

```
store(z, 1, rlz); || b := load(y, rlz); // 1
foo();           || if b = 1 then
store(y, 1, rlz); ||   c := bar(); // 1
                 ||   d := load(z, rlz); // 0
```

```
foo();           || b := load(y, acq); // 1
store(y, 1, rel); || if b = 1 then
                 ||   c := bar(); // 0
```

The annotated behaviors in both examples are allowed when the suggested implementation of `foo` and `bar` is used, but disallowed when the specification is used. Indeed, in the example on the left, since the specification ensures release/acquire synchronization when `bar` returns 1 (via the `rel` and `acq` annotations in the accesses to `x`), the write to `z` happens-before the read from `z`, and in this case the model disallows the read from `z` to observe the overwritten initial value. In turn, in the example on the right, since the *client* uses release/acquire accesses, the write in `foo` happens-before the read in `bar`, which similarly ensures that the read from `x` cannot observe the initial value.

To address this challenge, we need to make library histories more expressive, so we can avoid the (SC) pitfall that identifies the execution order and the synchronization order. Concretely, in the

¹Interestingly, it follows from our result that for a language that employs *only* RC11-style relaxed accesses, inclusion of sets of histories is a sound condition (like it is for SC). Such a language is, however, too weak to be considered useful.

example above, we need to expose the facts that: (1) `bar` returning 1 *entails* a happens-before relation from *the call of* `foo` *to the return of* `bar`; and (2) `bar` returning 0 *forbids* a happens-before relation from *the return of* `foo` *to the call of* `bar`. What does that mean in operational terms without talking about the happens-before partial order between call and return actions? Our solution consists of the following:

- We introduce a novel operational semantics of the memory model and show that it is equivalent to the original declarative model. The operational semantics makes *knowledge propagation* between threads explicit. Roughly speaking, an access that is executed by one thread is first unknown to the other threads, and later non-deterministically propagates to each other threads. Access modes impose certain constraints on the propagation order and the ability to read from (locally) unknown events. In the simplest case—for programs with only release/acquire accesses—the propagation order has to follow the program order and reads can only read from known writes.
- We include propagation of method invocations and responses in the memory trace and, in turn, in library histories. These steps have no effect on the outcomes of the operational semantics (i.e., what values can be read and when), but they serve as crucial “markers” in library history, making histories expressive enough for validating contextual refinement given the inclusion of the sets of histories. In other words, we keep the inclusion of the sets of library histories as the correctness criterion, and recover its soundness by including appropriate markers (call and return propagation among different threads) in the histories.

In particular, revisiting the example above, the proposed implementation will have a history in which the propagation of the call marker of `foo` is after the return of `bar` but `bar` returns 1; as well as a history in which the propagation of the return marker of `foo` is before the call of `bar` but `bar` returns 0. Both histories are impossible for the specification, so library correctness does not hold.

2.2 Specification under Relaxed Memory Concurrency

A second challenge that we address is related to specification of libraries that expose weak behaviors. A straightforward approach to specifying concurrent data structures is to take a simple sequential implementation of the data structure, and wrap every operation inside a per-object lock. In fact, correctness according to the classical linearizability criterion [Herlihy and Wing 1990] is equivalent to refinement with respect to such lock-based specifications [Bouajjani et al. 2015; Filipović et al. 2010]. Since a lock entails a total order on all operations, the resulting specification is typically easy to understand and enables reasoning similar to reasoning about sequential programs.

This approach, however, has major shortcomings under relaxed memory concurrency, as it identifies the execution order with the synchronization order, which is only justified assuming SC. For instance, consider the following program that uses a concurrent queue (inspired by [Mével and Jourdan 2021]):

```
x := 1;           || a := q.dequeue; // 1
q.enqueue(1);    || b := x; // 0
```

Assuming a lock-based specification of the queue, the client may easily conclude that the annotated behavior ($a = 1$ and $b = 0$, where 0 is the initial value of x and the queue is initially empty) is impossible. Indeed, the fact that the dequeue operation returned the enqueued value entails that the lock of the dequeue operation must have been acquired after the lock of the enqueue operation was released. In turn, since locks necessarily provide release/acquire synchronization, the write to x happens-before the read from x , and thus the read cannot observe the overwritten initial value. Now, while the implementation of a queue may provide such synchronization guarantees for its client, it is also possible that the queue is implemented using efficient C11 relaxed accesses,

and thus comes with weaker guarantees that do not allow the client to rely on library-induced synchronization (so the accesses to x above are racy and $a = 1$ but $b = 0$ is possible). In that case the lock-based specification would be too strong.

We note that such weak guarantees do not imply that the queue is not “linearizable”: there may still exist a total order on all queue operations that agrees with the execution order (which can be any total order extending the union of the program order and the reads-from relation), and the queue exhibits FIFO behavior w.r.t. that order. A weak queue may be useful transferring elements of base types, whereas, in order to be used for transferring ownership via pointers, the client is responsible to appropriately place fences before and after invoking the weak queue methods.

Remark 2. While release/acquire synchronization may seem natural between enqueue and dequeue of the same element, the lock-based specification implies synchronization also in other cases, which may or may not be supplied (the queue studied in [Mével and Jourdan 2021] provides a case in point). For example, a client relying on a lock-based specification can deduce that the following annotated outcomes are all disallowed (where $\perp$ denotes an empty queue):

$x := 1;$ $q.enqueue(1);$	$\parallel$	$q.enqueue(2);$ $a := x; \text{ // } 0$	$\parallel$	$b := q.dequeue; \text{ // } 1$ $c := q.dequeue; \text{ // } 2$
------------------------------	-------------	--	-------------	--

$x := 1;$ $a := q.dequeue; \text{ // } \perp$	$\parallel$	$q.enqueue(1);$ $b := x; \text{ // } 0$
--	-------------	--

Accordingly, the challenge lies in identifying specification constructs that can be used instead of standard locks and be sufficiently flexible to account for different synchronization guarantees for the client. Ideally, like locks, they should allow straightforward reasoning about the library behaviors. We observe that existing concurrency models, C11 in particular, lack such constructs.

To address this challenge, we propose to extend the language with specialized instructions. Concretely, we identify that there is a range of possible guarantees that lie between C11’s release/acquire accesses (which provide synchronization) and relaxed accesses (which do not), and introduce a novel type of memory accesses that lies on this spectrum. Roughly speaking, these memory accesses, which we call *partial release/acquire*, ensure synchronization only for certain variables (see §4 for the formal definition).² C11’s release/acquire accesses can be seen as partial release/acquire accesses that ensure synchronization for *all* variables, whereas C11’s relaxed accesses can be seen as partial release/acquire accesses that do *not* ensure synchronization for *any* variable.

Partial release/acquire accesses can be used to construct locks that provide synchronization *only* for library variables, and these special locks can be used to specify libraries with weak synchronization guarantees. In particular, in the queue example above, locks implemented by *partial* release/acquire accesses would behave just like standard locks from library perspective, but will not allow the clients to rely on their induced synchronization. More generally, these accesses allow a library-internal release/acquire synchronization, which is needed for correct behaviors of the library, but is considered invisible from the point of view of the client.

3 Concurrent Programs: Syntax and Memory-Independent Semantics

In this section we begin to present the formal preliminaries for our results. As standard in memory models, it is convenient to break the semantics into: a *program* semantics (a.k.a. thread subsystem) and a *memory* semantics. Next, we focus on the program part, presenting syntax (§3.1), semantics (§3.2), and the synchronization with a (parametric) memory system (§3.3).

²We generally require libraries and clients to operate on different “address spaces” and never access the same shared variable. Lifting this simplifying assumption (e.g., as was done for SC in [Gotsman and Yang 2013]) is beyond our current scope.

3.1 Program Syntax

We employ a simple programming language, which supports the distinction between clients and libraries variable spaces.³ We use the following domains (and metavariables ranging over them):

(variable) spaces	$X, Y \in \text{Space} \triangleq \{X, Y, \dots\}$	variables	$x, y \in \text{Loc} \triangleq \{x, y, \dots\}$
thread identifiers	$\tau, \pi \in \text{Tid} = \{T_1, T_2, \dots, T_N\}$	registers	$r \in \text{Reg} = \{a, b, \dots\}$
read modes	$o_R \in \text{Mod}_R \triangleq \{\text{na}, \text{rlx}, \text{pacq}, \text{acq}\}$	values	$v \in \text{Val} \triangleq \{0, 1, 2, \dots\}$
write modes	$o_W \in \text{Mod}_W \triangleq \{\text{na}, \text{rlx}, \text{prel}, \text{rel}\}$	method names	$f \in F \quad \text{main} \notin F$

In particular, as we will see below, every memory access is to a particular variable in a particular variable space (like address spaces in operating systems); registers are thread-local variables; and the name `main` is reserved for non-library operations. Memory instructions have access modes (a.k.a. memory orderings), which determine their “strength” (na for non-atomics, `rlx` for relaxed, `prel/pacq` for partial release/acquire, and `rel/acq` for release/acquire). We also assume a relation $\sqsubset$ that orders modes according to their strength (following the left-to-right enumeration of the sets Mod_R and Mod_W above).

The language provides the following constructs, inspired by C11 atomics:

- Expressions are constructed with arithmetic and boolean operations over registers and values. We use e to range over expressions, and leave the exact expression grammar parametric.
- Thread local instructions (which do not interact with the shared-memory system):
 - Assignments of the form $r := e$, used for storing an expression e in a register r .
 - Conditionals of the form `if` e `goto` $n_1 \mid \dots \mid n_m$ (where $n_1, \dots, n_m \in \mathbb{N}$), used to non-deterministically jump to some program counter among $\{n_1, \dots, n_m\}$ when e evaluates to non-zero or, otherwise, skipping.
 - Additional local instructions (standard loop constructs, unconditional non-deterministic choice, etc.). Below, we will use such instructions in our examples, and their semantics should be clear.
- Shared-memory single accesses:
 - Write (a.k.a. store) instructions to memory of the form `store`(X, x, e, o_W) for storing into a shared variable x in space X the value that e evaluates to with access mode o_W .
 - Read (a.k.a. load) instructions from memory of the form $r := \text{load}(X, x, o_R)$ for loading the value from a shared variable x in space X into a register r with access mode o_R .
- Read-modify-write (RMW) instructions:
 - Fetch-and-add (FAA) instructions of the form $r := \text{FADD}(X, x, e, o_R, o_W)$ for atomically incrementing a variable x in space X by the value of e with read mode o_R and write mode o_W .
 - Compare-and-swap (CAS) instructions of the form $r := \text{CAS}(X, x, e_R, e_W, o_R, o_W, o_R^{\text{fail}})$. This instruction atomically loads the value from x in space X into r , compares it to the value e_R , and overwrites it by the value of e_W in case if the loaded value coincides with the value e_R . The load part will have mode o_R if comparison succeeds and o_R^{fail} otherwise; and the store part (if it happens) has mode o_W .
- Library interaction: `call`(f) for calling a method f and `return` for returning to the caller. For simplicity, we do not provide any argument passing mechanism and we will use the full register store for that matter. (If needed, each component may store the values it needs in the memory, and reload them later on.)

To construct programs we introduce three syntactic categories, each of which builds on the previous one:

³The partition to variable spaces is only needed to support the “partial release/acquire” access mode in library specifications, and can be completely ignored for specifications without such accesses.

- *Instruction sequences* represent the (sequential) implementation of each method (including `main`). Formally, an instruction sequence I is a function from a non-empty finite domain of the form $\{0, \dots, n\}$ (representing the possible program counters) to the set of instructions. We say that an instruction sequence is *flat* if it does not include `call`($_$) instructions.
- *Sequential programs* consist of a “main” method accompanied with implementations of every method $f \in F$. Formally, a sequential program sPr is a function assigning an instruction sequence to every $f \in \{\text{main}\} \cup F$. To avoid modeling a call stack and simplify the framework, we require that $sPr(f)$ is a flat instruction sequence for every $f \in F$.
- *Concurrent programs*, which we often call *programs*, are top-level parallel compositions of sequential programs, all accompanied by the same method implementations. Formally, a (concurrent) program Pr is a mapping assigning a sequential program to every $\tau \in \text{Tid}$, with $Pr(\tau)(f) = Pr(\pi)(f)$ for every $\tau, \pi \in \text{Tid}$ and $f \in F$.

In our examples, we often write instruction sequences as sequences of instructions delimited by “;”, and concurrent programs using ‘||’ between the main method of each thread. We also refer to the program threads as $\tau_1, \tau_2, \dots$ following their left-to-right order in the program listing.

3.2 Program Semantics

We give semantics to the syntactic objects above using labeled transition systems.

Definition 3.1. A *labeled transition system* (LTS) is a tuple $A = \langle \Sigma, Q, q_0, T \rangle$, where Σ is a set of *transition labels*, Q is a set of *states*, $q_0 \in Q$ is the *initial state*, and $T \subseteq Q \times \Sigma \times Q$ is a set of *transitions*. We denote by $A.\Sigma$, $A.Q$, $A.q_0$, and $A.T$ the components of an LTS A . We write $q \xrightarrow{\sigma} q'$ to denote a transition $\langle q, \sigma, q' \rangle \in T$, and $\xrightarrow{\sigma}_A$ for the relation $\{\langle q, q' \rangle \mid q \xrightarrow{\sigma} q' \in A.T\}$, and $\rightarrow_A$ for $\bigcup_{\sigma \in \Sigma} \xrightarrow{\sigma}_A$. For a sequence $t \in A.\Sigma^*$, we write $\xrightarrow{t}_A$ for the composition $\xrightarrow{t(1)}_A ; \dots ; \xrightarrow{t(|t|)}_A$. A sequence $t \in A.\Sigma^*$ such that $A.q_0 \xrightarrow{t}_A q$ for some $q \in A.Q$ is called a *trace* of A . We denote by $\text{traces}(A)$ the set of all traces of A . A state $q \in A.Q$ is called *reachable* in A if $A.q_0 \xrightarrow{t}_A q$ for some $t \in \text{traces}(A)$. For a trace t and a set $\Theta \subseteq \Sigma$ of transition labels, we write $t|_{\Theta}$ for the longest subsequence of t over Θ .

Next, we define the LTSs induced by instruction sequences, sequential programs, and concurrent programs. We often identify the syntactic objects with the LTS they induce (e.g., when writing expressions like $sPr.Q$ for a sequential program sPr). The transition labels of these LTSs feature *action labels*, which represent the interactions that a program may have with the memory.

Definition 3.2. An *action label* l takes one of the following forms: a read $R(X, x, v_R, o_R)$, a write $W(X, x, v_W, o_W)$, a read-modify-write $RMW(X, x, v_R, v_W, o_R, o_W)$, a call $CALL(f, \phi)$, and a return $RET(\phi)$, where $X \in \text{Space}$, $x \in \text{Loc}$, $v_R, v_W \in \text{Val}$, $o_R \in \text{Mod}_R$, $o_W \in \text{Mod}_W$, $f \in F$, and $\phi : \text{Reg} \rightarrow \text{Val}$. We denote by Lab the set of all action labels. The functions typ , sp , loc , val_R , val_W , mod_R , mod_W , callee , and store respectively retrieve (when applicable) the type (R/W/...), space (X), variable (x), read value (v_R), written value (v_W), read mode (o_R), write mode (o_W), callee method name (f), and store (ϕ) of an action label.

Next, we define the LTS induced by an instruction sequence.

Definition 3.3. An *instruction sequence state* is a pair $\langle pc, \phi \rangle$, where $pc \in \mathbb{N}$, called *program counter*, stores the current instruction pointer inside the sequence, and $\phi : \text{Reg} \rightarrow \text{Val}$, called *local store*, records the values of the registers. Local stores are extended to apply on expressions in the standard way. The *LTS induced by an instruction sequence* I is an LTS over instruction sequence states with: $\text{Lab}_{\epsilon} \triangleq \text{Lab} \cup \{\epsilon\}$ as the set of transition labels (that is, the set of all action labels extended with ϵ for silent transitions); $\langle 0, \phi_{\text{init}} \rangle$ where $\phi_{\text{init}} \triangleq \lambda r. 0$ as the initial state; and the transitions as given in Fig. 1 (additional thread local instructions can be standardly added).

$$\begin{array}{c}
\frac{I(pc) = r := e \quad \phi' = \phi[r \mapsto \phi(e)]}{\langle pc, \phi \rangle \xrightarrow{\epsilon}_I \langle pc + 1, \phi' \rangle} \\
\\
\frac{I(pc) = \text{if } e \text{ goto } n_1 \mid \dots \mid n_m \quad \phi(e) \neq 0 \implies pc' \in \{n_1, \dots, n_m\} \quad \phi(e) = 0 \implies pc' = pc + 1}{\langle pc, \phi \rangle \xrightarrow{\epsilon}_I \langle pc', \phi \rangle} \\
\\
\frac{I(pc) = r := \text{load}(X, x, o_R) \quad l = R(X, x, v, o_R) \quad \phi' = \phi[r \mapsto v]}{\langle pc, \phi \rangle \xrightarrow{l}_I \langle pc + 1, \phi' \rangle} \\
\\
\frac{I(pc) = \text{store}(X, x, e, o_W) \quad l = W(X, x, \phi(e), o_W)}{\langle pc, \phi \rangle \xrightarrow{l}_I \langle pc + 1, \phi \rangle} \quad \frac{I(pc) = r := \text{FADD}(X, x, e, o_R, o_W) \quad l = \text{RMW}(X, x, v, v + \phi(e), o_R, o_W) \quad \phi' = \phi[r \mapsto v]}{\langle pc, \phi \rangle \xrightarrow{l}_I \langle pc + 1, \phi' \rangle} \\
\\
\frac{I(pc) = r := \text{CAS}(X, x, e_R, e_W, o_R, o_W, o_R^{\text{fail}}) \quad v = \phi(e_R) \implies l = \text{RMW}(X, x, v, v, \phi(e_W), o_R, o_W) \quad v \neq \phi(e_R) \implies l = R(X, x, v, o_R^{\text{fail}}) \quad \phi' = \phi[r \mapsto v]}{\langle pc, \phi \rangle \xrightarrow{l}_I \langle pc + 1, \phi' \rangle}
\end{array}$$

Fig. 1. Transitions of LTS induced by an instruction sequence

Recall that program semantics is separate from memory semantics, which is why the read and RMW transitions in Fig. 1 can observe any value. It is only important that each transition that interacts with the memory announces itself in the transition label. The `call`(`_`) and `return` instructions are not handled at the level of instruction sequences, but receive special semantics at the level of sequential programs, as defined next.

Definition 3.4. A *sequential program state* is a tuple $q = \langle pc, \phi, pc_s, f \rangle$, where: $\langle pc, \phi \rangle$ is an instruction sequence state storing the state of the sequence currently running; $pc_s \in \mathbb{N} \cup \{\perp\}$, called *the stored program counter*, is used to remember the program position to jump to when the current instruction sequence returns ($pc_s = \perp$ means that the main method is currently running); and $f \in F \cup \{\text{main}\}$, called *the active method*, tracks the method that is currently running. We denote by $q.pc$, $q.\phi$, $q.pc_s$, and $q.f$ the components of a sequential program state q .

Definition 3.5. The *LTS induced by a sequential program* sPr is given by:

- The set of transition labels is $\text{Lab}_\epsilon \times (F \cup \{\text{main}\})$. The functions `lab` and `method` respectively retrieve the action label (or ϵ) and method name of a transition label. All functions on action labels (`typ`, `sp`, `loc`, ...) are lifted to sequential program transition labels in the obvious way.
- The states are sequential program states, as defined in Def. 3.4.
- The initial state is $\langle 0, \phi_{\text{init}}, \perp, \text{main} \rangle$.
- The transitions are given by:

$$\begin{array}{c}
\frac{f \in \{\text{main}\} \cup F \quad \langle pc, \phi \rangle \xrightarrow{l_\epsilon}_{sPr(f)} \langle pc', \phi' \rangle}{\langle pc, \phi, pc_s, f \rangle \xrightarrow{l_\epsilon, f}_{sPr} \langle pc', \phi', pc_s, f \rangle} \\
\\
\frac{sPr(\text{main})(pc) = \text{call}(f) \quad l = \text{CALL}(f, \phi)}{\langle pc, \phi, \perp, \text{main} \rangle \xrightarrow{l, \text{main}}_{sPr} \langle 0, \phi, pc + 1, f \rangle} \quad \frac{sPr(f)(pc) = \text{return} \quad l = \text{RET}(\phi)}{\langle pc, \phi, pc_s, f \rangle \xrightarrow{l, f}_{sPr} \langle pc_s, \phi, \perp, \text{main} \rangle}
\end{array}$$

The first transition, which applies for any method (`main` or other), lifts the instruction-sequence transition to the level of sequential programs. The second transition passes control from the main method to some other method, jumping the program counter to the first instruction and storing the return point ($pc + 1$). Finally, the third transition passes control back using the stored return point. (We do not need to record a call stack since we assume that $sPr(f)$ is flat for every $f \in F$.)

Finally, the LTS induced by a concurrent program interleaves the thread transitions.

Definition 3.6. A (concurrent) program state $\bar{p}$ is a mapping assigning a sequential program state to every $\tau \in \text{Tid}$. The LTS induced by a program Pr is an LTS over program states, with: $\text{ProgLab} \triangleq \text{Tid} \times \text{Lab}_\epsilon \times (\text{F} \cup \{\text{main}\})$ as the set of transition labels; $\bar{p}_{\text{init}} \triangleq \lambda\tau. \langle 0, \phi_{\text{init}}, \perp, \text{main} \rangle$ as the initial state; and the following transitions:

$$\frac{\bar{p}(\tau) \xrightarrow{l_\epsilon f}_{Pr(\tau)} q'}{\bar{p} \xrightarrow{\tau, l_\epsilon f}_{Pr} \bar{p}[\tau \mapsto q']}$$

Below, for a program transition label $\alpha \in \text{ProgLab}$, the functions tid , lab , and method respectively retrieve the thread identifier (τ), the action label (or ϵ) (l_ϵ), and the method name (f) of α . Functions on action labels (typ , sp , loc , ...) are lifted to program transition labels in the obvious way.

3.3 Synchronizing Programs and Memories

To give semantics to programs running under a particular memory model, we synchronize the transitions of a program Pr with a memory system. For now, we leave the memory system parametric, and assume it is represented by some LTS $\mathcal{M}$ whose set of transition labels consists of non-silent program transition labels (elements of $\text{Tid} \times \text{Lab} \times (\text{F} \cup \{\text{main}\})$) as well as a (disjoint) set $\mathcal{M}.\Theta$ of internal memory actions (which we use later for non-deterministic propagation of knowledge between threads).

Definition 3.7. The composition of a program Pr and a memory system $\mathcal{M}$, denoted by $Pr \bowtie \mathcal{M}$, is the LTS whose transition labels are the elements of $\text{ProgLab} \cup \mathcal{M}.\Theta$; states are pairs $\langle \bar{p}, M \rangle \in Pr.Q \times \mathcal{M}.Q$; initial state is $\langle \bar{p}_{\text{init}}, \mathcal{M}.q_0 \rangle$; and transitions are given by:

$$\begin{array}{c} \frac{\alpha \in \text{Tid} \times \text{Lab} \times (\text{F} \cup \{\text{main}\}) \quad \bar{p} \xrightarrow{\alpha}_{Pr} \bar{p}' \quad M \xrightarrow{\alpha}_{\mathcal{M}} M'}{\langle \bar{p}, M \rangle \xrightarrow{\alpha}_{Pr \bowtie \mathcal{M}} \langle \bar{p}', M' \rangle} \quad \frac{\alpha \in \text{Tid} \times \{\epsilon\} \times (\text{F} \cup \{\text{main}\}) \quad \bar{p} \xrightarrow{\alpha}_{Pr} \bar{p}'}{\langle \bar{p}, M \rangle \xrightarrow{\alpha}_{Pr \bowtie \mathcal{M}} \langle \bar{p}', M \rangle} \quad \frac{\alpha \in \mathcal{M}.\Theta \quad M \xrightarrow{\alpha}_{\mathcal{M}} M'}{\langle \bar{p}, M \rangle \xrightarrow{\alpha}_{Pr \bowtie \mathcal{M}} \langle \bar{p}, M' \rangle} \end{array}$$

The above transitions are “synchronized transitions” of Pr and $\mathcal{M}$, using the labels to decide what to synchronize on. Both the program and the memory take the same step for transition labels that are common to both LTSs, only the program steps for transition labels that are program internal (i.e., with $\text{lab}(\alpha) = \epsilon$) and only the memory steps for transition labels that are memory internal.

Example 3.8. The most well-known memory system is the one of sequential consistency, denoted here by SC. This memory system simply tracks the most recent value written to each variable, and has no internal transitions ($\text{SC}.\Theta = \emptyset$). Formally, it is defined by $\text{SC}.Q \triangleq (\text{Space} \times \text{Loc}) \rightarrow \text{Val}$, $\text{SC}.q_0 \triangleq \lambda X, x. 0$, and $\rightarrow_{\text{SC}}$ is given by:

$$\begin{array}{c} \frac{l = R(X, x, v_R, _) \quad m(X, x) = v_R}{m \xrightarrow{\tau, l, f}_{\text{SC}} m} \quad \frac{l = W(X, x, v_W, _) \quad m' = m[(X, x) \mapsto v_W]}{m \xrightarrow{\tau, l, f}_{\text{SC}} m'} \quad \frac{l = \text{RMW}(X, x, v_R, v_W, _) \quad m(X, x) = v_R \quad m' = m[(X, x) \mapsto v_W]}{m \xrightarrow{\tau, l, f}_{\text{SC}} m'} \quad \frac{\text{typ}(l) \in \{\text{CALL}, \text{RET}\}}{m \xrightarrow{\tau, l, f}_{\text{SC}} m} \end{array}$$

4 The dRC11 Memory Model

In this section we introduce the weak memory model that we assume in this paper. This model, which we call dRC11, is a declarative (a.k.a. axiomatic) model forming an extension of (a fragment of) the RC11 model [Lahav et al. 2017] with specialized semantics for the novel `prel`/`pacq` accesses.

NOTATION 4.1 (RELATIONAL NOTATIONS). Given a (binary) relation R , $\text{dom}(R)$ and $\text{codom}(R)$ denote its domain and codomain, and $R^?$, R^+ , and R^* denote its reflexive, transitive, and reflexive-transitive closures. The inverse of a relation R is denoted by R^{-1} , and the (left) composition of two

relations R_1 and R_2 is denoted by $R_1 ; R_2$. We denote by $[A]$ the identity relation on a set A . In particular, $[A] ; R ; [B] = R \cap (A \times B)$. When A is finite, we write $[a_1, \dots, a_n]$ instead of $[\{a_1, \dots, a_n\}]$.

We start by defining *execution graphs*. Their nodes, called *events*, represent memory accesses, and their directed edges are of different kinds: *program order* represents the order imposed by the program; *reads-from* mapping maps each read event to the write event it obtains its value from; and *modification order* (a.k.a. *coherence order*) provides a total order on the writes to every variable. The precise definitions are given next.

Definition 4.2. An event e is a tuple $\langle \tau, s, l, f \rangle$, where $\tau \in \text{Tid} \uplus \{\perp\}$, called the event's *thread identifier* ($\perp$ is used for initialization events); $s \in \mathbb{N}$, called the event's *serial identifier*; $l \in \text{Lab}$, called the event's *label* (as defined in Def. 3.2), and $f \in F \cup \{\text{main}\}$, called the event's *method*. The functions tid , sn , lab , and method return the thread identifier (τ), identifier (s), action label (l), and method of an event (f). All functions on action labels (typ , sp , loc , ...) are lifted to events in the obvious way. We denote by E the set of all events, and define the following subsets:

$$\begin{aligned} R &\triangleq \{e \in E \mid \text{typ}(e) \in \{R, \text{RMW}\}\} & W &\triangleq \{e \in E \mid \text{typ}(e) \in \{W, \text{RMW}\}\} & \text{RMW} &\triangleq R \cap W \\ \text{CALL} &\triangleq \{e \in E \mid \text{typ}(e) = \text{CALL}\} & \text{RET} &\triangleq \{e \in E \mid \text{typ}(e) = \text{RET}\} & \text{CR} &\triangleq \text{CALL} \cup \text{RET} \end{aligned}$$

We employ subscripts and superscripts to restrict sets of events to certain properties, e.g., $W_X = \{w \in W \mid \text{sp}(w) = X\}$, $R_{X,x}^{\text{pacq}} = \{r \in R \mid \text{sp}(r) = X \wedge \text{loc}(r) = x \wedge \text{mod}_W(r) \sqsupseteq \text{pacq}\}$, $\text{CALL}_F = \{e \in \text{CALL} \mid \text{callee}(e) \in F\}$, $\text{RET}_F = \{e \in \text{RET} \mid \text{method}(e) \in F\}$, $\text{CR}_F = \text{CALL}_F \cup \text{RET}_F$, and $E^\tau = \{e \in E \mid \text{tid}(e) = \tau\}$ for any $E \subseteq E$. The set Init of *initialization events* is given by:

$$\text{Init} \triangleq \{\langle \perp, 0, W(X, x, 0, \text{rlx}), \text{main} \rangle \mid X \in \text{Space}, x \in \text{Loc}\}.$$

Definition 4.3. An *execution graph* G is a tuple $\langle E, po, rf, mo \rangle$, where:

- E is a finite set of events, such that $\text{Init} \subseteq E$ and $\text{tid}(e) \neq \perp$ for every $e \in E \setminus \text{Init}$.
- po is a *program order* for E , that is: $po = (\biguplus_{\tau \in \text{Tid}} po_\tau) \uplus (\text{Init} \times (E \setminus \text{Init}))$, for some relations po_τ , such that each po_τ is a strict total order on E^τ .
- rf is a *reads-from relation* for E , that is a relation on E satisfying:
 - If $\langle w, r \rangle \in rf$, then $w \in W$ and $r \in R$.
 - If $\langle w, r \rangle \in rf$, then $\text{sp}(w) = \text{sp}(r)$, $\text{loc}(w) = \text{loc}(r)$, and $\text{val}_W(w) = \text{val}_R(r)$.
 - $w_1 = w_2$ whenever $\langle w_1, r \rangle, \langle w_2, r \rangle \in rf$ (each read reads from at most one write).
 - $E \cap R \subseteq \text{codom}(rf)$ (each read reads from some write).
- mo is a *modification order* for E , that is: $mo = \biguplus_{X \in \text{Space}, x \in \text{Loc}} mo_{X,x}$, for some relations $mo_{X,x}$, such that each $mo_{X,x}$ is a strict total order on $E \cap W_{X,x}$.

We denote the components of G by $G.E$, $G.po$, $G.rf$, and $G.mo$. For any set $E' \subseteq E$, we write $G.E'$ for $G.E \cap E'$ (e.g., $G.W = G.E \cap W$).

To formally associate execution graphs with programs we use a memory system called FG (for “free graphs”), whose states are execution graphs. This system allows *all* possible program transitions, while recording them in its state, which is the current execution graph with (almost) arbitrary reads-from and modification order relations.

Definition 4.4. The memory system FG is the LTS whose transition labels are program transition labels (i.e., there are not any internal memory actions); states are execution graphs; initial state is G_{Init} , defined by $G_{\text{Init}}.E \triangleq \text{Init}$ and $G_{\text{Init}}.X \triangleq \emptyset$ for every other component of G ; and transitions are given in Fig. 2.

The transitions of FG are based on helper notations used to extend an execution graph G with a fresh event e at the end of the executing thread. For memory accesses, it requires to pick a write

$$\begin{array}{c}
\text{WRITE/READ/RMW} \\
\text{typ}(l) \in \{W, R, \text{RMW}\} \quad e = \text{NextEvent}(G.E, \tau, l, f) \\
w \in G.W_{\text{sp}(e), \text{loc}(e)} \quad e \in R \implies \text{val}_W(w) = \text{val}_R(e) \\
G' = \text{Add}(G, e, w) \\
\hline
G \xrightarrow{\tau, l, f}_{\text{FG}} G'
\end{array}
\qquad
\begin{array}{c}
\text{CALL/RETURN} \\
\text{typ}(l) \in \{\text{CALL}, \text{RET}\} \\
e = \text{NextEvent}(G.E, \tau, l, f) \\
G' = \text{Add}(G, e) \\
\hline
G \xrightarrow{\tau, l, f}_{\text{FG}} G'
\end{array}$$

Fig. 2. Transitions of FG system for generating all candidate execution graphs

event w in G , called the *write-predecessor* of e , that is: (1) the **rf**-source of e if e is a read; (2) the **mo**-immediate predecessor of e if e is a write; and (3) both the **rf**-source and the **mo**-immediate predecessor of e if e is an RMW. For that matter, we employ the following notations:

NOTATION 4.5. Given a set E of events, $\tau \in \text{Tid}$, $l \in \text{Lab}$, and $f \in F$, $\text{NextEvent}(E, \tau, l, f)$ denotes the event with thread identifier τ , label l , method f , and a minimal fresh serial identifier w.r.t. E , that is: $\text{NextEvent}(E, \tau, l, f) \triangleq \langle \tau, s, l, f \rangle$, where $s = \min\{n \in \mathbb{N} \mid \langle \tau, n, l, f \rangle \notin E\}$.

NOTATION 4.6. For an execution graph G and events e and w , $\text{Add}(G, e, w)$ denotes the tuple $\langle E', po', rf', mo' \rangle$, where:

$$\begin{aligned}
E' &= G.E \cup \{e\} & po' &\triangleq G.po \cup ((G.E^{\text{tid}(e)} \cup \text{Init}) \times \{e\}) \\
rf' &= \begin{cases} G.rf \cup \{\langle w, e \rangle\} & e \in R \\ G.rf & \text{otherwise} \end{cases} & mo' &= \begin{cases} G.mo \cup \text{dom}(G.mo^?; [w]) \times \{e\} \\ \cup \{e\} \times \text{codom}([w]; G.mo) & e \in W \\ G.mo & \text{otherwise} \end{cases}
\end{aligned}$$

Similarly, for an execution graph G and event e , $\text{Add}(G, e)$ denotes the tuple $\langle E', po', G.rf, G.mo \rangle$, where E' and po' are defined as above.

In the sequel, it will be useful to note that all execution graphs generated by FG when synchronized with a given program satisfy the following well-formedness property:

Definition 4.7. An execution graph G is *well-formed* if the following hold for every $f \in F$:

- $\{\{e \in E \mid \text{method}(e) \neq f\}\}; G.po; \{\{e \in E \mid \text{method}(e) = f\}\} \subseteq G.po^?; [\text{CALL}_{\{f\}}]; G.po.$
- $\{\{e \in E \mid \text{method}(e) = f\}\}; G.po; \{\{e \in E \mid \text{method}(e) \neq f\}\} \subseteq G.po^?; [\text{RET}_{\{f\}}]; G.po.$

PROPOSITION 4.8. If $\langle \bar{p}, G \rangle$ is reachable in $\text{Pr} \bowtie \text{FG}$ for some program Pr , then G is well-formed.

Now, to filter for *consistent* graphs among all candidate execution graphs generated by FG for a given program, we define several derived relations (some parametrized by $X \in \text{Space}$):

$$\begin{aligned}
G.\text{fr} &\triangleq (G.\text{rf}^{-1}; G.\text{mo}) \setminus [E] && \text{(from-read, a.k.a. read-before)} \\
G.\text{sw}_{\text{base}} &\triangleq [W^{\text{rel}}]; G.\text{rf}^+; [R^{\text{acq}}] && \text{(global synchronization)} \\
G.\text{sw}_X &\triangleq [W_X^{\neg \text{prel}}]; G.\text{rf}^+; [R_X^{\neg \text{pacq}}] && \text{(per-space synchronization)} \\
G.\text{hb}_X &\triangleq (G.po \cup G.\text{sw}_{\text{base}} \cup G.\text{sw}_X)^+ && \text{(per-space happens-before)}
\end{aligned}$$

The **fr** relation is standard in weak memory models [Alglave et al. 2014], relating every read (or RMW) to subsequent writes (or RMWs) as dictated by **rf** and **mo** (every write w that is **mo**-after the **rf**-source of a read r is **fr**-after r). The **sw**_{base} and **sw**_X relations formally capture synchronization: **sw**_{base} is for “global” synchronization which is formed by reads-from edges between **rel** and **acq** accesses (just like in C11 [Lahav et al. 2017]), whereas **sw**_X is for “space-internal” synchronization that affects only accesses to space X and is formed by reads-from edges *in space* X between **prel**

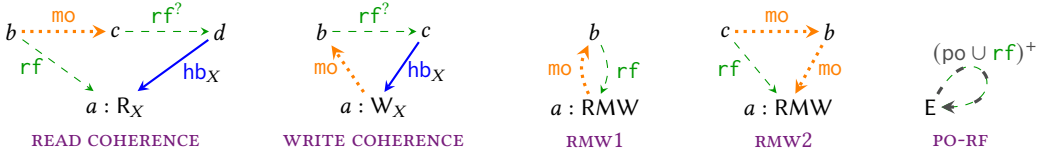


Fig. 3. Illustration of forbidden patterns in dRC11

and **pacq**. The use of rf^+ (rather than just rf) is for supporting *release sequences* as in C11, which ensures the synchronization between a release write w and an acquire read r also if there is chain of reads-from edges between them ($\langle w, u_1 \rangle, \langle u_1, u_2 \rangle, \dots, \langle u_{n-1}, u_n \rangle, \langle u_n, r \rangle \in rf$ where $u_1, \dots, u_n$ are RMWs, which can be relaxed).⁴ Finally, paths composed of the program order (po) and the sw_{base} and sw_X relations form the per-space *happens-before* relation, which again refines the C11 happens-before relation (defined as $(G.po \cup G.sw_{base})^+$).

Using the above definitions, the consistency of an execution graph is defined as follows.

Definition 4.9. An execution graph G is *dRC11-consistent* if the following hold:

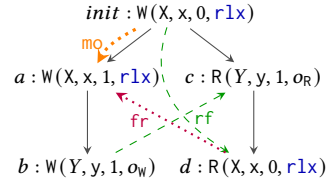
- For every $X \in \text{Space}$, $[R_X]$; $G.fr$; $G.rf^?$; $G.hb_X$ is irreflexive. (READ COHERENCE)
- For every $X \in \text{Space}$, $[W_X]$; $G.mo$; $G.rf^?$; $G.hb_X$ is irreflexive. (WRITE COHERENCE)
- $G.mo$; $G.rf$ is irreflexive. (RMW1)
- $G.fr$; $G.mo$ is irreflexive. (RMW2)
- $G.po \cup G.rf$ is acyclic. (PO-RF)

These constraints, depicted in Fig. 3, are variants of the ones in (R)C11, where the only essential difference is in the coherence constraints that here use $G.hb_X$ for restricting accesses to X .⁵ The **RMW1** and **RMW2** constraints are needed for ensuring the right behavior of RMWs including their atomicity. The **PO-RF** constraint is an addition of RC11 on top of C11, which is a conservative solution to the “out-of-thin-air” problem that arises if $po \cup rf$ -cycles are allowed [Batty et al. 2015; Kang et al. 2017].

The next example demonstrates the role of the coherence constraints:

Example 4.10. Consider the following standard “message-passing” litmus test (parametric in the space Y and the access modes o_W, o_R):

```
store(X, x, 1, rlx); || a := load(Y, y, o_R); // 1
store(Y, y, 1, o_W); || b := load(X, x, rlx); // 0
```



The annotated behavior is disallowed only if the synchronization on y (in space Y) is visible for the read of x (in space X). This will be the case only if (1) $o_W = \text{rel}$ and $o_R = \text{acq}$; or (2) $o_W \sqsupseteq \text{prel}$, $o_R \sqsupseteq \text{pacq}$, and $Y = X$ (i.e., the four accesses are to the same space). (In particular, if $Y \neq X$, then only $o_W = \text{rel}$ and $o_R = \text{acq}$ would forbid the annotated behavior.) To see how this follows from

⁴Another component of C11’s release sequence, which allows forming synchronization using a relaxed write *po*-after a release write to the same location, was omitted from the standard in C++20. (See https://en.cppreference.com/w/cpp/atomic/memory_order [Accessed July 2022].)

⁵There are some presentational differences w.r.t. the model in [Lahav et al. 2017]: (1) RC11 uses two events related by an *rmw*-edge to represent RMWs, so **RMW1** is not needed and **RMW2** has a different formulation (called **ATOMICITY** in [Lahav et al. 2017]); and (2) RC11 uses the “extended coherence order” (**eco**) which allows it to merge both coherence constraints.

the coherence constraints, note that these conditions are needed to ensure $\langle b, c \rangle \in \text{sw}_{\text{base}} \cup \text{sw}_x$, which in turn implies $\langle a, d \rangle \in \text{hb}_x$. Now, since $\langle \text{init}, a \rangle \in \text{po} \subseteq \text{hb}_x$, **WRITE COHERENCE** ensures that $\langle \text{init}, a \rangle \in \text{mo}$, and so $\langle d, a \rangle \in \text{fr}$. In turn, **READ COHERENCE** disallows hb_x from a to d .

By connecting the definition of candidate execution graphs for a given program (using the FG system) and dRC11-consistency, we define the allowed program behaviors under dRC11 (more precisely, possible reachable program states).

Definition 4.11. A program state $\bar{p}$ is *reachable for a program Pr under dRC11* if $\langle \bar{p}, G \rangle$ is reachable in $Pr \bowtie \text{FG}$ for some dRC11-consistent execution graph G .

Finally, again following (R)C11, data races on non-atomics are considered as programming errors (so non-atomic accesses can be heavily optimized by the hardware and the compiler). Accordingly, we define racy execution graphs and racy programs.

Definition 4.12. Two events e_1 and e_2 form a *race* in an execution graph G , if $e_1, e_2 \in G.W \cup G.R$, $e_1 \neq e_2$, $\text{sp}(e_1) = \text{sp}(e_2)$, $\text{loc}(e_1) = \text{loc}(e_2)$, $|\{e_1, e_2\} \cap W| \geq 1$, $|\{e_1, e_2\} \cap E^{\text{na}}| \geq 1$, and $\langle e_1, e_2 \rangle \notin G.\text{hb}_{\text{sp}(e_1)} \cup G.\text{hb}_{\text{sp}(e_1)}^{-1}$. An execution graph G is *racy* if some two events form a race in G .

Definition 4.13. A program Pr is *racy* under dRC11 if $\langle \bar{p}, G \rangle$ is reachable in $Pr \bowtie \text{FG}$ for some program state $\bar{p}$ and racy dRC11-consistent execution graph G .

5 The Operational Memory System: pRC11

In this section we introduce an operational version of the dRC11 memory model, called pRC11, which exposes knowledge propagation steps as internal memory steps, and is particularly suitable for library abstraction.

Since dRC11-consistency is *prefix-closed* w.r.t. $\text{po} \cup \text{rf}$ [Kokologiannakis et al. 2017], dRC11 can be easily made operational by adapting the LTS FG (Def. 4.4) to require dRC11-consistency after each step of the execution graph generation (rather than one time at the end). However, following §2.1, we opt for a more elaborate model with explicit point-to-point *propagation transitions* marking the steps in which some event of thread τ becomes visible to another thread π . In particular, this allows us to observe the propagation of the call/return events in memory traces, which is essential in our definition of the library correctness condition in §7. We expose propagation transitions in traces of pRC11 as internal memory steps labeled with *propagation labels* as defined next.

Definition 5.1. A *propagation label* is a triple, denoted by $p = \text{EP}(e, \tau, X)$, where $e \in E$ (propagated event), $\tau \in \text{Tid} \setminus \{\text{tid}(e)\}$ (destination thread identifier), and $X \in \text{Space}$ (destination space). We use E , ptid and psp to retrieve the components (e , τ , and X , respectively) of a propagation label p . All functions on events and action labels (tid , typ ...) are lifted to propagation labels in the obvious way. We denote by PLab the set of all propagation labels.

Then, the pRC11 memory system is defined as follows.

Definition 5.2. The pRC11 memory system is an LTS whose set of transition labels is $\text{ProgLab} \cup \text{PLab}$ (i.e., $\text{pRC11}.\Theta = \text{PLab}$); states are pairs of the form $M = \langle G, K \rangle$, where G is an execution graph and K is a *knowledge mapping* for G , that is a function in $\text{Tid} \rightarrow \text{Space} \rightarrow \mathcal{P}(G.E)$; initial state is $M_{\text{init}} \triangleq \langle G_{\text{init}}, K_{\text{init}} \rangle$, where $K_{\text{init}} \triangleq \lambda \tau. \lambda X. \text{Init}(G_{\text{init}})$ is defined in Def. 4.4; and transitions are as given in Fig. 4.

In addition to the current execution graph G , pRC11's states record a *knowledge mapping* K that records for each thread τ and space X all events that have already propagated to τ for space X (as well as all events associated with actions executed by τ itself). We refer to the set $K(\tau)(X)$ as the X -knowledge of thread τ .

$$\begin{array}{c}
\text{WRITE/READ/RMW} \\
\begin{array}{l}
\text{typ}(l) \in \{W, R, \text{RMW}\} \quad e = \text{NextEvent}(G.E, \tau, l, f) \quad X = \text{sp}(e) \\
w \in G.W_{X, \text{loc}(e)} \quad e \in R \implies \text{val}_W(w) = \text{val}_R(e) \\
w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(X)]) \quad e \in W \implies w \notin \text{dom}(G.\text{rf}; [\text{RMW}]) \\
e \in R^{\text{pacq}} \implies \text{dom}(G.\text{rf}^*; [w]) \subseteq K(\tau)(X) \quad e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G.\text{rf}^*; [w]) \subseteq K(\tau)(Y) \\
G' = \text{Add}(G, e, w) \quad k' = \lambda Y. K(\tau)(Y) \cup \{e\} \quad K' = K[\tau \mapsto k']
\end{array} \\
\hline
\langle G, K \rangle \xrightarrow{\tau, l, f}_{\text{pRC11}} \langle G', K' \rangle
\end{array}$$

$$\begin{array}{c}
\text{CALL/RETURN} \\
\begin{array}{l}
\text{typ}(l) \in \{\text{CALL}, \text{RET}\} \quad e = \text{NextEvent}(G.E, \tau, l, f) \\
G' = \text{Add}(G, e) \quad k' = \lambda Y. K(\tau)(Y) \cup \{e\} \quad K' = K[\tau \mapsto k']
\end{array} \\
\hline
\langle G, K \rangle \xrightarrow{\tau, l, f}_{\text{pRC11}} \langle G', K' \rangle
\end{array}$$

$$\begin{array}{c}
\text{PROPAGATE} \\
\begin{array}{l}
e \in G.E \setminus K(\tau)(X) \\
e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G.\text{hb}_X; [e]) \subseteq K(\tau)(X) \\
k' = K(\tau)[X \mapsto K(\tau)(X) \cup \{e\}] \quad K' = K[\tau \mapsto k']
\end{array} \\
\hline
\langle G, K \rangle \xrightarrow{\text{EP}(e, \tau, X)}_{\text{pRC11}} \langle G, K' \rangle
\end{array}$$

Fig. 4. Transitions of pRC11.

The WRITE/READ/RMW transition in Fig. 4 executes a memory access by thread τ , by adding a corresponding event e to the current execution graph while imposing certain conditions on w , the write-predecessor of e . Intuitively, the main imposed condition, $w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(X)])$, requires that w is not overwritten by any other write that τ is already aware of for the space of e . More precisely, if e is an access in space X , then τ should not have in its X -knowledge any write that is **mo**-later than w or any read that reads from a write that is **mo**-later than w . In addition:

- If e is a write (or RMW), then w should not be already read by an RMW. This condition is needed to ensure the atomicity of RMWs (corresponds to **rmw2** in Def. 4.9).
- If e is a read (or RMW) with **acq** or **pacq** mode, then w (the write that e reads from) has to be already present in the thread's knowledge: in its X -knowledge if e is **pacq**, or in Y -knowledge for all Y if e is **acq**. (Moreover, if w is an RMW, to account for release sequences, this should hold not only for w , but also for every event on the **rf**-chain entering w .)

Finally, every thread certainly knows about its own actions, so in addition to extending G , this step also extends the knowledge of τ by adding the event e to all spaces.

The CALL/RETURN transition is simple: it adds a corresponding event e to the current execution graph and extends the knowledge of the executing thread to include e .

The PROPAGATE transition is a non-deterministic internal memory step that extends the threads' knowledge. It picks some event e that is not in τ 's X -knowledge and adds e to $K(\tau)(X)$. When e is a **prel**-write to space X , a **rel**-write (to any space), or a call or return marker, then the propagation of e to the X -knowledge of thread τ can be done only after all accesses to X , as well as all call and return markers, that are $G.\text{hb}_X$ -before e have propagated to the X -knowledge of thread τ . What is not constrained is equally important: for instance, relaxed writes can propagate “out-of-order”.

Example 5.3. pRC11's transitions are best understood via the “message-passing” litmus test presented in Example 4.10. Consider first the case that $o_W = \text{rel}$ and $o_R = \text{acq}$. Then, due to the constraints on acquire reads and release writes: (1) since the read of y is acquire, it has to read

from a write that is in T_2 's X -knowledge hold for every space X : and (2) since the write of y is release, this write can only propagate to the X -knowledge of T_2 (for every X) after the po-earlier write to x has propagated there. This means that if the read of y retrieves 1, then T_2 already has the write to x in its X -knowledge, so when later reading x it cannot read from the overwritten initial value. Similarly, if $o_W = \text{prel}$, $o_R = \text{pacq}$, and $Y = X$, then, (i) and (ii) apply for $X = X$, and the same reasoning holds. Nevertheless, in every other case, T_2 can read the overwritten initial value: because either (1) the *read* from y is too weak and it allows to read from an event that has not yet propagated to the thread's X -knowledge; or (2) the *write* to y is too weak and it can propagate to the thread's X -knowledge before the po-earlier write to X has propagated.

Example 5.4. Being equivalent to dRC11, pRC11 provides “per-location-SC” (a.k.a. coherence). As a concrete example, $a = 1 \wedge b = 0$ in the following program on the right is disallowed.

To see this in pRC11, note that read events are also included in the threads' knowledge. Concretely, after executing the first read, the second thread has its own read in its X -knowledge. Then, reading later from the (implicit) initialization write is forbidden by pRC11 since that write and the first read will be have $G.\text{mo}$; $G.\text{rf}$ between them.

```
store(X, x, 1, rlx); || a := load(X, x, rlx);
                    || b := load(X, x, rlx);
```

Remark 3. It is also instructive to consider a simplified fragment with only one variable space and without na/*rlx* accesses. In this fragment, we can simply talk about the knowledge of each thread (instead of the X -knowledge for each X), and observe that: reads by thread τ can only read from writes that thread τ already knows about; and the propagation order respects **hb** (in particular, it follows the program order). This actually makes the reads deterministic: they have to read-from the **mo**-latest write among all known writes to the relevant variable. The resulting model is similar to message passing models for causal consistency [Beillahi et al. 2021].

Based on the pRC11 memory system, we define reachable program states and racy programs.

Definition 5.5. A program state $\bar{p}$ is *reachable* for a program Pr under pRC11 if $\langle \bar{p}, \langle G, K \rangle \rangle$ is reachable in $Pr \bowtie \text{pRC11}$ for some $\langle G, K \rangle \in \text{pRC11.Q}$. A program Pr is *racy* under pRC11 if $\langle \bar{p}, \langle G, K \rangle \rangle$ is reachable in $Pr \bowtie \text{pRC11}$ for some program state $\bar{p}$ and $\langle G, K \rangle \in \text{pRC11.Q}$ such that G is a racy execution graph (see Def. 4.12).

5.1 Equivalence of pRC11 and dRC11

We state our equivalence result between pRC11 and dRC11, relating Def. 5.5 to Definitions 4.11 and 4.13.

THEOREM 5.6 (EQUIVALENCE OF THE MODELS). *A program state $\bar{p}$ is reachable for a program Pr under dRC11 iff it is reachable for Pr under pRC11. Furthermore, Pr is racy under dRC11 iff it is racy under pRC11.*

Next, we describe the main steps in the proof (the full proof is given in §A.1). First, for the right-to-left directions, it suffices to establish the following invariants on reachable pRC11-states.

Definition 5.7. A knowledge mapping K is *well-formed* for an execution graph G if the following hold for every $\tau \in \text{Tid}$ and $X \in \text{Space}$:

- (1) $G.E^\tau \subseteq K(\tau)(X)$.
- (2) $(E_X \cup \text{CR}) \cap \text{dom}(G.\text{hb}_X ; [(W_X^{\text{prel}} \cup W^{\text{rel}}) \cap K(\tau)(X)]) \subseteq K(\tau)(X)$.
- (3) $(E_X \cup \text{CR}) \cap \text{dom}(G.\text{hb}_X ; [E^\tau]) \subseteq K(\tau)(X)$.

PROPOSITION 5.8. *If $\langle G, K \rangle$ is reachable in pRC11, then K is well-formed for G .*

LEMMA 5.9. *If $\langle G, K \rangle$ is reachable in pRC11, then G is dRC11-consistent.*

For the converse, one starts with an dRC11-consistent execution graph G , and has to traverse its events (following the program order, so it can be synchronized with the program), and intersperse propagation actions to make it a valid trace of pRC11. To construct this traversal, we define the following relations, where $P \triangleq \{p \in \text{PLab} \mid E(p) \in G.E\}$:

$$R_{\text{prop}} \triangleq \{\langle e, p \rangle \in G.E \times P \mid E(p) = e\}$$

$$T \triangleq \{\langle p, p' \rangle \in P \times P \mid \text{ptid}(p) = \text{ptid}(p') \wedge \text{psp}(p) = \text{psp}(p') \wedge \langle E(p), E(p') \rangle \in [E_{\text{psp}(p)} \cup \text{CR}] ; G.\text{hb}_{\text{psp}(p)} ; [W_{\text{psp}(p)}^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}]\}$$

$$R_{\text{rfp}} \triangleq \{\langle p, e \rangle \in P \times G.E \mid \langle E(p), e \rangle \in G.\text{rf}^+ \wedge \text{ptid}(p) = \text{tid}(e) \wedge e \in R_{\text{psp}(p)}^{\text{pacq}}\}$$

$$R_{\text{rf}} \triangleq \{\langle p, e \rangle \in P \times G.E \mid \langle E(p), e \rangle \in G.\text{rf}^+ \wedge \text{ptid}(p) = \text{tid}(e) \wedge e \in R^{\text{acq}}\}$$

$$R_{\text{fr}} \triangleq \{\langle r, p \rangle \in G.E \times P \mid \langle r, E(p) \rangle \in G.\text{fr} ; G.\text{rf}^? \wedge \text{ptid}(p) = \text{tid}(r) \wedge \text{psp}(p) = \text{sp}(r)\}$$

$$R_{\text{mo}} \triangleq \{\langle w, p \rangle \in G.E \times P \mid \langle w, E(p) \rangle \in G.\text{mo} ; G.\text{rf}^? \wedge \text{ptid}(p) = \text{tid}(w) \wedge \text{psp}(p) = \text{sp}(w)\}$$

$$R \triangleq G.\text{po} \cup G.\text{rf} \cup R_{\text{prop}} \cup T \cup R_{\text{rfp}} \cup R_{\text{rf}} \cup R_{\text{fr}} \cup R_{\text{mo}}$$

Then, the proof proceeds by showing that R is acyclic, and that every total order of $G.E \cup P$ extending R induces a trace of pRC11. Indeed, the relations above are in one-to-one correspondence with the conditions in the steps of pRC11.

6 Libraries and Their Clients

In this section we list the necessary definitions for the library abstraction theorem, and state the key properties that are used in its proof.

Client-library composition. A library L is a function mapping a set $\text{dom}(L) \subseteq \text{F}$ of method names to flat instruction sequences representing the method bodies. We only consider the case where libraries and their clients never access the same variable space. To formally define this syntactic restriction, we use the following notations for spaces used by libraries and their clients:

- $\text{Space}(I)$ denotes the set of variable spaces mentioned in an instruction sequence I .
- For a library L , $\text{Space}(L) \triangleq \bigcup_{f \in \text{dom}(L)} \text{Space}(L(f))$.
- For a program Pr and a set $F \subseteq \text{F}$, $\text{Space}(Pr \setminus F) \triangleq \bigcup_{\tau \in \text{Tid}, f \in (\text{FU} \setminus \{\text{main}\}) \setminus F} \text{Space}(Pr(\tau)(f))$.

Then, client-library composition is defined as follows.

Definition 6.1. A library L is *safe* for a program Pr if $\text{Space}(L) \cap \text{Space}(Pr \setminus \text{dom}(L)) = \emptyset$. When L is safe for Pr , we write $Pr[L]$ for the program obtained from Pr by setting $Pr(\tau)(f) = L(f)$ for every $\tau \in \text{Tid}$ and $f \in \text{dom}(L)$.

Next, we observe that the safety condition above ensures that execution graphs generated by $Pr[L]$ satisfy the following conditions relating the location spaces and the invoked methods of the graph events.

PROPOSITION 6.2. *Let L be a library that is safe for a program Pr . Suppose that $\langle \bar{p}, G \rangle$ is reachable in $Pr[L] \bowtie \text{FG}$. Then, the following hold for every $e \in G.R \cup G.W$:*

- If $\text{sp}(e) \in \text{Space}(L)$, then $\text{method}(e) \in \text{dom}(L)$.
- If $\text{sp}(e) \in \text{Space}(Pr \setminus \text{dom}(L))$, then $\text{method}(e) \notin \text{dom}(L)$.

Definition 6.3. A set $F \subseteq \text{F}$ is *encapsulated in an execution graph G* if for every $e_1, e_2 \in G.R \cup G.W$ with $\text{method}(e_1) \in F$ and $\text{method}(e_2) \notin F$, we have $\text{sp}(e_1) \neq \text{sp}(e_2)$.

From **Prop. 6.2**, we obtain the following.

PROPOSITION 6.4. *Let L be a library that is safe for a program Pr . If $\langle \bar{p}, G \rangle$ is reachable in $Pr[L] \bowtie FG$, then $\text{dom}(L)$ is encapsulated in G .*

The notions of well-formedness and encapsulated set of methods are lifted to memory states in the obvious way:

Definition 6.5. Let $M = \langle G, K \rangle$ be a memory state.

- $M = \langle G, K \rangle$ is *well-formed* if G is well-formed (Def. 4.7) and K is well-formed for G (Def. 5.7).
- F is *encapsulated in M* if it is encapsulated in G (Def. 6.3).

Since every reachable state in $Pr[L] \bowtie pRC11$ is also reachable in $Pr[L] \bowtie FG$, the following is an immediate consequence of Propositions 4.8, 5.8 and 6.4.

PROPOSITION 6.6. *Let L be a library that is safe for a program Pr . If $\langle \bar{p}, M \rangle$ is reachable in $Pr[L] \bowtie pRC11$, then M is well-formed and $\text{dom}(L)$ is encapsulated in M .*

Client-library program states. We define the composition of a program state $\bar{p}_{cl}$ representing a client state and a program state $\bar{p}_{lib}$ representing a library state as follows.

Definition 6.7. The *composition of two program states $\bar{p}_{cl}$ and $\bar{p}_{lib}$ w.r.t. a set $F \subseteq F$* , denoted by $\bar{p}_{cl}[F \mapsto \bar{p}_{lib}]$, is given by:

$$\bar{p}_{cl}[F \mapsto \bar{p}_{lib}] \triangleq \lambda \tau. \begin{cases} \langle \bar{p}_{lib}(\tau).pc, \bar{p}_{lib}(\tau).\phi, \bar{p}_{cl}(\tau).pc_s, \bar{p}_{cl}(\tau).f \rangle & \bar{p}_{cl}(\tau).f \in F \\ \bar{p}_{cl}(\tau) & \text{otherwise} \end{cases}$$

This definition uses $\bar{p}_{cl}$ for threads that are not currently inside a method in F , and $\bar{p}_{lib}$, but with the stored program counter and active method of $\bar{p}_{cl}$, for threads that are inside a method in F .

Histories. Histories record the interactions between libraries and clients. Formally, a *history* h of a library L is a sequence of transition labels representing a call to a method of L , a return from a method of L , or propagation of these call and return events. To define the history induced by a program, we employ the following notations (for every $F \subseteq F$):

$$\begin{aligned} \text{Call}_F &\triangleq \{\alpha \in \text{ProgLab} \mid \text{typ}(\alpha) = \text{CALL} \wedge \text{callee}(\alpha) \in F\} & \text{CP}_F &\triangleq \{p \in \text{PLab} \mid E(p) \in \text{CALL}_F\} \\ \text{Ret}_F &\triangleq \{\alpha \in \text{ProgLab} \mid \text{typ}(\alpha) = \text{RET} \wedge \text{method}(\alpha) \in F\} & \text{RP}_F &\triangleq \{p \in \text{PLab} \mid E(p) \in \text{RET}_F\} \\ \text{HLib}_F &\triangleq \text{Call}_F \cup \text{Ret}_F \cup \text{CP}_F \cup \text{RP}_F \end{aligned}$$

Definition 6.8. Let $F \subseteq F$. The *F-history* induced by a trace t of $Pr \bowtie pRC11$ for some program Pr , denoted by $H_F(t)$, is given by $H_F(t) \triangleq t|_{\text{HLib}_F}$. This notion is extended to sets of traces in the obvious way. The set of *F-histories of a program Pr* , denoted by $H_F(Pr)$, is given by $H_F(Pr) \triangleq H_F(\text{traces}(Pr \bowtie pRC11))$.

Client-library trace restrictions. We extract library and client transitions from a given trace as follows.

Definition 6.9. For $F \subseteq F$, the *F-restriction* and the $\bar{F}$ -restriction of a trace t of $Pr \bowtie pRC11$, denoted by $t|_F$ and $t|_{\bar{F}}$ (respectively), are given by:

$$t|_F \triangleq t|_{\{\alpha \in \text{ProgLab} \cup \text{PLab} \mid \text{method}(\alpha) \in F\} \cup \text{Call}_F \cup \text{CP}_F} \quad t|_{\bar{F}} \triangleq t|_{\{\alpha \in \text{ProgLab} \cup \text{PLab} \mid \text{method}(\alpha) \notin F\} \cup \text{Ret}_F \cup \text{RP}_F}$$

Note that both the *F-restriction* and the $\bar{F}$ -restriction of a trace t contain the *F-history* induced by t as a subsequence.

Restricting memory states. Similarly, we will need to restrict memory states to client/library, as defined next.

Definition 6.10. The *restriction of an execution graph* G w.r.t. a set $E \subseteq \mathcal{E}$, denoted by $G|_E$, is defined by: $G|_E.E \triangleq E \cup \text{Init}$ and $G|_E.X \triangleq [G|_E.E] ; G.X ; [G|_E.E]$ for every other component (i.e., $X \in \{\text{po}, \text{rf}, \text{mo}\}$).

Definition 6.11. The *restriction of a memory state* $M = \langle G, K \rangle$ w.r.t. a set $E \subseteq \mathcal{E}$, denoted by $M|_E$, is given by $M|_E \triangleq \langle G|_E, K|_E \rangle$, where $K|_E \triangleq \lambda \tau. \lambda X. K(\tau)(X) \cap E$.

Definition 6.12. Let $F \subseteq \mathcal{F}$. The F -events and the $\bar{F}$ -events denoted by E_F and $E_{\bar{F}}$ (respectively), are given by:

$$E_F \triangleq \{e \in \mathcal{E} \mid \text{method}(e) \in F\} \cup \text{CALL}_F \quad E_{\bar{F}} \triangleq \{e \in \mathcal{E} \mid \text{method}(e) \notin F\} \cup \text{RET}_F$$

The F -restriction and the $\bar{F}$ -restriction of a memory state $M = \langle G, K \rangle$, denoted by $M|_F$ and $M|_{\bar{F}}$ (respectively), are given by $M|_F \triangleq M|_{E_F}$ and $M|_{\bar{F}} \triangleq M|_{E_{\bar{F}}}$.

Again, we note that both the F -restriction and the $\bar{F}$ -restriction of a memory M contain the call events invoking methods in F and the return events that complete these invocations.

Restriction and merge properties. The following lemmas summarize the critical properties of the pRC11 memory system that allow us to compose memory traces of clients and libraries. In the proof of the abstraction theorem below we rely only on these properties of pRC11.

LEMMA 6.13 (RESTRICTION-1). Suppose that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ and let $F \subseteq \mathcal{F}$ that is encapsulated in M' . Then, the following hold:

- (1) If $\text{method}(\alpha) \in F$ or $\alpha \in \text{HLab}_F$, then $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'|_F$.
- (2) If $\text{method}(\alpha) \notin F$ or $\alpha \in \text{HLab}_F$, then $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M'|_{\bar{F}}$.

LEMMA 6.14 (RESTRICTION-2). Suppose that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ and let $F \subseteq \mathcal{F}$. Then, the following hold:

- (1) If $\text{method}(\alpha) \notin F$ and $\alpha \notin \text{HLab}_F$, then $M|_F = M'|_F$.
- (2) If $\text{method}(\alpha) \in F$ and $\alpha \notin \text{HLab}_F$, then $M|_{\bar{F}} = M'|_{\bar{F}}$.

LEMMA 6.15 (MERGE). Suppose that $F \subseteq \mathcal{F}$ is encapsulated in a well-formed memory state $M = \langle G, K \rangle$. Then, the following hold:

- (1) Let α be an pRC11 transition label with $\text{method}(\alpha) \in F$. Suppose that if $\alpha \in \text{ProgLab}$ and $\text{typ}(\alpha) \in \{\text{W}, \text{R}, \text{RMW}\}$, then $\text{method}(e) \in F$ for every $e \in G.E$ with $\text{sp}(e) = \text{sp}(\alpha)$. Then, $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'_F$ implies that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' such that $M'|_F = M'_F$.
- (2) Let α be an pRC11 transition label with $\text{method}(\alpha) \notin F$. Suppose that if $\alpha \in \text{ProgLab}$ and $\text{typ}(\alpha) \in \{\text{W}, \text{R}, \text{RMW}\}$, then $\text{method}(e) \notin F$ for every $e \in G.E$ with $\text{sp}(e) = \text{sp}(\alpha)$. Then, $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M'_{\bar{F}}$ implies that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' such that $M'|_{\bar{F}} = M'_{\bar{F}}$.

7 The Library Abstraction Theorem

In this section we state and prove the library abstraction theorem. This theorem assumes two libraries implementing the same set of methods, an *implementation* L and a *specification* $L^\#$ with $F = \text{dom}(L) = \text{dom}(L^\#)$, and two programs, a *client* Pr and a *most general client* MGC . The latter is representing the library's calling policy.

Example 7.1. For a library with no restrictions whatsoever on its clients (beyond the separation of variable spaces) one can use a most general client that repeatedly invokes arbitrary library methods with arbitrary stores. We denote this client by *MGC*_{free}. On the right we present the code of the main method in each thread τ in *MGC*_{free} for $\text{dom}(L) = \{f_1, \dots, f_n\}$. We use **havoc** for arbitrarily modifying all registers.

```
BEGIN : havoc; goto f1 | ... | fn;
f1 : call(f1); goto BEGIN;
...
fn : call(fn); goto BEGIN;
```

Example 7.2. Alternatively, a policy that requires to call the library methods in a race-free fashion is captured by a most general client that uses for each thread the code on the right. Here, we hold a lock L while executing every method. We assume a standard lock implementation using release/acquire accesses (see §8).

```
BEGIN : havoc; goto f1 | ... | fn;
f1 : acquire(L); call(f1);
      release(L); goto BEGIN;
...
fn : acquire(L); call(fn);
      release(L); goto BEGIN;
```

Beyond syntactic separation of variable spaces between libraries and their clients (see Def. 6.1), the library abstraction theorem has two conditions:

- L should *refine* $L^\#$ w.r.t. *MGC*, denoted by $L \sqsubseteq_{MGC} L^\#$, and formally defined by:

$$L \sqsubseteq_{MGC} L^\# \iff H_F(MGC[L]) \subseteq H_F(MGC[L^\#]) \wedge MGC[L] \text{ is not racy under pRC11}$$

- Pr should *adhere to MGC* w.r.t. $L^\#$, denoted by $Pr \sqsubseteq_{L^\#} MGC$, and formally defined by:

$$Pr \sqsubseteq_{L^\#} MGC \iff H_F(Pr[L^\#]) \subseteq H_F(MGC[L^\#]) \wedge Pr[L^\#] \text{ is not racy under pRC11}$$

The first condition is an obligation of the library developer. It requires that all histories generated for the most general client using the implementation are also generated with the specification. Additionally, the implementation should not have data races when run by the most general client. Importantly, both conditions do not mention Pr : library developers have to be able to verify their implementations “once and for all” without access to a particular client program.

Dually, the second condition is an obligation of the client. It requires that the client adheres to the library policy (otherwise, the blame is on the client), which means that all histories generated by the client program should be also generated by the most general client, which expresses that policy. Additionally, the client program should not have data races. Importantly, both conditions do not mention L : clients should be able to apply the abstraction theorem without access to the library implementation. In fact, it suffices to assume that Pr uses $L^\#$ both when checking for adherence to the library’s calling policy and for checking for data-race freedom (this can be important for both checks if the calling policy relies on the values returned by the library).

Now, what should the theorem guarantee? Intuitively, we want all client behaviors observable when using L to be observable when using $L^\#$. Thus, for every trace t generated by $Pr[L]$ that reaches a program state $\bar{p}$ and a memory state M , there should exist a corresponding trace $t^\#$ generated by $Pr[L^\#]$ that reaches a program state $\bar{p}^\#$ and a memory state $M^\#$, and a client should not be able to observe the difference between t and $t^\#$, $\bar{p}$ and $\bar{p}^\#$, and M and $M^\#$. This, however, does not mean that there is not any difference between these objects: L and $L^\#$ may perform different operations (leading to different traces), use different variables (leading to different memories), and internally use different registers (leading to different program states). We capture “client-equivalence” by requiring that: (1) the $\bar{F}$ -parts of t and $t^\#$ coincide (using Def. 6.9); (2) the state $\bar{p}^\#$ is obtained from $\bar{p}$ by only changing the instruction sequence state for threads that are currently inside a method in F (using Def. 6.7); and (3) the $\bar{F}$ -restriction of M and $M^\#$ coincide (using Def. 6.12). Finally, we also want to ensure that $Pr[L]$ is not racy.

All in all, we reach the following statement of the abstraction theorem.

THEOREM 7.3 (LIBRARY ABSTRACTION). *Let L and $L^\#$ be libraries implementing the same set F of methods. Let MGC and Pr be programs, such that both L and $L^\#$ are safe for both MGC and Pr . Suppose that $L \sqsubseteq_{MGC} L^\#$ and $Pr \sqsubseteq_{L^\#} MGC$. Then, the following hold:*

- *If $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}, M \rangle$, then there exist $t^\#$ and $\langle \bar{p}^\#, M^\# \rangle$ such that the following hold:*
 - (1) $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t^\#} \langle \bar{p}^\#, M^\# \rangle$; (2) $t^\#|_{\bar{F}} = t|_{\bar{F}}$; (3) $\bar{p}^\# = \bar{p}[F \mapsto \bar{p}_{\text{lib}}]$ for some $\bar{p}_{\text{lib}}$ (in particular, $\bar{p}^\#(\tau) = \bar{p}(\tau)$ whenever $\bar{p}(\tau).f \notin F$); and (4) $M^\#|_{\bar{F}} = M|_{\bar{F}}$.
- *$Pr[L]$ is not racy under pRC11 .*

In particular, the conditions of [Thm. 7.3](#) ensure that $L \sqsubseteq_{Pr} L^\#$. The following property, which allows compositional verification of a library consisting of several (non-interacting) libraries, is obtained as a corollary of the abstraction theorem as in [\[Khyzha and Lahav 2021\]](#).

COROLLARY 7.4 (COMPOSITIONALITY). *Let $L_1, \dots, L_n$ be libraries implementing pairwise disjoint sets of methods, such that $\text{Space}(L_1), \dots, \text{Space}(L_n), \text{Space}(L_1^\#), \dots, \text{Space}(L_n^\#)$, and $\text{Space}(MGC \setminus \text{dom}(L_1 \uplus \dots \uplus L_n))$ are pairwise disjoint. Suppose that for every $1 \leq i \leq n$, we have $L_i \sqsubseteq_{MGC_i} L_i^\#$ for $MGC_i = MGC[L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_n^\#]$. Then, $L_1 \uplus \dots \uplus L_n \sqsubseteq_{MGC} L_1^\# \uplus \dots \uplus L_n^\#$.*

The rest of this section is devoted to sketch the main steps in the proof of the abstraction theorem. First, we prove a general “composition lemma” which allows us to take a client’s portion from one trace and glue it together with a library’s portion from another trace, provided that the two traces induce the same history. The proof of this lemma is by induction on the length of the traces, where [Lemmas 6.13](#) to [6.15](#) provide the main tools for the induction step.

LEMMA 7.5 (COMPOSITION). *Let L and L' be libraries implementing the same set F of methods such that both are safe for a program Pr , and L is also safe for a program Pr' . Suppose that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t_{\text{cl}}} \langle \bar{p}_{\text{cl}}, M_{\text{cl}} \rangle$ and $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t_{\text{lib}}} \langle \bar{p}_{\text{lib}}, M_{\text{lib}} \rangle$, with $H_F(t_{\text{cl}}) = H_F(t_{\text{lib}})$. Then, $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$ for some trace t and memory state M such that $t|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$, $t|_F = t_{\text{lib}}|_F$, $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$, and $M|_F = M_{\text{lib}}|_F$.*

Now, using the composition lemma, we are able to show the following key property.

LEMMA 7.6. *Under the conditions of [Thm. 7.3](#), $H_F(Pr[L]) \subseteq H_F(MGC[L^\#])$.*

PROOF (OUTLINE). Assume otherwise, and let h be a shortest history in $H_F(Pr[L]) \setminus H_F(MGC[L^\#])$. Let t be a shortest trace of $Pr[L] \bowtie \text{pRC11}$ with $H_F(t) = h$. Since ε (the empty history) is clearly in $H_F(MGC[L^\#])$, we know that t is non-empty. Consider the last transition label α in t , and let t' such that $t = t' \cdot \alpha$. The minimality of t ensures that α must be an element of HLab_F (i.e., call, return, call propagation or return propagation of some method in F). Let $h' = H_F(t')$. The minimality of h further ensures that $h' \in H_F(MGC[L^\#])$. Let $t'_\#$ and $\langle \bar{p}'_\#, M'_\# \rangle$ such that $H_F(t'_\#) = h'$ and $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'_\#} \langle \bar{p}'_\#, M'_\# \rangle$. Let $\langle \bar{p}', M' \rangle$ be such that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'} \langle \bar{p}', M' \rangle$. We consider the following cases:

- (1) $\alpha \in \text{Call}_F \cup \text{CP}_F$: Using [Lemma 7.5](#) (applied with $L := L^\#$, $L' := L$, $Pr := Pr$, and $Pr' := MGC$), there exist $t''_\#$ and $M''_\#$ such that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t''_\#} \langle \bar{p}'[F \mapsto \bar{p}'_\#], M''_\# \rangle$, $H_F(t''_\#) = h'$, and $M''_\#|_{\bar{F}} = M'|_{\bar{F}}$. Then, using [Lemmas 6.13](#) and [6.15](#), it is straightforward to show that α is enabled in $\langle \bar{p}'[F \mapsto \bar{p}'_\#], M''_\# \rangle$, and it follows that $h = h' \cdot \alpha = H_F(t''_\#) \cdot \alpha \in H_F(Pr[L^\#])$. Since $Pr \sqsubseteq_{L^\#} MGC$, we obtain $h \in H_F(MGC[L^\#])$, which contradicts our assumption.

- (2) $\alpha \in \text{Ret}_F \cup \text{RP}_F$: Using [Lemma 7.5](#) (applied with $L := L$, $L' := L^\#$, $Pr := \text{MGC}$, and $Pr' := Pr$), there exist t'' and M'' such that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{\text{MGC}[L] \bowtie \text{pRC11}} \langle \bar{p}'[F \mapsto \bar{p}'], M'' \rangle$, $H_F(t'') = h'$, and $M''|_F = M'|_F$. Then, using [Lemmas 6.13](#) and [6.15](#), it is straightforward to show that α is enabled in $\langle \bar{p}'[F \mapsto \bar{p}'], M'' \rangle$, and it follows that $h = h' \cdot \alpha = H_F(t'') \cdot \alpha \in H_F(\text{MGC}[L])$. Since $L \sqsubseteq_{\text{MGC}} L^\#$, we obtain $h \in H_F(\text{MGC}[L^\#])$, which contradicts our assumption. $\square$

Then, to prove the abstraction theorem we need one additional lemma.

LEMMA 7.7. *Suppose that $M_{\text{Init}} \xrightarrow{t}_{\text{pRC11}} \langle G, K \rangle$. Then, there exist a trace t' and a knowledge mapping K' for G such that the following hold:*

- $M_{\text{Init}} \xrightarrow{t'}_{\text{pRC11}} \langle G, K' \rangle$.
- $t'|_{\text{ProgLab}} = t|_{\text{ProgLab}}$.
- $K'(\tau, X) \cap (W_X^{\text{prel}} \cup W_X^{\text{rel}} \cup \text{CR}) \subseteq \text{dom}(G.\text{hb}_X^2; [E^\tau])$ for every $\tau \in \text{Tid}$ and $X \in \text{Space}$.

The abstraction theorem is proved as follows.

PROOF OF THM. 7.3 (SKETCH). By [Lemma 7.6](#), we have $H_F(Pr[L]) \subseteq H_F(\text{MGC}[L^\#])$. Then, the first part of the claim directly follows using [Lemma 7.5](#) by letting $L := L^\#$, $L' := L$, $Pr := Pr$, and $Pr' := \text{MGC}$. For the second part, suppose that $Pr[L]$ is racy under pRC11. Let $\langle \bar{p}, \langle G, K \rangle \rangle$ be a reachable state in $Pr[L] \bowtie \text{pRC11}$ such that G is racy, and let e_1, e_2 be two events that form a race in G . Let $E = \text{dom}((G.\text{po} \cup G.\text{rf})^*; [e_1, e_2])$ and $G' = G|_E$. Clearly, e_1 and e_2 form a race in G' . It is also easy to see that there exist $\bar{p}'$ and K' such that $\langle \bar{p}', \langle G', K' \rangle \rangle$ is reachable in $Pr[L] \bowtie \text{pRC11}$. By [Lemma 7.7](#), there exists a knowledge mapping K'' for G' such that $\langle \bar{p}', \langle G', K'' \rangle \rangle$ is reachable in $Pr[L] \bowtie \text{pRC11}$ as well and $K''(\tau, X) \cap \text{CR} \subseteq \text{dom}(G'.\text{hb}_X^2; [E^\tau])$ for every $\tau \in \text{Tid}$ and $X \in \text{Space}$. By [Prop. 6.6](#), F is encapsulated in $\langle G', K'' \rangle$. Thus, since $\text{sp}(e_1) = \text{sp}(e_2)$, either (i) $e_1 \in E_{\bar{F}} \wedge e_2 \in E_{\bar{F}}$ or (ii) $e_1 \in E_F \wedge e_2 \in E_F$. If (i) holds, we obtain a contradiction by obtaining a reachable state in $Pr[L^\#] \bowtie \text{pRC11}$ in which e_1 and e_2 form a race. If (ii) holds, we obtain a contradiction by obtaining a reachable state in $\text{MGC}[L] \bowtie \text{pRC11}$ in which e_1 and e_2 form a race. $\square$

8 Illustrative Applications

In this section we present three examples of applications of library abstraction of different kinds. To simplify the presentation, we use an extended programming language with standard loops and conditionals, mechanisms for passing arguments in method calls/returns, arrays, etc. While these are not directly supported in our development above, the intention should be clear and the translation to programs as defined in §3 is straightforward (but tedious). Each of the libraries below has its own private variable space denoted by X_L . For simplicity, we use macros for lock operations:

```

acquire(x)  $\triangleq$  LOCK : 1 := CAS( $X_L, x, 0, 1$ , acq, rlx, rlx); if 1  $\neq$  0 goto LOCK
release(x)  $\triangleq$  store( $X_L, x, 0$ , rel)
pacquire(x)  $\triangleq$  LOCK : 1 := CAS( $X_L, x, 0, 1$ , pacq, rlx, rlx); if 1  $\neq$  0 goto LOCK
prelease(x)  $\triangleq$  store( $X_L, x, 0$ , prel)

```

8.1 Read-Copy-Update Synchronization

Read-copy-update (RCU) is a synchronization mechanism, heavily used in the Linux kernel, that allows a (single) writer to safely manipulate a data structure while multiple readers are concurrently accessing it [[Mckenney 2004](#)]. For that matter, it provides *RCU critical sections* for the readers (delimited by `rcu_read_lock` and `rcu_read_unlock`) and a synchronization method (`synchronize_rcu`) for the writer that waits for all readers currently in critical sections to exit

them. Readers should access the structure inside a critical section and the writer should not perform destructive updates before calling the synchronization method [Desnoyers et al. 2012].

Intuitively, every invocation of `synchronize_rcu` begins another ‘grace period’, and the implementation ensures that “read-side critical sections cannot span grace periods”.⁶ Following various long-lasting informal discussions among developers, Alglave et al. [2018] formalized this guarantee in the context of their proposed Linux Kernel memory model. They provided specialized ad hoc semantics to the RCU primitives in the form of declarative consistency constraints.

Next, we use our framework to demonstrate a simple specification of RCU under weak memory that is based on *standard locks*. We believe that our specification has the advantage of being more parsimonious and amenable to formal verification of client programs using RCU: it can be understood by relying solely on the semantics of locks, and allows verification using techniques and tools that already support reasoning about locks.

Our specification is given below. It assumes an MGC in which a particular thread (the writer) is repeatedly calling `synchronize_rcu`, and each other thread is a reader that interleaves invocations of `rcu_read_lock` and `rcu_read_unlock` with its own thread identifier, thus acquiring and releasing a per-reader lock $l[\tau]$. The set *Readers* consist of all thread identifiers except for the writer’s identifier, and the `foreach` loop iterates over this set in an arbitrary order (which may vary between invocations).

```
rcu_read_lock( $\tau$ ) :
  acquire( $l[\tau]$ );
  return();
```

```
rcu_read_unlock( $\tau$ ) :
  release( $l[\tau]$ );
  return();
```

```
synchronize_rcu :
  foreach  $\tau \in \text{Readers}$ 
    acquire( $l[\tau]$ );
    release( $l[\tau]$ );
  return();
```

Each RCU critical section is protected by a per-reader lock. For synchronization, the writer acquires and immediately releases each of the reader locks. This ensures that RCU critical sections do not span grace periods. For example, using the specification, it is easy to conclude that the following behavior of the client program is disallowed (as usual, 0 is the initial value of all variables):

<pre>store(X, x, 1, rlx); synchronize_rcu(); store(X, y, 1, rlx);</pre>		<pre>rcu_read_lock(T_2); a := load(X, x, rlx); // 0 b := load(X, y, rlx); // 1 rcu_read_unlock(T_2);</pre>
---	--	--

Indeed, roughly speaking, since the read of x returns 0, we know that the writer must have acquired the reader’s lock after the reader has released it. In turn, the value 1 has been written to y only after the reader exited its critical section.

Remark 4. Gotsman et al. [2013] introduced a verification technique *assuming SC* that is based on a simpler RCU specification. In their specification each reader thread τ sets a flag `racs`[τ] to 1 when entering the critical section, and resets it to 0 when exiting the critical section. The writer waits for all reader flags to be 0, which directly means that a critical section cannot span over more than one grace period. Using release/acquire accesses, their specification is as follows:

⁶See <https://lwn.net/Articles/573497/> [Accessed July 2022].

```
rcu_read_lock( $\tau$ ) :
  store( $X_L$ , rcs[ $\tau$ ], 1, rel);
  return();
```

```
rcu_read_unlock( $\tau$ ) :
  store( $X_L$ , rcs[ $\tau$ ], 0, rel);
  return();
```

```
synchronize_rcu :
  foreach  $\tau \in \text{Readers}$ 
    repeat
      a := load( $X_L$ , rcs[ $\tau$ ], acq);
      until a = 0;
  return();
```

Interestingly, this specification generates the same invocation-response sequences (i.e., same histories as employed in standard linearizability) as the one we propose. Nevertheless, it is *too weak* under weak memory concurrency, and, in fact, it allows the annotated behavior of the client above.

Following [Alglave et al. 2018; Desnoyers et al. 2012], we present a simplified implementation of the RCU library with weak memory constructs. It uses a global and per-reader ‘phase bits’, phase and rphase[τ], as well as per-reader flags, rcs[τ]. For correctness, it employs barriers (smp_mb in Linux or atomic_thread_fence(memory_order_seq_cst) in C11). While these are not included in the model we consider, they can be implemented here by acquire-release RMWs to an otherwise unused variable [Lahav et al. 2016]. Thus, fence() below is syntactic sugar for $r := \text{FADD}(X_L, f, 0, \text{acq}, \text{rel})$.

```
rcu_read_lock( $\tau$ ) :
  a := load( $X_L$ , phase, acq);
  store( $X_L$ , rphase[ $\tau$ ], a, rel);
  store( $X_L$ , rcs[ $\tau$ ], 1, rel);
  fence();
  return();
```

```
rcu_read_unlock( $\tau$ ) :
  store( $X_L$ , rcs[ $\tau$ ], 0, rel);
  return();
```

```
synchronize_rcu :
  fence();
  store( $X_L$ , phase, 1, rel);
  foreach  $\tau \in \text{Readers}$ 
    repeat
      a := load( $X_L$ , rcs[ $\tau$ ], acq);
      b := load( $X_L$ , rphase[ $\tau$ ], acq);
      until a = 0  $\vee$  b = 1;
  store( $X_L$ , phase, 0, rel);
  foreach  $\tau \in \text{Readers}$ 
    repeat
      a := load( $X_L$ , rcs[ $\tau$ ], acq);
      b := load( $X_L$ , rphase[ $\tau$ ], acq);
      until a = 0  $\vee$  b = 0;
  return();
```

In this implementation, for entering the critical section each reader τ stores phase in rphase[τ] and announces it enters the critical section by setting rcs[τ] to 1. When exiting the critical section, the reader resets rcs[τ] to 0. For the writer synchronization, the writer switches phase from 0 to 1, waits for each reader until it sees rcs[τ] = 0 or rphase[τ] = 1, switches phase back to 0, and waits again for each reader until it sees rcs[τ] = 0 or rphase[τ] = 0.

To verify refinement (i.e., to show $H_F(MGC[L]) \subseteq H_F(MGC[L^\#])$) we used the FDR refinement checker [Gibson-Robinson et al. 2014], which has been used before for automatic linearizability checking [Lowe 2017], and gives us guarantees up to a certain bound. Concretely, we managed to complete the automatic check for two readers with at most four unpropagated events between every two threads and for three readers with at most one unpropagated event. Nevertheless, we note that the length of the traces being checked is unbounded. A standard limitation here is the modeling gap between the paper definitions and FDR’s encoding, which uses communicating sequential processes (CSP). Using the fact that the library specification and implementation employ solely rel/acq accesses, which implies that propagation has to follow the program order and reads only read from propagated writes (see Remark 3), we can model pRC11’s memory state using FIFO buffers, which are supported by FDR. (But, using Thm. 7.3, refinement is guaranteed also when

non-`rel/acq` accesses are used by clients as in the example above.) We leave a systematic FDR encoding, and possibly the verification of more efficient variants of the above implementation that use `rlx` accesses to future work.

Finally, we note that while every history of the implementation is also a history of the specification, the implementation has important advantages over the specification. For instance, consider a scenario where a reader is repeatedly entering critical sections and the writer invokes `synchronize_rcu`. With the specification, to complete its invocation, the writer has to be “lucky” and catch a free lock between two reader’s sections. In turn, with the implementation, the writer is able to complete its invocation using the phase bit, without any particular “lucky” scheduling.

8.2 A Relaxed Concurrent Queue

Our second example continues the discussion in §2.2 and presents a specification and an implementation of a weak concurrent queue. The queue object is specified as follows:

```
enqueue(v) :
  pacquire(LT);
  tp := load(XL, T, rlx);
  store(XL, q[tp], v, rlx);
  store(XL, T, tp + 1, rlx);
  prelease(LT);
  return();
```

```
dequeue :
  pacquire(LH);
  hp := load(XL, H, rlx);
  hc := load(XL, q[hp], rlx);
  if hc ≠ ⊥ then
    store(XL, H, hp + 1, rlx);
  prelease(LH);
  return(hc);
```

The specification code uses an unbounded array `q` (represented by infinitely many variables `q[0]`, `q[1]`, ...) for the queue elements that initially contains only $\perp$. The variables `H` and `T` store the index of the current head and tail of the queue. Then, the specification wraps a sequential implementation inside locks, making three choices (aiming to demonstrate the framework’s flexibility):

- Except for the locks, all accesses are *relaxed*, so synchronization is only induced by the locks.
- The locks are accessed using *partial* acquire/release. This means that the queue does not expose its internal synchronization to the client, thus allowing all behaviors demonstrated in §2.2.
- The `enqueue` and `dequeue` methods use *two different locks*, which allows for certain “non-linearizable” behaviors. For example, after `enqueue` returns in one thread, a `dequeue` by another thread can retrieve $\perp$ (that signifies an empty queue). Indeed, a relaxed read from `q[k]` may return either the value that was written to `q[k]` or the initial value. To avoid these behaviors, clients, if they want, may form their own synchronization between enqueues and dequeues (and analyze possible behaviors of the queue using the above specification).

A possible implementation that avoids the locks and uses only relaxed accesses is shown next.

```
enqueue(v) :
  BEGIN : tp := load(XL, T, rlx);
  if tp ≠ 0 then
    tc := load(XL, q[tp - 1], rlx);
    if tc = ⊥ goto BEGIN;
  tp' := CAS(XL, T, tp, tp + 1, rlx, rlx, rlx);
  if tp' ≠ tp goto BEGIN;
  store(XL, q[tp'], v, rlx);
  return();
```

```
dequeue :
  BEGIN : hp := load(XL, H, rlx);
  hc := load(XL, q[hp], rlx);
  if hc = ⊥ then
    return(⊥);
  else
    hp' := CAS(XL, H, hp, hp + 1, rlx, rlx, rlx);
    if hp' ≠ hp goto BEGIN;
  return(hc);
```

This implementation uses CAS instructions to update `H` and `T`. It exploits the freedom allowed by the specification: no synchronization is ever induced between queue operations, and certain “non-linearizable” behaviors are allowed (when the read of `q[hp]` returns the overwritten $\perp$ value). We

note that it is important that `enqueue` waits for the previous cell being filled, which corresponds to all `enqueue`'s being totally ordered in the specification. In §A.5 we provide a sketch of the refinement proof, which generally follows a standard simulation argument albeit assuming non-SC memory and considering call/return propagation as observable transitions.

8.3 Local Data-Race-Freedom as an Instance of Library Abstraction

Our third application relies on the library policy component (*MGC*) of the abstraction theorem to derive a *local data-race-freedom guarantee* (LDRF, for short) for the underlying memory model (a fragment of RC11).⁷ Local data-race-freedom guarantees, introduced in [Dolan et al. 2018] and further developed in [Cho et al. 2021], generalize the more well-known (global) data-race-freedom guarantee, by ensuring strong semantics for locations accessed by non-racy executions. Crucially, the premise of these guarantees (i.e., what locations are racy) is checked assuming the strong semantics for the specified locations, which makes LDRF particularly useful for modular reasoning [Cho et al. 2021].

In our case, in the role of the strong semantics we have release/acquire (RA, for short) semantics (thus we establish what is called LDRF-RA in [Cho et al. 2021]), and we want to show that if a program avoids races on a set S of variables, then it is safe to assume that the accesses to S have RA semantics (as if they have `rel/acq` access modes). Moreover, it suffices to establish the premise, namely the avoidance of races on S , assuming the accesses to S have RA semantics.

To formulate this intricate property in terms of refinement, consider a library implementing methods `writex` and `readx` for every variable $x \in S$. The library specification ($L^\#$) is as follows:

```
writex(v) : store( $\chi_L$ , x, v, rel); return();
```

```
readx : a := load( $\chi_L$ , x, acq); return(a);
```

The implementation (L) employs non-atomic accesses (it could use any access mode):

```
writex(v) : store( $\chi_L$ , x, v, na); return();
```

```
readx : a := load( $\chi_L$ , x, na); return(a);
```

Then, the assurance provided by LDRF is precisely (conditioned) contextual refinement between L and $L^\#$. Now, to express the data-race-freedom premise we take *MGCdrf* to be a program that repeatedly and non-deterministically calls the library methods with arbitrary arguments, but avoids races by properly using standard per-variable readers-writer lock (which allows concurrent read operations but write operations require exclusive access). That is, before calling `writex` (respectively, `readx`) *MGCdrf* takes a lock on x with a write-mode (respectively, read-mode).

Adherence to the calling policy as specified by *MGCdrf* (i.e., $Pr \sqsubseteq_{L^\#} \text{MGCdrf}$) is equivalent to the LDRF premise. In particular, $H_{dom(L)}(Pr[L^\#]) \subseteq H_{dom(L)}(\text{MGCdrf}[L^\#])$ means that in histories in $H_{dom(L)}(Pr[L^\#])$ for every variable $x \in S$: (1) before calling `writex` by thread τ , the return markers of all previous invocations of `writex` and `readx` have propagated to thread τ ; and (2) before calling `readx` by thread τ , the return markers of all previous invocations of `writex` have propagated to thread τ . In addition, adherence to the policy in our abstraction theorem assumes $L^\#$ (rather than L) just like LDRF's premise assuming RA semantics for accesses to S .

Accordingly, by applying the abstraction theorem we can establish LDRF-RA by proving our library correctness criterion $L \sqsubseteq_{\text{MGCdrf}} L^\#$ for the libraries above. This is straightforward: the conditions on propagation events of calls and returns that arise by $L^\#$ are already covered by *MGCdrf*. For instance, if an invocation of `readx` reads from an invocation of `writex` in an execution of *MGCdrf*[$L^\#$], then due to the `rel/acq` accesses by $L^\#$, the call `writex` has to propagate before `readx` returns, but this already holds in traces of *MGCdrf*[$L^\#$] due to *MGCdrf*.

⁷To the best of our knowledge, we are the first to observe that (local) data-race-freedom guarantee is an instance of library abstraction.

9 Related Work

We have adopted several key instruments proposed in earlier work to achieve library abstraction. Concretely, [Burckhardt et al. 2012] addressed this challenge for the x86-TSO architecture [Owens et al. 2009], and achieved library abstraction by including call/return markers in the thread-local store buffers, and exposing the propagation of these markers to the main memory in library histories. We adapt this idea to allow point-to-point communication and apply it in a much weaker memory model. (Moreover, [Burckhardt et al. 2012] does not consider libraries that rely on synchronization by the client or not meant to expose their internal synchronization to the client, since these issues do not arise on x86-TSO.) In turn, from [Khyzha and Lahav 2022], which studied library abstraction under non-volatile memory, we adopt the general formalism, the proof strategy of the abstraction theorem by relying on a general composition lemma (Lemma 7.5), and the modeling of a library’s calling policy as a client program (MGC).

Another closely related work is [Batty et al. 2013], which provides the first library abstraction result for C11. For having simple specifications (which they want, like us, to be pieces of code) Batty et al. [2013] extended (a fragment of) C11 with specialized “atomic block” constructs, whose semantics is similar to the semantics of software transactions. In contrast, we use locks with “partial release/acquire” semantics, and have no need in including transactional features, which may be harder for client reasoning. Other crucial differences is that [Batty et al. 2013] does not support library’s calling policies, and their correctness condition is based on partially ordered execution traces (with ‘guarantee’ and ‘deny’ relations), while we opt for considering only totally ordered histories, which, we believe, are easier to grasp. We also note that the relaxed atomics employed in [Batty et al. 2013] have the original C11 semantics (*without* the **PO-RR** constraint, and thus *with* “out-of-thin-air” behaviors), which, as they show, does not allow fully compositional reasoning.

Library abstraction under weak memory was also studied from a *declarative* point of view in [Dongol et al. 2018; Raad et al. 2019], using partial orders for exposing the externally visible synchronization induced by concurrent objects. Concretely, in the framework of [Raad et al. 2019], libraries are specified as collections of execution graphs, and the abstraction condition relates the graphs produced by the implementation to those in the specification. Their specification framework does not depend on the language constructs, and it is sufficiently expressive for having direct specifications of non-standard non-linearizable objects (e.g., queues with certain non-FIFO behaviors). We believe that there is a price for the generality: rich declarative specifications are often hard to understand and informally apply, and require non-standard methods for client reasoning. In addition, [Raad et al. 2019]’s ‘well-formedness’ requirement (which corresponds to our “policy adherence”) is defined in terms of the implementation library (unlike in our abstraction theorem).

Smith et al. [2020] studied linearizability on a general hardware memory model, and related it to a certain refinement notion, which they call “object refinement”. In contrast, we are aiming for standard contextual refinement between an implementation and its specification code. Other works developed correctness conditions for concurrent objects under weakly consistent memory, but were not formally related to a refinement notion (e.g., [Doherty et al. 2018]).

Weakly consistent objects were also studied by Emmi and Enea [2019] who proposed a specification framework based on extending standard histories with a *visibility* relation that is subject to varying constraints. Their approach was applied to non-linearizable Java concurrent objects, and allows verification by a generalization of the concept of linearization points [Krishna et al. 2020].

Other recent works propose *logical* approaches to library specification and verification under weak memory [Dalvandi and Dongol 2021a,b; Dang et al. 2022; Mével and Jourdan 2021]. In contrast to our work, their specifications are given as Hoare triples and refinement is understood from a program logic perspective. Specifically, for a fragment of the model we study (without non-atomics),

Dalvandi and Dongol [2021a,b] propose an Owicki-Gries-style Hoare logic and specify libraries with view-based object semantics. Their approach is, however, limited to clients who synchronize only through the library operations. In turn, Mével and Jourdan [2021] and Dang et al. [2022] are based on the idea of *logical atomicity* [Birkedal et al. 2021; Jung 2019], and provide mechanized verification of certain library implementations in the Iris framework [Jung et al. 2015]. The weak memory model in [Dang et al. 2022] is similar to the one we study (RC11), while the model in [Mével and Jourdan 2021] is Multicore OCaml, which provides much stronger accesses than RC11. The library studied in [Mével and Jourdan 2021] is a queue that is stronger than our example in §8.2. It exposes synchronization to the client between from an enqueue of some element to the dequeue of that specific element, but still, unlike a lock-based queue, does not expose synchronization in all other cases (e.g., between two different enqueues).

Finally, the idea of having explicit steps for propagating knowledge between threads appeared before in semantics of relaxed memory concurrency, particularly in an operational model for the POWER architecture [Sarkar et al. 2011], but we are not aware of a similar existing model for RC11.

10 Conclusion

We established a library abstraction for (a fragment of) the RC11 model extended with specialized partial release/acquire accesses. The latter are used to have simple lock-based specifications of libraries that do not expose internal synchronization to their clients or libraries that rely on synchronization by the client. The condition for library abstraction is based on inclusion of sets of totally ordered library histories. To achieve this, we developed an operational version of the model with explicit steps for propagation of knowledge between threads, and included the propagation of method invocation and responses in library histories.

Our work has several limitations, which are interesting to lift and address in future work. In particular: (1) The fragment of RC11 that we consider excludes fences and SC-atomics; (2) We only consider partial correctness and finite histories ignoring liveness issues and thus our refinement notion is not termination preserving (see [Gotsman and Yang 2011] for a possible approach to lift this limitation); (3) We disallow nested method calls (in particular, recursion), as common in works relating linearizability to refinement (also under SC); and (4) We disallow libraries and their clients to transfer ownership of data structures among themselves or to run in a shared address space (see [Gotsman and Yang 2013] for a possible approach to lift this limitation).

Verification techniques for both library development and client reasoning are beyond the scope of the current work. In particular, it is interesting to see how our specification constructs can be supported by a program logic, which will pave the way to foundational mechanized refinement proofs as in [Dang et al. 2022], as well as by model checking tools, such as GenMC [Kokologiannakis et al. 2019] and C11Tester [Luo and Demsky 2021], that can be used for client reasoning.

Finally, we are interested in developing a view-based operational model (see, e.g., [Kaiser et al. 2017; Kang et al. 2017]) with explicit propagation steps that will be equivalent to pRC11, but possibly more intuitive. We also believe that pRC11 (in particular, its *rel/acq* fragment, see Remark 3) may be of independent interest for developing verification techniques, by relying on the analogy between pRC11's knowledge propagation and well-studied x86-TSO store propagation steps (see e.g., [Bouajjani et al. 2018; Enea and Farzan 2016]).

Acknowledgments

We thank the anonymous reviewers for their valuable feedback. This work was supported by the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation programme (grant agreement no. 851811) and the Israel Science Foundation (grant numbers 1566/18 and 814/22).

References

- Jade Alglave, Luc Maranget, Paul E. McKenney, Andrea Parri, and Alan Stern. 2018. Frightening Small Children and Disconcerting Grown-ups: Concurrency in the Linux Kernel. In *ASPLOS*. ACM, New York, NY, USA, 405–418. doi:10.1145/3173162.3177156
- Jade Alglave, Luc Maranget, and Michael Tautschnig. 2014. Herding Cats: Modelling, Simulation, Testing, and Data Mining for Weak Memory. *ACM Trans. Program. Lang. Syst.* 36, 2, Article 7 (July 2014), 74 pages. doi:10.1145/2627752
- Mark Batty, Mike Dodds, and Alexey Gotsman. 2013. Library Abstraction for C/C++ Concurrency. In *POPL*. ACM, New York, NY, USA, 235–248. doi:10.1145/2429069.2429099
- Mark Batty, Kayvan Memarian, Kyndylan Nienhuis, Jean Pichon-Pharabod, and Peter Sewell. 2015. The Problem of Programming Language Concurrency Semantics. In *ESOP*. Springer, Berlin, Heidelberg, 283–307. http://dx.doi.org/10.1007/978-3-662-46669-8_12
- Sidi Mohamed Beillahi, Ahmed Bouajjani, and Constantin Enea. 2021. Robustness Against Transactional Causal Consistency. *Logical Methods in Computer Science* Volume 17, Issue 1 (Feb. 2021). doi:10.23638/LMCS-17(1:12)2021
- Lars Birkedal, Thomas Dinsdale-Young, Armaël Guéneau, Guilhem Jaber, Kasper Svendsen, and Nikos Tzevelekos. 2021. Theorems for Free from Separation Logic Specifications. *Proc. ACM Program. Lang.* 5, ICFP, Article 81 (Aug. 2021), 29 pages. doi:10.1145/3473586
- Hans-Juergen Boehm and Brian Demsky. 2014. Outlawing Ghosts: Avoiding Out-of-thin-air Results. In *MSPC*. ACM, New York, NY, USA, Article 7, 6 pages. doi:10.1145/2618128.2618134
- Ahmed Bouajjani, Michael Emmi, Constantin Enea, and Jad Hamza. 2015. Tractable Refinement Checking for Concurrent Objects. In *POPL*. ACM, New York, NY, USA, 651–662. <https://doi.org/10.1145/2676726.2677002>
- Ahmed Bouajjani, Constantin Enea, Suha Orhun Mutluergil, and Serdar Tasiran. 2018. Reasoning About TSO Programs Using Reduction and Abstraction. In *CAV*. Springer International Publishing, Cham, 336–353. doi:10.1007/978-3-319-96142-2_21
- Sebastian Burckhardt, Alexey Gotsman, Madanlal Musuvathi, and Hongseok Yang. 2012. Concurrent Library Correctness on the TSO Memory Model. In *ESOP*. Springer, Berlin, Heidelberg, 87–107. doi:10.1007/978-3-642-28869-2_5
- Minki Cho, Sung-Hwan Lee, Chung-Kil Hur, and Ori Lahav. 2021. Modular Data-Race-Freedom Guarantees in the Promising Semantics. In *PLDI*. ACM, New York, NY, USA, 867–882. <https://doi.org/10.1145/3453483.3454082>
- Sadeh Dalvandi and Brijesh Dongol. 2021a. Verifying C11-Style Weak Memory Libraries. In *PPoPP*. ACM, New York, NY, USA, 451–453. doi:10.1145/3437801.3441619
- Sadeh Dalvandi and Brijesh Dongol. 2021b. Verifying C11-Style Weak Memory Libraries via Refinement. *CoRR* abs/2108.06944 (2021). arXiv:2108.06944 <https://arxiv.org/abs/2108.06944>
- Hoang-Hai Dang, Jaehwang Jung, Jaemin Choi, Duc-Thuan Nguyen, William Mansky, Jeehoon Kang, and Derek Dreyer. 2022. Compass: Strong and Compositional Library Specifications in Relaxed Memory Separation Logic. In *PLDI*. ACM, New York, NY, USA, 792–808. doi:10.1145/3519939.3523451
- Mathieu Desnoyers, Paul E. McKenney, Alan S. Stern, Michel R. Dagenais, and Jonathan Walpole. 2012. User-Level Implementations of Read-Copy Update. *IEEE Trans. Parallel Distrib. Syst.* 23, 2 (2012), 375–382. doi:10.1109/TPDS.2011.159
- Simon Doherty, Brijesh Dongol, Heike Wehrheim, and John Derrick. 2018. Making Linearizability Compositional for Partially Ordered Executions. In *iFM*. Springer International Publishing, Cham, 110–129. doi:10.1007/978-3-319-98938-9_7
- Stephen Dolan, KC Sivaramakrishnan, and Anil Madhavapeddy. 2018. Bounding Data Races in Space and Time. In *PLDI*. ACM, New York, NY, USA, 242–255. doi:10.1145/3192366.3192421
- Brijesh Dongol, Radha Jagadeesan, James Riely, and Alasdair Armstrong. 2018. On Abstraction and Compositionality for Weak-Memory Linearisability. In *VMCAI*. Springer International Publishing, Cham, 183–204. doi:10.1007/978-3-319-73721-8_9
- Michael Emmi and Constantin Enea. 2019. Weak-Consistency Specification via Visibility Relaxation. *Proc. ACM Program. Lang.* 3, POPL, Article 60 (Jan. 2019), 28 pages. doi:10.1145/3290373
- Constantin Enea and Azadeh Farzan. 2016. On Atomicity in Presence of Non-atomic Writes. In *TACAS*. Springer, Berlin, Heidelberg, 497–514. doi:10.1007/978-3-662-49674-9_29
- Ivana Filipović, Peter O’Hearn, Noam Rinetzy, and Hongseok Yang. 2010. Abstraction for concurrent objects. *Theoretical Computer Science* 411, 51 (2010), 4379–4398. <https://www.sciencedirect.com/science/article/pii/S0304397510005001>
- Thomas Gibson-Robinson, Philip Armstrong, Alexandre Boulgakov, and Andrew W. Roscoe. 2014. FDR3 — A Modern Refinement Checker for CSP. In *TACAS*. Springer, Berlin, Heidelberg, 187–201. doi:10.1007/978-3-642-54862-8_13
- Alexey Gotsman, Noam Rinetzy, and Hongseok Yang. 2013. Verifying Concurrent Memory Reclamation Algorithms with Grace. In *ESOP*. Springer, Berlin, Heidelberg, 249–269. doi:10.1007/978-3-642-37036-6_15
- Alexey Gotsman and Hongseok Yang. 2011. Liveness-Preserving Atomicity Abstraction. In *ICALP*. Springer, Berlin, Heidelberg, 453–465. doi:10.1007/978-3-642-22012-8_36
- Alexey Gotsman and Hongseok Yang. 2013. Linearizability with Ownership Transfer. *Logical Methods in Computer Science* Volume 9, Issue 3 (Sept. 2013). <https://lmcs.episciences.org/931>

- Maurice P. Herlihy and Jeannette M. Wing. 1990. Linearizability: A Correctness Condition for Concurrent Objects. *ACM Trans. Program. Lang. Syst.* 12, 3 (July 1990), 463–492. doi:10.1145/78969.78972
- Ralf Jung. 2019. Logical atomicity in Iris: The good, the bad, and the ugly. In *Iris Workshop*. <https://people.mpi-sws.org/~jung/iris/talk-iris2019.pdf>
- 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*. ACM, New York, NY, USA, 637–650. doi:10.1145/2676726.2676980
- Jan-Oliver Kaiser, Hoang-Hai Dang, Derek Dreyer, Ori Lahav, and Viktor Vafeiadis. 2017. Strong Logic for Weak Memory: Reasoning About Release-Acquire Consistency in Iris. In *ECOOP*. Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik, Dagstuhl, Germany, 17:1–17:29. doi:10.4230/LIPICs.ECOOP.2017.17
- Jeehoon Kang, Chung-Kil Hur, Ori Lahav, Viktor Vafeiadis, and Derek Dreyer. 2017. A Promising Semantics for Relaxed-Memory Concurrency. In *POPL*. ACM, New York, NY, USA, 175–189. doi:10.1145/3009837.3009850
- Artem Khyzha and Ori Lahav. 2021. Taming x86-TSO Persistency. *Proc. ACM Program. Lang.* 5, POPL, Article 47 (Jan. 2021), 29 pages. <https://doi.org/10.1145/3434328>
- Artem Khyzha and Ori Lahav. 2022. Abstraction for Crash-Resilient Objects. In *ESOP*. Springer International Publishing, Cham, 262–289. doi:10.1007/978-3-030-99336-8_10
- Michalis Kokologiannakis, Ori Lahav, Konstantinos Sagonas, and Viktor Vafeiadis. 2017. Effective Stateless Model Checking for C/C++ Concurrency. *Proc. ACM Program. Lang.* 2, POPL, Article 17 (Dec. 2017), 32 pages. doi:10.1145/3158105
- Michalis Kokologiannakis, Azalea Raad, and Viktor Vafeiadis. 2019. Model Checking for Weakly Consistent Libraries. In *PLDI*. ACM, New York, NY, USA, 96–110. doi:10.1145/3314221.3314609
- Siddharth Krishna, Michael Emmi, Constantin Enea, and Dejan Jovanović. 2020. Verifying Visibility-Based Weak Consistency. In *ESOP*. Springer International Publishing, Cham, 280–307. doi:10.1007/978-3-030-44914-8_11
- Ori Lahav, Nick Giannarakis, and Viktor Vafeiadis. 2016. Taming Release-Acquire Consistency. In *POPL*. ACM, New York, NY, USA, 649–662. doi:10.1145/2837614.2837643
- Ori Lahav and Roy Margalit. 2019. Robustness Against Release/Acquire Semantics. In *PLDI*. ACM, New York, NY, USA, 126–141. doi:10.1145/3314221.3314604
- Ori Lahav, Viktor Vafeiadis, Jeehoon Kang, Chung-Kil Hur, and Derek Dreyer. 2017. Repairing Sequential Consistency in C/C++11. In *PLDI*. ACM, New York, NY, USA, 618–632. doi:10.1145/3062341.3062352
- Leslie Lamport. 1979. How to Make a Multiprocessor Computer That Correctly Executes Multiprocess Programs. *IEEE Trans. Computers* 28, 9 (Sept. 1979), 690–691. doi:10.1109/TC.1979.1675439
- Gavin Lowe. 2017. Analysing Lock-Free Linearizable Datatypes Using CSP. In *Concurrency, Security, and Puzzles - Essays Dedicated to Andrew William Roscoe on the Occasion of His 60th Birthday*. Springer, 162–184. doi:10.1007/978-3-319-51046-0_9
- Weiyu Luo and Brian Demsky. 2021. C11Tester: A Race Detector for C/C++ Atomics. In *ASPLOS*. ACM, New York, NY, USA, 630–646. doi:10.1145/3445814.3446711
- Paul E. McKenney. 2004. *Exploiting Deferred Destruction: An Analysis of Read-Copy-Update Techniques in Operating System Kernels*. Ph. D. Dissertation. Oregon Health & Science University.
- Glen Mével and Jacques-Henri Jourdan. 2021. Formal Verification of a Concurrent Bounded Queue in a Weak Memory Model. *Proc. ACM Program. Lang.* 5, ICFP, Article 66 (Aug. 2021), 29 pages. doi:10.1145/3473571
- Peizhao Ou and Brian Demsky. 2018. Towards Understanding the Costs of Avoiding Out-of-Thin-Air Results. *Proc. ACM Program. Lang.* 2, OOPSLA, Article 136 (Oct. 2018), 29 pages. doi:10.1145/3276506
- Scott Owens, Susmit Sarkar, and Peter Sewell. 2009. A Better x86 Memory Model: x86-TSO. In *TPHOLs*. Springer, Berlin, Heidelberg, 391–407. doi:10.1007/978-3-642-03359-9_27
- Azalea Raad, Marko Doko, Lovro Rožić, Ori Lahav, and Viktor Vafeiadis. 2019. On Library Correctness under Weak Memory Consistency: Specifying and Verifying Concurrent Libraries under Declarative Consistency Models. *Proc. ACM Program. Lang.* 3, POPL, Article 68 (Jan. 2019), 31 pages. doi:10.1145/3290381
- Susmit Sarkar, Peter Sewell, Jade Alglave, Luc Maranget, and Derek Williams. 2011. Understanding POWER Multiprocessors. In *PLDI*. ACM, New York, NY, USA, 175–186. doi:10.1145/1993498.1993520
- Graeme Smith, Kirsten Winter, and Robert J. Colvin. 2020. Linearizability on Hardware Weak Memory Models. *Form. Asp. Comput.* 32, 1 (Feb. 2020), 1–32. doi:10.1007/s00165-019-00499-8

A Proofs

In this appendix we provide proof outlines that were not included in the main text. We skip proofs of propositions that straightforwardly follow from our definitions.

A.1 Proofs for Section 5

LEMMA 5.9. *If $\langle G, K \rangle$ is reachable in pRC11, then G is dRC11-consistent.*

PROOF. The proof is by induction on the length of the trace reaching $\langle G, K \rangle$. For the empty trace, the claim is obvious. Consider an pRC11-step transitioning from a reachable state $\langle G, K \rangle$ to $\langle G', K' \rangle$. By the induction hypothesis, G is dRC11-consistent. For the CALL/RETURN and PROPAGATE steps of pRC11, the dRC11-consistency of G' easily follows from the fact that G is dRC11-consistent. Consider a WRITE/READ/RMW step: $\langle G, K \rangle \xrightarrow{\tau, l, f}_{\text{pRC11}} \langle G', K' \rangle$ with $\text{typ}(l) \in \{W, R, \text{RMW}\}$. Let $e = \text{NextEvent}(G.E, \tau, l, f)$, and let $w \in G.W_{\text{sp}(e), \text{loc}(e)}$ such that the following hold:

- $G' = \text{Add}(G, e, w)$.
- $w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(\text{sp}(e))])$.
- $e \in W \implies w \notin \text{dom}(G.\text{rf}; [\text{RMW}])$.
- $K' = K' = K[\tau \mapsto k']$, where $k' = \lambda Y. K(\tau)(Y) \cup \{e\}$.

By Prop. 5.8, we also know that K' is a well-formed knowledge mapping for G' . We prove each of the consistency conditions for G' :

- **READ COHERENCE:** Let $X \in \text{Space}$. We show that $[R_X]; G'.\text{fr}; G'.\text{rf}^?; G'.\text{hb}_X$ is irreflexive. Since e is $G'.\text{po} \cup G'.\text{rf}$ -maximal, it suffices to show that $\langle e, e \rangle \notin [R_X]; G'.\text{fr}; G'.\text{rf}^?; G'.\text{hb}_X$. Suppose otherwise, and let $e' \in G'.E$ such that $\langle e, e' \rangle \in [R_X]; G'.\text{fr}; G'.\text{rf}^?$ and $\langle e', e \rangle \in G'.\text{hb}_X$. In this case, we have $e \in R_X$. Then, the fact that $\langle e, e' \rangle \in G'.\text{fr}; G'.\text{rf}^?$ implies that $\langle w, e' \rangle \in G.\text{mo}; G.\text{rf}^?$. In addition, since $\langle e', e \rangle \in G'.\text{hb}_X$, the fact that K' is a well-formed knowledge mapping for G' implies that $e' \in K'(\tau)(X)$, which entails that $e' \in K(\tau)(X)$ (since $e \neq e'$). This contradicts the fact that $w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(X)])$.
- **WRITE COHERENCE:** Let $X \in \text{Space}$. We show that $[W_X]; G'.\text{mo}; G'.\text{rf}^?; G'.\text{hb}_X$ is irreflexive. Since e is $G'.\text{po} \cup G'.\text{rf}$ -maximal, it suffices to show that $\langle e, e \rangle \notin [W_X]; G'.\text{mo}; G'.\text{rf}^?; G'.\text{hb}_X$. Suppose otherwise, and let $e' \in G'.E$ such that $\langle e, e' \rangle \in [W_X]; G'.\text{mo}; G'.\text{rf}^?$ and $\langle e', e \rangle \in G'.\text{hb}_X$. In this case, we have $e \in W_X$. Then, the fact that $\langle e, e' \rangle \in G'.\text{mo}; G'.\text{rf}^?$ implies that $\langle w, e' \rangle \in G.\text{mo}; G.\text{rf}^?$. In addition, since $\langle e', e \rangle \in G'.\text{hb}_X$, the fact that K' is a well-formed knowledge mapping for G' implies that $e' \in K'(\tau)(X)$, which entails that $e' \in K(\tau)(X)$ (since $e \neq e'$). This contradicts the fact that $w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(X)])$.
- **RMW1:** We show that $G'.\text{mo}; G'.\text{rf}$ is irreflexive. Suppose otherwise, and let $e_1, e_2 \in G'.E$ such that $\langle e_1, e_2 \rangle \in G'.\text{mo}$ and $\langle e_2, e_1 \rangle \in G'.\text{rf}$. This implies that $e_1 \in \text{RMW}$. Since e has no outgoing **rf**-edges in G' , we have $e_2 \neq e$. Now, if $e_1 \neq e$, then we have $\langle e_1, e_2 \rangle \in G.\text{mo}$ and $\langle e_2, e_1 \rangle \in G.\text{rf}$, which contradicts the fact that G satisfies **RMW1**. On the other hand, we cannot have $e_1 = e$ as well, since when adding the **rf**-edge to an RMW event, we make its source to be the **mo**-immediate predecessor of the added event (see definition of $\text{Add}(G, e, w)$).
- **RMW2:** We show that $G'.\text{fr}; G'.\text{mo}$ is irreflexive. Suppose otherwise, and let $e_1, e_2 \in G'.E$ such that $\langle e_1, e_2 \rangle \in G'.\text{fr}$ and $\langle e_2, e_1 \rangle \in G'.\text{mo}$. This implies that $e_1 \in \text{RMW}$. Let e_3 such that $\langle e_3, e_1 \rangle \in G'.\text{rf}$ and $\langle e_3, e_2 \rangle \in G'.\text{mo}$. Since e has no outgoing **rf**-edges in G' , we have $e_3 \neq e$. In addition, we have $e_1 \neq e$ as well, since when adding the **rf**-edge to an RMW event, we make its source to be the **mo**-immediate predecessor of the added event. Now, if $e_2 \neq e$, we obtain $\langle e_1, e_2 \rangle \in G.\text{fr}$ and $\langle e_2, e_1 \rangle \in G.\text{mo}$, which contradicts the fact that G satisfies **RMW2**. Otherwise, we have $e_2 = e$. Then, $\langle e_2, e_1 \rangle \in G'.\text{mo}$ implies that $\langle w, e_1 \rangle \in G.\text{mo}$, and $\langle e_3, e \rangle \in G'.\text{mo}$ implies that $\langle e_3, w \rangle \in G.\text{mo}^?$. Now, if $e_3 \neq w$, we obtain $\langle e_1, w \rangle \in G.\text{fr}$ and $\langle w, e_1 \rangle \in G.\text{mo}$, which again contradicts the fact that G

satisfies **RMW2**. On the other hand, we cannot have $e_3 = w$, since this contradicts the fact that $w \notin \text{dom}(G.\text{rf}; [\text{RMW}])$ when $e \in W$.

- **PO-RF**: Easily follows by the facts that G satisfies **PO-RF** and that e is $G'.\text{po} \cup G'.\text{rf}$ -maximal. $\square$

THEOREM 5.6 (EQUIVALENCE OF THE MODELS). *A program state $\bar{p}$ is reachable for a program Pr under dRC11 iff it is reachable for Pr under pRC11. Furthermore, Pr is racy under dRC11 iff it is racy under pRC11.*

PROOF. It suffices to show that for every program state $\bar{p}$ and execution graph G the following are equivalent:

- (a) $\langle \bar{p}, G \rangle$ is reachable in $Pr \bowtie \text{FG}$ and G is dRC11-consistent.
- (b) $\langle \bar{p}, \langle G, K \rangle \rangle$ is reachable in $Pr \bowtie \text{pRC11}$ for some knowledge mapping K .

Given **Lemma 5.9**, (b) $\implies$ (a) is easy: Since for every $\alpha \in \text{ProgLab}$, $\langle G, K \rangle \xrightarrow{\alpha}_{\text{pRC11}} \langle G', K' \rangle$ implies $G \xrightarrow{\alpha}_{\text{FG}} G'$, it suffices to consider an pRC11 trace and show that the execution graph component in the final state of that run is dRC11-consistent, which follows by **Lemma 5.9**.

Next, we prove that (a) $\implies$ (b). Suppose that $\langle \bar{p}, G \rangle$ is reachable in $Pr \bowtie \text{FG}$ and G is dRC11-consistent. Relying on commutativity of operations of different threads in the program semantics (i.e., in the LTS induced by Pr), it suffices to demonstrate an enumeration $e_1, \dots, e_n$ of $G.E \setminus \text{Init}$ that respects $G.\text{po}$, such that

$$\langle \text{tid}(e_1), \text{lab}(e_1), \text{method}(e_1) \rangle, \dots, \langle \text{tid}(e_n), \text{lab}(e_n), \text{method}(e_n) \rangle$$

is obtained as a restriction of some trace of pRC11 to ProgLab. We construct this enumeration by taking an arbitrary total order extending several relations as defined next:

$$P \triangleq \{p \in \text{PLab} \mid E(p) \in G.E\}$$

$$R_{\text{prop}} \triangleq \{\langle e, p \rangle \in G.E \times P \mid E(p) = e\}$$

$$T \triangleq \{\langle p, p' \rangle \in P \times P \mid \begin{array}{l} \text{ptid}(p) = \text{ptid}(p') \wedge \text{psp}(p) = \text{psp}(p') \wedge \\ \langle E(p), E(p') \rangle \in [E_{\text{psp}(p)} \cup \text{CR}] ; G.\text{hb}_{\text{psp}(p)} ; [W_{\text{psp}(p)}^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}] \end{array}\}$$

$$R_{\text{rfp}} \triangleq \{\langle p, e \rangle \in P \times G.E \mid \langle E(p), e \rangle \in G.\text{rf}^+ \wedge \text{ptid}(p) = \text{tid}(e) \wedge e \in R_{\text{psp}(p)}^{\text{pacq}}\}$$

$$R_{\text{rf}} \triangleq \{\langle p, e \rangle \in P \times G.E \mid \langle E(p), e \rangle \in G.\text{rf}^+ \wedge \text{ptid}(p) = \text{tid}(e) \wedge e \in R^{\text{acq}}\}$$

$$R_{\text{fr}} \triangleq \{\langle r, p \rangle \in G.E \times P \mid \langle r, E(p) \rangle \in G.\text{fr} ; G.\text{rf}^? \wedge \text{ptid}(p) = \text{tid}(r) \wedge \text{psp}(p) = \text{sp}(r)\}$$

$$R_{\text{mo}} \triangleq \{\langle w, p \rangle \in G.E \times P \mid \langle w, E(p) \rangle \in G.\text{mo} ; G.\text{rf}^? \wedge \text{ptid}(p) = \text{tid}(w) \wedge \text{psp}(p) = \text{sp}(w)\}$$

$$R \triangleq G.\text{po} \cup G.\text{rf} \cup R_{\text{prop}} \cup T \cup R_{\text{rfp}} \cup R_{\text{rf}} \cup R_{\text{fr}} \cup R_{\text{mo}}$$

Claim A.0.1: T is transitive.

Claim A.0.2: $R_{\text{prop}} ; (R_{\text{rfp}} \cup R_{\text{rf}}) \subseteq \bigcup_X G.\text{rf}^+$.

Claim A.0.3: $R_{\text{prop}} ; T^+ ; (R_{\text{rfp}} \cup R_{\text{rf}}) \subseteq \bigcup_X G.\text{hb}_X$.

Claim A.0.4: $(R_{\text{fr}} \cup R_{\text{mo}}) ; (R_{\text{rfp}} \cup R_{\text{rf}}) \subseteq \bigcup_X [E_X] ; (G.\text{fr} \cup G.\text{mo}) ; G.\text{rf}^+$.

Claim A.0.5: $(R_{\text{fr}} \cup R_{\text{mo}}) ; T^+ ; (R_{\text{rfp}} \cup R_{\text{rf}}) \subseteq \bigcup_X [E_X] ; (G.\text{fr} \cup G.\text{mo}) ; G.\text{rf}^? ; G.\text{hb}_X$.

Claim A.0.6: $(R_{\text{fr}} \cup R_{\text{mo}}) ; T^* ; (R_{\text{rfp}} \cup R_{\text{rf}}) \subseteq G.\text{po}$.

Proof. Let $\langle e_1, e_2 \rangle \in (R_{\text{fr}} \cup R_{\text{mo}}) ; T^* ; (R_{\text{rfp}} \cup R_{\text{rf}})$. First, note that we must have $\text{tid}(e_1) = \text{tid}(e_2)$, and so we either have $\langle e_1, e_2 \rangle \in G.\text{po}$ or $\langle e_2, e_1 \rangle \in G.\text{po}$. Assume, by way of contradiction, that $\langle e_2, e_1 \rangle \in G.\text{po}$. By the previous two claims, we have $\langle e_1, e_2 \rangle \in \bigcup_X [E_X] ; (G.\text{fr} \cup G.\text{mo}) ; (G.\text{rf}^? ; G.\text{hb}_X \cup G.\text{rf}^+) ; G.\text{po}$, which entails a loop in $[E_X] ; (G.\text{fr} \cup G.\text{mo}) ; (G.\text{rf}^? ; G.\text{hb}_X \cup G.\text{rf}^+) ; G.\text{po}$, while both are forbidden by dRC11-consistency. $\triangleleft$

Claim A.0.7: $[E]; R^+; [E] \subseteq (G.\text{po} \cup G.\text{rf})^+.$

Proof. It is easy to see that $[E]; R^+; [E] \subseteq (G.\text{po} \cup G.\text{rf} \cup (R_{\text{prop}} \cup R_{\text{rf}} \cup R_{\text{mo}}); T^*; (R_{\text{rfp}} \cup R_{\text{rf}}))^+.$ Then, the claim easily follows from the previous claims. $\triangleleft$

Claim A.0.8: R is acyclic.

Proof. R -cycles between propagation actions are T -loops (since T is transitive), whereas R -cycles between events are $G.\text{po} \cup G.\text{rf}$ -cycles (using the previous claim). Both are forbidden by the **PO-RF** constraint in dRC11-consistency. $\triangleleft$

Now, let $\alpha_1, \dots, \alpha_n$ be an enumeration of $(G.E \setminus \text{Init}) \cup P$ that respects R . We prove that $\beta_1, \dots, \beta_n$ is a trace of pRC11, where:

$$\beta_i = \begin{cases} \langle \text{tid}(\alpha_i), \text{lab}(\alpha_i), \text{method}(\alpha_i) \rangle & \alpha_i \in G.E \\ \alpha_i & \alpha_i \in P \end{cases}$$

For that matter, let $G^0 = G_{\text{Init}}$ and $K^0 = K_{\text{Init}}$, and for every $1 \leq i \leq n$, let: $E_i = E \cap \{\alpha_1, \dots, \alpha_i\}$, $G^i = G|_{E_i}$, and $K^i = \lambda\tau. \lambda X. \text{Init} \cup G.E^\tau \cup \{e \in E \mid \exists k \leq i. \alpha_k = \text{EP}(e, \tau, X)\}$. We prove by induction on i that $\langle G^i, K^i \rangle \xrightarrow{\beta_i}_{\text{pRC11}} \langle G^{i+1}, K^{i+1} \rangle$ for every $0 \leq i \leq n-1$:

- Case 1— $\alpha_i \in G.E$ and $\text{typ}(\alpha_i) \in \{W, R, \text{RMW}\}$: Let $X = \text{sp}(\alpha_i)$, $\tau = \text{tid}(\alpha_i)$, and $k = K^i(\tau)$. Let w be the immediate $G^{i+1}.\text{mo}$ -predecessor of α_i if $\alpha_i \in W$, or the $G^{i+1}.\text{rf}$ -source of α_i if $\alpha_i \in R$. Then, we have $\alpha_i = \text{NextEvent}(G^i.E, \tau, \text{lab}(\alpha_i), \text{method}(\alpha_i))$, $w \in G^i.W_{X, \text{loc}(\alpha_i)}$, $\alpha_i \in R \implies \text{val}_W(w) = \text{val}_R(\alpha_i)$, $G^{i+1} = \text{Add}(G^i, \alpha_i, w)$, and $K^{i+1} = K^i[\tau \mapsto k \sqcup (\lambda Y. \{\alpha_i\})]$. Hence, $\langle G^i, K^i \rangle \xrightarrow{\beta_i}_{\text{pRC11}} \langle G^{i+1}, K^{i+1} \rangle$ follows from the following claims:
 - $w \notin \text{dom}(G^i.\text{mo}; G^i.\text{rf}^?; [k(X)])$: To see this, suppose otherwise. Let $e \in k(X)$ such that $\langle w, e \rangle \in G^i.\text{mo}; G^i.\text{rf}^?$. Then, we have $\langle \alpha_i, e \rangle \in (G^i.\text{mo} \cup G^i.\text{rf}); G^i.\text{rf}^?$, and so we also have $\langle \alpha_i, e \rangle \in (G.\text{mo} \cup G.\text{rf}); G.\text{rf}^?$. Let $p = \text{EP}(e, \tau, X)$, and let $k \leq i$ such that $\alpha_k = p$. It follows that $\langle \alpha_i, p \rangle \in R_{\text{mo}} \cup R_{\text{rf}}$. Since the enumeration respects $R_{\text{mo}} \cup R_{\text{rf}}$, this contradicts the fact that $k \leq i$.
 - If $\alpha_i \in W$, then $w \notin \text{dom}(G^i.\text{rf}; [\text{RMW}])$: To see this, suppose otherwise. Let $u \in \text{RMW}$ such that $\langle w, u \rangle \in G^i.\text{rf}$. Hence, we also have $\langle w, u \rangle \in G.\text{rf}$. However, our construction ensures that $\langle w, \alpha_i \rangle \in G.\text{mo}$ and $\langle \alpha_i, u \rangle \in G.\text{mo}$, which contradicts the fact that G satisfies the **RMW2** condition.
 - If $\alpha_i \in R^{\text{pacq}}$, then $\text{dom}(G^i.\text{rf}^*; [w]) \subseteq k(X)$: Suppose that $\alpha_i \in R^{\text{pacq}}$, and let w' such that $\langle w', w \rangle \in G^i.\text{rf}^*$. Then, $\langle w', \alpha_i \rangle \in G.\text{rf}^+$. Let $p = \text{EP}(w', \tau, X)$. Then, we have $\langle p, \alpha_i \rangle \in R_{\text{rfp}}$. Since the enumeration respects R_{rfp} , we have $\alpha_k = p$ for some $k < i$. The latter ensures that $w' \in k(X)$.
 - If $\alpha_i \in R^{\text{acq}}$, then $\forall Y. \text{dom}(G^i.\text{rf}^*; [w]) \subseteq k(Y)$: Let $Y \in \text{Space}$. Suppose that $\alpha_i \in R^{\text{acq}}$, and let w' such that $\langle w', w \rangle \in G^i.\text{rf}^*$. Then, $\langle w', \alpha_i \rangle \in G.\text{rf}^+$. Let $p = \text{EP}(w', \tau, Y)$. Then, we have $\langle p, \alpha_i \rangle \in R_{\text{rf}}$. Since the enumeration respects R_{rf} , we have $\alpha_k = p$ for some $k < i$. The latter ensures that $w' \in k(Y)$.
- Case 2— $\alpha_i \in G.E$ and $\text{typ}(\alpha_i) \in \{\text{CALL}, \text{RET}\}$: Let $\tau = \text{tid}(\alpha_i)$ and $k = K^i(\tau)$. Then, we have $\alpha_i = \text{NextEvent}(G^i.E, \tau, \text{lab}(\alpha_i), \text{method}(\alpha_i))$, $G^{i+1} = \text{Add}(G^i, \alpha_i)$, and $K^{i+1} = K^i[\tau \mapsto k \sqcup (\lambda Y. \{\alpha_i\})]$. Hence, $\langle G^i, K^i \rangle \xrightarrow{\beta_i}_{\text{pRC11}} \langle G^{i+1}, K^{i+1} \rangle$.
- Case 3— $\alpha_i = \text{EP}(e, \tau, X) \in P$: In this case, $\langle G^i, K^i \rangle \xrightarrow{\beta_i}_{\text{pRC11}} \langle G^{i+1}, K^{i+1} \rangle$ follows from the following:
 - $e \in G^i.E \setminus K^i(\tau)(X)$: Since the enumeration respects R_{prop} , we have that $\alpha_k = e$ for some $k < i$, and so $e \in G^i.E$. Since $\alpha_i = \text{EP}(e, \tau, X)$, we cannot have $\alpha_j = \text{EP}(e, \tau, X)$ for some $j < i$, and so $e \notin K^i(\tau)(X)$.

– $e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G^i.\text{hb}_X; [e]) \subseteq K^i(\tau)(X)$: Suppose that $e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}$, and let $e' \in E_X \cup \text{CR}$ such that $\langle e', e \rangle \in G^i.\text{hb}_X$. Let $p' = \text{EP}(e', \tau, X)$. Then, we have $\langle p', \alpha_i \rangle \in T$. Since the enumeration respects T , we have $\alpha_k = p'$ for some $k < i$. The latter ensures that $e' \in K^i(\tau)(X)$. $\square$

A.2 Properties of Program Semantics and Program State Composition

In this appendix we state and prove several claims on program semantics (§3) and the program state composition (Def. 6.7). Henceforth, we let $F \subseteq \mathbb{F}$ be an arbitrary set of methods.

PROPOSITION A.1. *Let L be a library that is safe for programs Pr and Pr' . Suppose that $\bar{p}' \xrightarrow{\alpha}_{Pr'[L]} \bar{p}$ with $\text{method}(\alpha) \in \text{dom}(L)$. Then, $\bar{p}' \xrightarrow{\alpha}_{Pr[L]} \bar{p}$.*

PROPOSITION A.2. *Let L and L' be libraries with the same domain that are safe for a program Pr . Suppose that $\bar{p}' \xrightarrow{\alpha}_{Pr[L']} \bar{p}$ with $\text{method}(\alpha) \notin \text{dom}(L)$. Then, $\bar{p}' \xrightarrow{\alpha}_{Pr[L]} \bar{p}$.*

PROPOSITION A.3. *Suppose that $\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{lib}}$. Let $\bar{p}_{\text{cl}}$ be a program state with $\bar{p}_{\text{cl}}(\text{tid}(\alpha)).f \notin F$. Then, $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. Let $\tau = \text{tid}(\alpha)$. We prove the claim by establishing $\forall \pi \in \text{Tid}. \bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. Let $\pi \in \text{Tid}$. If $\bar{p}_{\text{cl}}(\pi).f \notin F$, then $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}(\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. Otherwise, if $\bar{p}_{\text{cl}}(\pi).f \in F$, we have $\pi \neq \tau$, and hence $\bar{p}_{\text{lib}}(\pi) = \bar{p}'_{\text{lib}}(\pi)$. Therefore, we have $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \langle \bar{p}_{\text{lib}}(\pi).pc, \bar{p}_{\text{lib}}(\pi).f, \bar{p}_{\text{cl}}(\pi).pc_s, \bar{p}_{\text{cl}}(\pi).f \rangle = \langle \bar{p}'_{\text{lib}}(\pi).pc, \bar{p}'_{\text{lib}}(\pi).f, \bar{p}_{\text{cl}}(\pi).pc_s, \bar{p}_{\text{cl}}(\pi).f \rangle = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. $\square$

PROPOSITION A.4. *Suppose that $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$ with $\text{method}(\alpha) \in F$ and $\text{typ}(\alpha) \neq \text{RET}$. Let $\bar{p}_{\text{lib}}$ be a program state. Then, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}] = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. Let $\tau = \text{tid}(\alpha)$. We prove the claim by establishing $\forall \pi \in \text{Tid}. \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. Let $\pi \in \text{Tid}$. If $\pi \neq \tau$, then $\bar{p}_{\text{cl}}(\pi) = \bar{p}'_{\text{cl}}(\pi)$ and it is easy to see (irrespective of whether $\bar{p}_{\text{cl}}(\pi).f \in F$ or $\bar{p}_{\text{cl}}(\pi).f \notin F$) that $\bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi) = \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. When $\pi = \tau$, we have $\bar{p}'_{\text{cl}}(\tau).f \in F$, and since $\bar{p}'_{\text{cl}}(\tau).f \in F$ and $\text{typ}(\alpha) \neq \text{RET}$, we have $\bar{p}_{\text{cl}}(\tau).pc_s = \bar{p}'_{\text{cl}}(\tau).pc_s$ and $\bar{p}_{\text{cl}}(\tau).f = \bar{p}'_{\text{cl}}(\tau).f$. Therefore, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \langle \bar{p}_{\text{lib}}(\tau).pc, \bar{p}_{\text{lib}}(\tau).f, \bar{p}'_{\text{cl}}(\tau).pc_s, \bar{p}'_{\text{cl}}(\tau).f \rangle = \langle \bar{p}_{\text{lib}}(\tau).pc, \bar{p}_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, \bar{p}_{\text{cl}}(\tau).f \rangle = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$. $\square$

PROPOSITION A.5. *Suppose that $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$ with $\text{method}(\alpha) \notin F$ and $\alpha \notin \text{Call}_F$. Let $\bar{p}_{\text{lib}}$ be a program state. Then, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}] \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. Let $\alpha = \langle \tau, l_e, f \rangle$. For every $\pi \neq \tau$, we have $\bar{p}'_{\text{cl}}(\pi) = \bar{p}_{\text{cl}}(\pi)$, and so $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. It remains to show that $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) \xrightarrow{l_e, f}_{Pr(\tau)} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$. Since $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$, we have $\bar{p}'_{\text{cl}}(\tau).f = f$ and $\bar{p}_{\text{cl}}(\tau).f \notin F$. Hence, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \bar{p}'_{\text{cl}}(\tau)$ and $\bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \bar{p}_{\text{cl}}(\tau)$, and so $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) \xrightarrow{l_e, f}_{Pr(\tau)} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$. $\square$

PROPOSITION A.6. *Suppose that $\bar{p}'_{\text{lib}} \xrightarrow{\tau, l_e, f}_{Pr} \bar{p}_{\text{lib}}$ with $f \in F$ and $\text{typ}(l_e) \neq \text{RET}$. Let $\bar{p}_{\text{cl}}$ be a program state with $\bar{p}_{\text{cl}}(\tau).f = f$. Then, $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\tau, l_e, f}_{Pr} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. For every $\pi \neq \tau$, we have $\bar{p}'_{\text{lib}}(\pi) = \bar{p}_{\text{lib}}(\pi)$, and so $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. It remains to show that $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) \xrightarrow{l_e, f}_{Pr(\tau)} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$. Since $\bar{p}_{\text{cl}}(\tau).f = f \in F$, we have $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) = \langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle$ and $\bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \langle \bar{p}_{\text{lib}}(\tau).pc, \bar{p}_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle$. Since $\bar{p}'_{\text{lib}} \xrightarrow{\tau, l_e, f}_{Pr} \bar{p}_{\text{lib}}$, we have $\bar{p}'_{\text{lib}}(\tau) \xrightarrow{l_e, f}_{Pr(\tau)} \bar{p}_{\text{lib}}(\tau)$. Since $f \in$

F and $\text{typ}(l_\epsilon) \neq \text{RET}$, it must be the case that $\langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f \rangle \xrightarrow{l_\epsilon}_{Pr(\tau)(f)} \langle \bar{p}_{\text{lib}}(\tau).pc, \bar{p}_{\text{lib}}(\tau).f \rangle$. Then, by definition, we have $\langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle \xrightarrow{l_\epsilon f}_{Pr(\tau)} \langle \bar{p}_{\text{lib}}(\tau).pc, \bar{p}_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle$. $\square$

PROPOSITION A.7. *Let $f \in F$. Suppose that $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$ and $\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr'} \bar{p}_{\text{lib}}$ with $\alpha = \langle \tau, \text{CALL}(f, \phi), \text{main} \rangle$. Then, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. For every $\pi \neq \tau$, we have $\bar{p}'_{\text{cl}}(\pi) = \bar{p}_{\text{cl}}(\pi)$ and $\bar{p}'_{\text{lib}}(\pi) = \bar{p}_{\text{lib}}(\pi)$, and so $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. It remains to show that $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) \xrightarrow{l, \text{main}}_{Pr(\tau)} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$ where $l = \text{CALL}(f, \phi)$. Since $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$, we have $Pr(\tau)(\text{main})(\bar{p}'_{\text{cl}}(\tau).pc) = \text{call}(f)$, $\bar{p}'_{\text{cl}}(\tau).f = \phi$, $\bar{p}'_{\text{cl}}(\tau).pc_s = \perp$, $\bar{p}'_{\text{cl}}(\tau).f = \text{main}$, $\bar{p}_{\text{cl}}(\tau).pc_s = \bar{p}'_{\text{cl}}(\tau).pc + 1$, and $\bar{p}_{\text{cl}}(\tau).f = f$. Since $\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr'} \bar{p}_{\text{lib}}$, we have $\bar{p}_{\text{lib}}(\tau).pc = 0$ and $\bar{p}_{\text{lib}}(\tau).f = \phi$. Hence, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) = \langle \bar{p}'_{\text{cl}}(\tau).pc, \phi, \perp, \text{main} \rangle$ and $\bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \langle 0, \phi, \bar{p}'_{\text{cl}}(\tau).pc + 1, f \rangle$. By definition, we have $\langle \bar{p}'_{\text{cl}}(\tau).pc, \phi, \perp, \text{main} \rangle \xrightarrow{l, \text{main}}_{Pr(\tau)} \langle 0, \phi, \bar{p}'_{\text{cl}}(\tau).pc + 1, f \rangle$. $\square$

PROPOSITION A.8. *Let $f \in F$. Suppose that $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$ and $\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr'} \bar{p}_{\text{lib}}$ with $\alpha = \langle \tau, \text{RET}(\phi), f \rangle$. Then, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\alpha}_{Pr'} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$.*

PROOF. For every $\pi \neq \tau$, we have $\bar{p}'_{\text{cl}}(\pi) = \bar{p}_{\text{cl}}(\pi)$ and $\bar{p}'_{\text{lib}}(\pi) = \bar{p}_{\text{lib}}(\pi)$, and so $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\pi) = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\pi)$. It remains to show that $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) \xrightarrow{l, f}_{Pr'(\tau)} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau)$ where $l = \text{RET}(\phi)$. Since $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr} \bar{p}_{\text{cl}}$, we have $\bar{p}'_{\text{cl}}(\tau).pc_s = \bar{p}_{\text{cl}}(\tau).pc$, $\bar{p}'_{\text{cl}}(\tau).f = f$, $\bar{p}_{\text{cl}}(\tau).f = \phi$, $\bar{p}_{\text{cl}}(\tau).pc_s = \perp$, and $\bar{p}_{\text{cl}}(\tau).f = \text{main}$. Since $\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr'} \bar{p}_{\text{lib}}$, we have $Pr'(\tau)(f)(\bar{p}'_{\text{lib}}(\tau).pc) = \text{return}$ and $\bar{p}'_{\text{lib}}(\tau).f = \phi$. Hence, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) = \langle \bar{p}'_{\text{lib}}(\tau).pc, \phi, \bar{p}_{\text{cl}}(\tau).pc, f \rangle$ and $\bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \langle \bar{p}_{\text{cl}}(\tau).pc, \phi, \perp, \text{main} \rangle$. By definition, we have $\langle \bar{p}'_{\text{lib}}(\tau).pc, \phi, \bar{p}_{\text{cl}}(\tau).pc, f \rangle \xrightarrow{l, f}_{Pr'(\tau)} \langle \bar{p}_{\text{cl}}(\tau).pc, \phi, \perp, \text{main} \rangle$. $\square$

PROPOSITION A.9. *Suppose that α is enabled in $\bar{p}'_{\text{cl}}$ in Pr . Let $\bar{p}_{\text{lib}}$ be a program state. If $\text{method}(\alpha) \notin F$, then α is enabled in $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$ in Pr .*

PROOF. Let $\alpha = \langle \tau, l_\epsilon, f \rangle$. Since α is enabled in $\bar{p}'_{\text{cl}}$ in Pr , we have $\bar{p}'_{\text{cl}}(\tau).f = f \notin F$. Hence, $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}](\tau) = \bar{p}'_{\text{cl}}(\tau)$, from which the claim follows. $\square$

PROPOSITION A.10. *Suppose that α is enabled in $\bar{p}'_{\text{lib}}$ in Pr . Let $\bar{p}_{\text{cl}}$ be a program state with $\bar{p}_{\text{cl}}(\text{tid}(\alpha)).f = \text{method}(\alpha) \in F$. Then, α is enabled in $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}]$ in Pr .*

PROOF. Let $\alpha = \langle \tau, l_\epsilon, f \rangle$. Since $\bar{p}_{\text{cl}}(\tau).f \in F$, we have $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}](\tau) = \langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle$. Since α is enabled in $\bar{p}'_{\text{lib}}$ in Pr , $\langle l_\epsilon, f \rangle$ is enabled in $\bar{p}'_{\text{lib}}(\tau)$ in $Pr(\tau)$. Then, either (i) l_ϵ is enabled in $\langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f \rangle$ in $Pr(\tau)(f)$; or (ii) $l_\epsilon = \text{RET}(\phi)$, $Pr(\tau)(f)(\bar{p}'_{\text{lib}}(\tau).pc) = \text{return}$ and $\bar{p}'_{\text{lib}}(\tau).f = \phi$. In both cases, by definition, α is enabled in $\langle \bar{p}'_{\text{lib}}(\tau).pc, \bar{p}'_{\text{lib}}(\tau).f, \bar{p}_{\text{cl}}(\tau).pc_s, f \rangle$. $\square$

PROPOSITION A.11. *Suppose that $\langle \bar{p}, M \rangle \xrightarrow{t_1}_{Pr \bowtie \text{RC11}} \langle \bar{p}_1, M_1 \rangle$ and $\langle \bar{p}, M \rangle \xrightarrow{t_2}_{Pr' \bowtie \text{RC11}} \langle \bar{p}_2, M_2 \rangle$. If $H_F(t_1) = H_F(t_2)$, then the following hold for every $\tau \in \text{Tid}$:*

- If $\bar{p}_1(\tau).f \in F$, then $\bar{p}_2(\tau).f = \bar{p}_1(\tau).f$.
- If $\bar{p}_1(\tau).f \notin F$, then $\bar{p}_2(\tau).f \notin F$.

A.3 Proofs for Section 6

PROPOSITION A.12. *Let G be an execution graph, $F \subseteq F$, $\tau \in \text{Tid}$, $l \in \text{Lab}$, and $f \in F$.*

- If $f \in F$ or $\text{typ}(l) = \text{CALL}$ and $\text{callee}(l) \in F$, then $\text{NextEvent}(G|_{E_F}, \tau, l, f) = \text{NextEvent}(G.E, \tau, l, f)$.
- If $f \notin F$ or $\text{typ}(l) = \text{RET}$, then $\text{NextEvent}(G|_{E_F}, \tau, l, f) = \text{NextEvent}(G.E, \tau, l, f)$.

PROPOSITION A.13. Suppose that a set $F \subseteq \mathcal{F}$ is encapsulated in an execution graph G . Then, $G.\mathbf{SW}_{\text{base}} = G|_{E_F}.\mathbf{SW}_{\text{base}} \cup G|_{E_{\bar{F}}}.\mathbf{SW}_{\text{base}}$ and $G.\mathbf{SW}_X = G|_{E_F}.\mathbf{SW}_X \cup G|_{E_{\bar{F}}}.\mathbf{SW}_X$ for every $X \in \text{Space}$.

PROPOSITION A.14. Suppose that $F \subseteq \mathcal{F}$ is encapsulated in a well-formed execution graph G . Then, for every $X \in \text{Space}$:

- $G.\mathbf{hb}_X ; [\{e \in E \mid \text{method}(e) \in F\}] \subseteq (G.\mathbf{hb}_X ; [\mathbf{CR}_F])^? ; G|_{E_F}.\mathbf{hb}_X$.
- $G.\mathbf{hb}_X ; [\{e \in E \mid \text{method}(e) \notin F\}] \subseteq (G.\mathbf{hb}_X ; [\mathbf{CR}_F])^? ; G|_{E_{\bar{F}}}.\mathbf{hb}_X$.
- $[\{e \in E \mid \text{method}(e) \in F\}] ; G.\mathbf{hb}_X \subseteq G|_{E_F}.\mathbf{hb}_X ; ([\mathbf{CR}_F] ; G.\mathbf{hb}_X)^?$.
- $[\{e \in E \mid \text{method}(e) \notin F\}] ; G.\mathbf{hb}_X \subseteq G|_{E_{\bar{F}}}.\mathbf{hb}_X ; ([\mathbf{CR}_F] ; G.\mathbf{hb}_X)^?$.

LEMMA 6.15 (MERGE). Suppose that $F \subseteq \mathcal{F}$ is encapsulated in a well-formed memory state $M = \langle G, K \rangle$. Then, the following hold:

- (1) Let α be an pRC11 transition label with $\text{method}(\alpha) \in F$. Suppose that if $\alpha \in \text{ProgLab}$ and $\text{typ}(\alpha) \in \{\mathbf{W}, \mathbf{R}, \mathbf{RMW}\}$, then $\text{method}(e) \in F$ for every $e \in G.E$ with $\text{sp}(e) = \text{sp}(\alpha)$. Then, $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'_F$ implies that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' such that $M'|_F = M'_F$.
- (2) Let α be an pRC11 transition label with $\text{method}(\alpha) \notin F$. Suppose that if $\alpha \in \text{ProgLab}$ and $\text{typ}(\alpha) \in \{\mathbf{W}, \mathbf{R}, \mathbf{RMW}\}$, then $\text{method}(e) \notin F$ for every $e \in G.E$ with $\text{sp}(e) = \text{sp}(\alpha)$. Then, $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M'_{\bar{F}}$ implies that $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' such that $M'|_{\bar{F}} = M'_{\bar{F}}$.

PROOF. We prove each of the claims:

- (1) Suppose that $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'_F$. Let $M'_F = \langle G'_F, K'_F \rangle$. We consider the three possible pRC11-steps:
 - $\alpha = \langle \tau, l, f \rangle$ and $\text{typ}(l) \in \{\mathbf{W}, \mathbf{R}, \mathbf{RMW}\}$: Let $e = \text{NextEvent}(G|_{E_F}.E, \tau, l, f)$ and $X = \text{sp}(e)$. Then, there exists an event $w \in G|_{E_F}.\mathbf{W}_{X, \text{loc}(e)}$ such that the following hold:
 - $e \in \mathbf{R} \implies \text{val}_W(w) = \text{val}_R(e)$.
 - $w \notin \text{dom}(G|_{E_F}.\mathbf{mo} ; G|_{E_F}.\mathbf{rf}^? ; [K|_{E_F}(\tau)(X)])$.
 - $e \in \mathbf{W} \implies w \notin \text{dom}(G|_{E_F}.\mathbf{rf} ; [\mathbf{RMW}])$.
 - $e \in \mathbf{R}^{\text{pacq}} \implies \text{dom}(G|_{E_F}.\mathbf{rf}^* ; [w]) \subseteq K|_{E_F}(\tau)(X)$.
 - $e \in \mathbf{R}^{\text{acq}} \implies \forall Y. \text{dom}(G|_{E_F}.\mathbf{rf}^* ; [w]) \subseteq K|_{E_F}(\tau)(Y)$.
 - $G'_F = \text{Add}(G|_{E_F}, e, w)$.
 - $K'_F = K|_{E_F}[\tau \mapsto k'_F]$ where $k'_F = \lambda Y. K|_{E_F}(\tau)(Y) \cup \{e\}$.
 In addition, the assumption on α ensures that $f \in F$ and $\text{method}(w) \in F$ (since $\text{sp}(e) = \text{sp}(w)$). Let $G' = \text{Add}(G, e, w)$ and $K' = K[\tau \mapsto k']$, where $k' = \lambda Y. K(\tau)(Y) \cup \{e\}$. Let $M' = \langle G', K' \rangle$. It is easy to see that $G'|_{E_F} = G'_F$ (in particular, $e = \text{NextEvent}(G.E, \tau, l, f)$ by [Prop. A.12](#)) and $K'|_{E_F} = K'_F$. Then, $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ follows from the following:
 - $w \notin \text{dom}(G.\mathbf{mo} ; G.\mathbf{rf}^? ; [K(\tau)(X)])$: Suppose that $\langle w, a \rangle \in G.\mathbf{mo}$ and $\langle a, b \rangle \in G.\mathbf{rf}^?$ for some $a \in G.E$ and $b \in K(\tau)(X)$. Then, by definition, we have $\text{sp}(w) = \text{sp}(a) = \text{sp}(b)$. Since F is encapsulated in G and $\text{method}(w) \in F$, we must have $\text{method}(a) \in F$ and $\text{method}(b) \in F$. Hence, we have $\langle w, a \rangle \in G|_{E_F}.\mathbf{mo}$, $\langle a, b \rangle \in G|_{E_F}.\mathbf{rf}^?$, and $a \in K|_{E_F}(\tau)(X)$. This contradicts the fact that $w \notin \text{dom}(G|_{E_F}.\mathbf{mo} ; G|_{E_F}.\mathbf{rf}^? ; [K|_{E_F}(\tau)(X)])$.
 - $e \in \mathbf{W} \implies w \notin \text{dom}(G.\mathbf{rf} ; [\mathbf{RMW}])$: Suppose that $e \in \mathbf{W}$ but $\langle w, a \rangle \in G.\mathbf{rf}$ for some $a \in \mathbf{RMW}$. Then, by definition, we have $\text{sp}(w) = \text{sp}(a)$. Since F is encapsulated in G and $\text{method}(w) \in F$, we must have $\text{method}(a) \in F$. Hence, we have $\langle w, a \rangle \in G|_{E_F}.\mathbf{rf}$. This contradicts the fact that $e \in \mathbf{W} \implies w \notin \text{dom}(G|_{E_F}.\mathbf{rf} ; [\mathbf{RMW}])$.
 - $e \in \mathbf{R}^{\text{pacq}} \implies \text{dom}(G.\mathbf{rf}^* ; [w]) \subseteq K(\tau)(X)$: Suppose that $e \in \mathbf{R}^{\text{pacq}}$ but $\langle a_1, a_2 \rangle, \dots, \langle a_{n-1}, a_n \rangle \in G.\mathbf{rf}$ and $\langle a_n, w \rangle \in G.\mathbf{rf}^?$ for some $n \geq 1$ and $a_1, \dots, a_n$ such that $a_1 \notin K(\tau)(X)$. Then, by definition, we have $\text{sp}(a_1) = \dots = \text{sp}(a_n) = \text{sp}(w)$. Since F is encapsulated in G and $\text{method}(w) \in F$, we must have $\text{method}(a_i) \in F$ for every $1 \leq i \leq n$. Hence, we have $\langle a_1, w \rangle \in G|_{E_F}.\mathbf{rf}^*$. This contradicts the fact that $e \in \mathbf{R}^{\text{pacq}} \implies \text{dom}(G|_{E_F}.\mathbf{rf}^* ; [w]) \subseteq K|_{E_F}(\tau)(X)$.

- $e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G.\text{rf}^*; [w]) \subseteq K(\tau)(Y)$: Suppose that $e \in R^{\text{acq}}$ but there exist $Y \in \text{Space}$ and $a_1 \notin K(\tau)(Y)$ such that $\langle a_1, w \rangle \in G.\text{rf}^*$. Let $n \geq 1$ and $a_2, \dots, a_n$ such that $\langle a_1, a_2 \rangle, \dots, \langle a_{n-1}, a_n \rangle \in G.\text{rf}$ and $\langle a_n, w \rangle \in G.\text{rf}^?$. Then, by definition, we have $\text{sp}(a_1) = \dots = \text{sp}(a_n) = \text{sp}(w)$. Since F is encapsulated in G and $\text{method}(w) \in F$, we must have $\text{method}(a_i) \in F$ for every $1 \leq i \leq n$. Hence, we have $\langle a_1, w \rangle \in G|_{E_F}.\text{rf}^*$. This contradicts the fact that $e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G|_{E_F}.\text{rf}^*; [w]) \subseteq K|_{E_F}(\tau)(Y)$.
- $\alpha = \langle \tau, l, f \rangle$ and $\text{typ}(l) \in \{\text{CALL}, \text{RET}\}$: Let $e = \text{NextEvent}(G|_{E_F}.\text{E}, \tau, l, f)$. Then, we have $G'_F = \text{Add}(G|_{E_F}, e)$ and $K'_F = K|_{E_F}[\tau \mapsto k'_F]$, where $k'_F = \lambda Y. K|_{E_F}(\tau)(Y) \cup \{e\}$. Let $G' = \text{Add}(G, e)$ and $K' = K[\tau \mapsto k']$, where $k' = \lambda Y. K(\tau)(Y) \cup \{e\}$. Let $M' = \langle G', K' \rangle$. It is easy to see that $M \xrightarrow{\alpha}_{\text{pRC11}} M', G'|_{E_F} = G'_F$ (in particular, $e = \text{NextEvent}(G.\text{E}, \tau, l, f)$ by [Prop. A.12](#)), and $K'|_{E_F} = K'_F$.
- $\alpha = \text{EP}(e, \tau, X) \in \text{PLab}$: Then, the following hold:
 - $e \in G|_{E_F}.\text{E} \setminus K|_{E_F}(\tau)(X)$.
 - $e \in W_X^{\text{prel}} \cup W_X^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G|_{E_F}.\text{hb}_X; [e]) \subseteq K|_{E_F}(\tau)(X)$.
 - $G'_F = G|_{E_F}$ and $K'_F = K|_{E_F}[\tau \mapsto k'_F]$, where $k'_F = K|_{E_F}(\tau)[X \mapsto K|_{E_F}(\tau)(X) \cup \{e\}]$.
 Since $\text{method}(\alpha) \in F$, we clearly have that $e \in G.\text{E} \setminus K(\tau)(X)$. In addition the following holds:

Claim A.14.1: $e \in W_X^{\text{prel}} \cup W_X^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G.\text{hb}_X; [e]) \subseteq K(\tau)(X)$.

Proof. Suppose that $e \in W_X^{\text{prel}} \cup W_X^{\text{rel}} \cup \text{CR}$ (and so, $(E_X \cup \text{CR}) \cap \text{dom}(G|_{E_F}.\text{hb}_X; [e]) \subseteq K|_{E_F}(\tau)(X)$), and let $e' \in E_X \cup \text{CR}$ such that $\langle e', e \rangle \in G.\text{hb}_X$. By [Prop. A.14](#), $\langle e', e \rangle \in G.\text{hb}_X$ implies that $\langle e', e \rangle \in (G.\text{hb}_X; [\text{CR}_F])^?; G|_{E_F}.\text{hb}_X$. If $\langle e', e \rangle \in G|_{E_F}.\text{hb}_X$, then we have $e' \in K|_{E_F}(\tau)(X)$, and so $e' \in K(\tau)(X)$. Otherwise, let $e'' \in \text{CR}_F$ such that $\langle e', e'' \rangle \in G.\text{hb}_X$ and $\langle e'', e \rangle \in G|_{E_F}.\text{hb}_X$. Then, we have $e'' \in K|_{E_F}(\tau)(X)$, and so $e'' \in K(\tau)(X)$. Since K is well-formed for G , $e'' \in K(\tau)(X)$ and $\langle e', e'' \rangle \in G.\text{hb}_X$ together imply that $e' \in K(\tau)(X)$. $\blacktriangleleft$

Now, let $M' = \langle G', K' \rangle$, where $G' = G$ and $K' = K[\tau \mapsto k']$ for $k' = K(\tau)[X \mapsto K(\tau)(X) \cup \{e\}]$. Then, by definition we have $M \xrightarrow{\alpha}_{\text{pRC11}} M'$, and it is also easy to see that $G'|_{E_F} = G'_F$ and $K'|_{E_F} = K'_F$.

- (2) Suppose that $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M'_F$. Let $M'_F = \langle G'_F, K'_F \rangle$. We consider the three possible pRC11-steps:
- $\alpha = \langle \tau, l, f \rangle$ and $\text{typ}(l) \in \{\text{W}, \text{R}, \text{RMW}\}$: Let $e = \text{NextEvent}(G|_{\bar{E}_F}.\text{E}, \tau, l, f)$ and $X = \text{sp}(e)$. Then, there exists an event $w \in G|_{\bar{E}_F}.\text{W}_{X, \text{loc}(e)}$ such that the following hold:
 - $e \in R \implies \text{val}_W(w) = \text{val}_R(e)$.
 - $w \notin \text{dom}(G|_{\bar{E}_F}.\text{mo}; G|_{\bar{E}_F}.\text{rf}^?; [K|_{\bar{E}_F}(\tau)(X)])$.
 - $e \in W \implies w \notin \text{dom}(G|_{\bar{E}_F}.\text{rf}; [\text{RMW}])$.
 - $e \in R^{\text{pacq}} \implies \text{dom}(G|_{\bar{E}_F}.\text{rf}^*; [w]) \subseteq K|_{\bar{E}_F}(\tau)(X)$.
 - $e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G|_{\bar{E}_F}.\text{rf}^*; [w]) \subseteq K|_{\bar{E}_F}(\tau)(Y)$.
 - $G'_F = \text{Add}(G|_{\bar{E}_F}, e, w)$.
 - $K'_F = K|_{\bar{E}_F}[\tau \mapsto k'_F]$ where $k'_F = \lambda Y. K|_{\bar{E}_F}(\tau)(Y) \cup \{e\}$.
 In addition, the assumption on α ensures that $f \notin F$ and $\text{method}(w) \notin F$ (since $\text{sp}(e) = \text{sp}(w)$). Let $G' = \text{Add}(G, e, w)$ and $K' = K[\tau \mapsto k']$, where $k' = \lambda Y. K(\tau)(Y) \cup \{e\}$. Let $M' = \langle G', K' \rangle$. It is easy to see that $G'|_{\bar{E}_F} = G'_F$ (in particular, $e = \text{NextEvent}(G.\text{E}, \tau, l, f)$ by [Prop. A.12](#)) and $K'|_{\bar{E}_F} = K'_F$. Then, $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ follows from the following:
 - $w \notin \text{dom}(G.\text{mo}; G.\text{rf}^?; [K(\tau)(X)])$: Suppose that $\langle w, a \rangle \in G.\text{mo}$ and $\langle a, b \rangle \in G.\text{rf}^?$ for some $a \in G.\text{E}$ and $b \in K(\tau)(X)$. Then, by definition, we have $\text{sp}(w) = \text{sp}(a) = \text{sp}(b)$. Since F is encapsulated in G and $\text{method}(w) \notin F$, we must have $\text{method}(a) \notin F$ and $\text{method}(b) \notin F$. Hence, we have $\langle w, a \rangle \in G|_{\bar{E}_F}.\text{mo}$, $\langle a, b \rangle \in G|_{\bar{E}_F}.\text{rf}^?$, and $a \in K|_{\bar{E}_F}(\tau)(X)$. This contradicts the fact that $w \notin \text{dom}(G|_{\bar{E}_F}.\text{mo}; G|_{\bar{E}_F}.\text{rf}^?; [K|_{\bar{E}_F}(\tau)(X)])$.

- $e \in W \implies w \notin \text{dom}(G.\text{rf}; [\text{RMW}])$: Suppose that $e \in W$ but $\langle w, a \rangle \in G.\text{rf}$ for some $a \in \text{RMW}$. Then, by definition, we have $\text{sp}(w) = \text{sp}(a)$. Since F is encapsulated in G and $\text{method}(w) \notin F$, we must have $\text{method}(a) \notin F$. Hence, we have $\langle w, a \rangle \in G|_{E_F}.\text{rf}$. This contradicts the fact that $e \in W \implies w \notin \text{dom}(G|_{E_F}.\text{rf}; [\text{RMW}])$.
- $e \in R^{\text{pacq}} \implies \text{dom}(G.\text{rf}^*; [w]) \subseteq K(\tau)(X)$: Suppose that $e \in R^{\text{pacq}}$ but $\langle a_1, a_2 \rangle, \dots, \langle a_{n-1}, a_n \rangle \in G.\text{rf}$ and $\langle a_n, w \rangle \in G.\text{rf}^?$ for some $n \geq 1$ and $a_1, \dots, a_n$ such that $a_1 \notin K(\tau)(X)$. Then, by definition, we have $\text{sp}(a_1) = \dots = \text{sp}(a_n) = \text{sp}(w)$. Since F is encapsulated in G and $\text{method}(w) \notin F$, we must have $\text{method}(a_i) \notin F$ for every $1 \leq i \leq n$. Hence, we have $\langle a_1, w \rangle \in G|_{E_F}.\text{rf}^*$. This contradicts the fact that $e \in R^{\text{pacq}} \implies \text{dom}(G|_{E_F}.\text{rf}^*; [w]) \subseteq K|_{E_F}(\tau)(X)$.
- $e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G.\text{rf}^*; [w]) \subseteq K(\tau)(Y)$: Suppose that $e \in R^{\text{acq}}$ but there exist $Y \in \text{Space}$ and $a_1 \notin K(\tau)(Y)$ such that $\langle a_1, w \rangle \in G.\text{rf}^*$. Let $n \geq 1$ and $a_2, \dots, a_n$ such that $\langle a_1, a_2 \rangle, \dots, \langle a_{n-1}, a_n \rangle \in G.\text{rf}$ and $\langle a_n, w \rangle \in G.\text{rf}^?$. Then, by definition, we have $\text{sp}(a_1) = \dots = \text{sp}(a_n) = \text{sp}(w)$. Since F is encapsulated in G and $\text{method}(w) \notin F$, we must have $\text{method}(a_i) \notin F$ for every $1 \leq i \leq n$. Hence, we have $\langle a_1, w \rangle \in G|_{E_F}.\text{rf}^*$. This contradicts the fact that $e \in R^{\text{acq}} \implies \forall Y. \text{dom}(G|_{E_F}.\text{rf}^*; [w]) \subseteq K|_{E_F}(\tau)(Y)$.
- $\alpha = \langle \tau, l, f \rangle$ and $\text{typ}(l) \in \{\text{CALL}, \text{RET}\}$: Let $e = \text{NextEvent}(G|_{E_F}.E, \tau, l, f)$. Then, we have $G'_F = \text{Add}(G|_{E_F}, e)$ and $K'_F = K|_{E_F}[\tau \mapsto k'_F]$, where $k'_F = \lambda Y. K|_{E_F}(\tau)(Y) \cup \{e\}$. Let $G' = \text{Add}(G, e)$ and $K' = K[\tau \mapsto k']$, where $k' = \lambda Y. K(\tau)(Y) \cup \{e\}$. Let $M' = \langle G', K' \rangle$. It is easy to see that $M \xrightarrow{\alpha}_{\text{pRC11}} M', G'|_{E_F} = G'_F$ (in particular, $e = \text{NextEvent}(G.E, \tau, l, f)$ by [Prop. A.12](#)), and $K'|_{E_F} = K'_F$.
- $\alpha = \text{EP}(e, \tau, X) \in \text{PLab}$: Then, the following hold:
 - $e \in G|_{E_F}.E \setminus K|_{E_F}(\tau)(X)$.
 - $e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G|_{E_F}.\text{hb}_X; [e]) \subseteq K|_{E_F}(\tau)(X)$.
 - $G'_F = G|_{E_F}$ and $K'_F = K|_{E_F}[\tau \mapsto k'_F]$, where $k'_F = K|_{E_F}(\tau)[X \mapsto K|_{E_F}(\tau)(X) \cup \{e\}]$.
 Since $\text{method}(\alpha) \notin F$, we clearly have that $e \in G.E \setminus K(\tau)(X)$. In addition the following holds:

Claim A.14.2: $e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR} \implies (E_X \cup \text{CR}) \cap \text{dom}(G.\text{hb}_X; [e]) \subseteq K(\tau)(X)$.

Proof. Suppose that $e \in W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}$ (and so, $(E_X \cup \text{CR}) \cap \text{dom}(G|_{E_F}.\text{hb}_X; [e]) \subseteq K|_{E_F}(\tau)(X)$), and let $e' \in E_X \cup \text{CR}$ such that $\langle e', e \rangle \in G.\text{hb}_X$. By [Prop. A.14](#), $\langle e', e \rangle \in G.\text{hb}_X$ implies that $\langle e', e \rangle \in (G.\text{hb}_X; [\text{CR}_F])^?; G|_{E_F}.\text{hb}_X$. If $\langle e', e \rangle \in G|_{E_F}.\text{hb}_X$, then we have $e' \in K|_{E_F}(\tau)(X)$, and so $e' \in K(\tau)(X)$. Otherwise, let $e'' \in \text{CR}_F$ such that $\langle e', e'' \rangle \in G.\text{hb}_X$ and $\langle e'', e \rangle \in G|_{E_F}.\text{hb}_X$. Then, we have $e'' \in K|_{E_F}(\tau)(X)$, and so $e'' \in K(\tau)(X)$. Since K is well-formed for G , $e'' \in K(\tau)(X)$ and $\langle e', e'' \rangle \in G.\text{hb}_X$ together imply that $e' \in K(\tau)(X)$. $\triangleleft$

Now, let $M' = \langle G', K' \rangle$, where $G' = G$ and $K' = K[\tau \mapsto k']$ for $k' = K(\tau)[X \mapsto K(\tau)(X) \cup \{e\}]$. Then, by definition we have $M \xrightarrow{\alpha}_{\text{pRC11}} M'$, and it is also easy to see that $G'|_{E_F} = G'_F$ and $K'|_{E_F} = K'_F$. $\square$

PROPOSITION A.15 (DETERMINISM). *If $(\alpha \in \text{ProgLab}$ and $\text{typ}(\text{lab}(\alpha)) \in \{\text{CALL}, \text{RET}\})$ or $\alpha \in \text{PLab}$, then $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ and $M \xrightarrow{\alpha}_{\text{pRC11}} M''$ imply that $M' = M''$.*

PROPOSITION A.16. *Suppose that $F \subseteq F$ is encapsulated in a well-formed memory state M . If $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'_F$ and $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M''_F$ for $\alpha \in \text{HLab}_F$, then $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' such that $M'|_F = M'_F$ and $M'|_{\bar{F}} = M''_F$.*

PROOF. Suppose first that $\text{method}(\alpha) \in F$. Then, by [Lemma 6.15 \(Item 1\)](#), we have $M \xrightarrow{\alpha}_{\text{pRC11}} M'$ for some M' . By [Lemma 6.13](#), we have $M|_F \xrightarrow{\alpha}_{\text{pRC11}} M'|_F$ and $M|_{\bar{F}} \xrightarrow{\alpha}_{\text{pRC11}} M'|_{\bar{F}}$. By [Prop. A.15](#), it

follows that $M'|_F = M'_F$ and $M'|\bar{F} = M'_F$. The case that $\text{method}(\alpha) \notin F$ is similar using [Item 2 in Lemma 6.15](#). $\square$

A.4 Proof for Section 7

LEMMA 7.5 (COMPOSITION). *Let L and L' be libraries implementing the same set F of methods such that both are safe for a program Pr , and L is also safe for a program Pr' . Suppose that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t_{\text{cl}}}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}, M_{\text{cl}} \rangle$ and $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t_{\text{lib}}}_{Pr'[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{lib}}, M_{\text{lib}} \rangle$, with $H_F(t_{\text{cl}}) = H_F(t_{\text{lib}})$. Then, $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$ for some trace t and memory state M such that $t|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$, $t|_F = t_{\text{lib}}|_F$, $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$, and $M|_F = M_{\text{lib}}|_F$.*

PROOF. We prove the claim by induction on the sum of lengths of t_{cl} and t_{lib} . In the base case, $|t_{\text{cl}}| + |t_{\text{lib}}| = 0$, and we can simply take $t = \epsilon$ and $M = M_{\text{Init}}$. Suppose now that the claim holds for every t'_{cl} and t'_{lib} with $|t'_{\text{cl}}| + |t'_{\text{lib}}| < |t_{\text{cl}}| + |t_{\text{lib}}|$. We split the rest of the proof into the following three cases (which exhaust all the possibilities for t_{cl} and t_{lib}):

- (1) $t_{\text{cl}} = t'_{\text{cl}} \cdot \alpha_{\text{cl}}$ for $\alpha_{\text{cl}} \notin \text{HLab}_F$: Consider a state $\langle \bar{p}'_{\text{cl}}, M'_{\text{cl}} \rangle$ for which the following holds:

$$\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'_{\text{cl}}}_{Pr[L'] \bowtie \text{pRC11}} \langle \bar{p}'_{\text{cl}}, M'_{\text{cl}} \rangle \xrightarrow{\alpha_{\text{cl}}}_{Pr[L'] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}, M_{\text{cl}} \rangle$$

By the induction hypothesis, $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M' \rangle$ for some t' and M' , such that $t'|_{\bar{F}} = t'_{\text{cl}}|_{\bar{F}}$, $t'|_F = t_{\text{lib}}|_F$, $M'|_{\bar{F}} = M'_{\text{cl}}|_{\bar{F}}$, and $M'|_F = M_{\text{lib}}|_F$. We consider two cases:

- (a) $\text{method}(\alpha_{\text{cl}}) \in F$: Let $t \triangleq t'$ and $M \triangleq M'$. Then, the claim follows from the following:

- $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}] = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$: If $\alpha_{\text{cl}} \in \text{PLab}$ (i.e., memory internal step), then $\bar{p}'_{\text{cl}} = \bar{p}_{\text{cl}}$, and we are done. Otherwise, $\bar{p}'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{Pr[L']} \bar{p}_{\text{cl}}$, and the claim follows by [Prop. A.4](#).
- $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$: Since $M = M'$ and $M'|_{\bar{F}} = M'_{\text{cl}}|_{\bar{F}}$, it suffices to show that $M'_{\text{cl}}|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$. If $\text{lab}(\alpha_{\text{cl}}) = \epsilon$ (i.e., program internal step), then $M'_{\text{cl}} = M_{\text{cl}}$, and we are done. Otherwise, $M'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M_{\text{cl}}$, and by [Lemma 6.13](#) we have $M'_{\text{cl}}|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$.

- (b) $\text{method}(\alpha_{\text{cl}}) \notin F$: Let $t \triangleq t' \cdot \alpha_{\text{cl}}$. It is easy to see that $t|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$ and $t|_F = t_{\text{lib}}|_F$. It suffices to show that $\langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{cl}}}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$ for some M such that $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$ and $M|_F = M_{\text{lib}}|_F$. Consider the following cases:

- When $\text{lab}(\alpha_{\text{cl}}) = \epsilon$ (i.e., program internal step), we have $M'_{\text{cl}} = M_{\text{cl}}$ and $\bar{p}'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{Pr[L']} \bar{p}_{\text{cl}}$. Let $M \triangleq M'$. Then, we clearly have $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$ and $M|_F = M_{\text{lib}}|_F$. Since $\bar{p}'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{Pr[L']} \bar{p}_{\text{cl}}$, using [Propositions A.2](#) and [A.5](#) we have $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}] \xrightarrow{\alpha_{\text{cl}}}_{Pr[L]} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$. Hence, we also have $\langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{cl}}}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$.
- Otherwise, we have $M'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M_{\text{cl}}$. By [Lemma 6.13](#), we have $M'_{\text{cl}}|_{\bar{F}} \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M_{\text{cl}}|_{\bar{F}}$. Since $M'|_{\bar{F}} = M'_{\text{cl}}|_{\bar{F}}$, we have $M'|_{\bar{F}} \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M_{\text{cl}}|_{\bar{F}}$. In addition, since $M'|_{\bar{F}} \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M_{\text{cl}}|_{\bar{F}}$, using [Lemma 6.15-Item 2](#) (and establishing its precondition using [Prop. 6.2](#)), we have $M' \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M$ for some M such that $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$. By [Lemma 6.14](#), $M' \xrightarrow{\alpha_{\text{cl}}}_{\text{pRC11}} M$ implies that $M|_F = M'|_F$, and so we also have $M|_F = M_{\text{lib}}|_F$. To show that $\langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{cl}}}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$, consider each case:
 - $\alpha_{\text{cl}} \in \text{PLab}$ (i.e., memory internal step): In this case $\bar{p}'_{\text{cl}} = \bar{p}_{\text{cl}}$, and the claim easily follows.
 - Otherwise, we have $\bar{p}'_{\text{cl}} \xrightarrow{\alpha_{\text{cl}}}_{Pr[L']} \bar{p}_{\text{cl}}$. By [Propositions A.2](#) and [A.5](#) we have $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}] \xrightarrow{\alpha_{\text{cl}}}_{Pr[L]} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$, from which the claim follows.

- (2) $t_{\text{lib}} = t'_{\text{lib}} \cdot \alpha_{\text{lib}}$ for $\alpha_{\text{lib}} \notin \text{HLab}_F$: Consider a state $\langle \bar{p}'_{\text{lib}}, M'_{\text{lib}} \rangle$ for which the following holds:

$$\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'_{\text{lib}}}_{Pr'[L] \bowtie \text{pRC11}} \langle \bar{p}'_{\text{lib}}, M'_{\text{lib}} \rangle \xrightarrow{\alpha_{\text{lib}}}_{Pr'[L] \bowtie \text{pRC11}} \langle \bar{p}_{\text{lib}}, M_{\text{lib}} \rangle$$

By the induction hypothesis, $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle$ for some t' and M' , such that $t'|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$, $t'|_F = t'_{\text{lib}}|_F$, $M'|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$, and $M'|_F = M'_{\text{lib}}|_F$. We consider two cases:

(a) $\text{method}(\alpha_{\text{lib}}) \notin F$: Let $t \triangleq t'$ and $M \triangleq M'$. Then, the claim follows from the following:

- $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$: If $\alpha_{\text{lib}} \in \text{PLab}$ (i.e., memory internal step), then $\bar{p}'_{\text{lib}} = \bar{p}_{\text{lib}}$, and we are done. Otherwise, $\bar{p}'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{Pr'[L]} \bar{p}_{\text{lib}}$. Let $\tau = \text{tid}(\alpha)$. Since $\text{method}(\alpha_{\text{lib}}) \notin F$, we have $\bar{p}'_{\text{lib}}(\tau) \notin F$. By [Prop. A.11](#), we also have $\bar{p}_{\text{cl}}(\tau) \notin F$. Hence, the claim follows by [Prop. A.3](#).
- $M|_F = M_{\text{lib}}|_F$: Since $M = M'$ and $M'|_F = M'_{\text{lib}}|_F$, it suffices to show that $M'_{\text{lib}}|_F = M_{\text{lib}}|_F$. If $\text{lab}(\alpha_{\text{lib}}) = \epsilon$ (i.e., program internal step), then $M'_{\text{lib}} = M_{\text{lib}}$, and we are done. Otherwise, $M'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M_{\text{lib}}$, and by [Lemma 6.13](#) we have $M'_{\text{lib}}|_F = M_{\text{lib}}|_F$.

(b) $\text{method}(\alpha_{\text{lib}}) \in F$: Let $t \triangleq t' \cdot \alpha_{\text{lib}}$. It is easy to see that $t|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$ and $t|_F = t_{\text{lib}}|_F$. It suffices to show that $\langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{lib}}}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$ for some M such that $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$ and $M|_F = M_{\text{lib}}|_F$. Consider the following cases:

- When $\text{lab}(\alpha_{\text{lib}}) = \epsilon$ (i.e., program internal step), we have $M'_{\text{lib}} = M_{\text{lib}}$ and $\bar{p}'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{Pr'[L]} \bar{p}_{\text{lib}}$. Let $M \triangleq M'$. Then, we clearly have $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$ and $M|_F = M_{\text{lib}}|_F$. Since $\bar{p}'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{Pr'[L]} \bar{p}_{\text{lib}}$, using [Propositions A.1, A.6](#) and [A.11](#) we have $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\alpha_{\text{lib}}}_{Pr[L]} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$. Hence, we also have $\langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{lib}}}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$.
- Otherwise, we have $M'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M_{\text{lib}}$. By [Lemma 6.13](#), we have $M'_{\text{lib}}|_F \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M_{\text{lib}}|_F$. Since $M'|_F = M'_{\text{lib}}|_F$, we have $M'|_F \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M_{\text{lib}}|_F$. In addition, since $M'|_F \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M_{\text{lib}}|_F$, using [Lemma 6.15-Item 1](#) (and establishing its precondition using [Prop. 6.2](#)), we have $M' \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M$, for some M such that $M|_F = M_{\text{lib}}|_F$. By [Lemma 6.14](#), $M' \xrightarrow{\alpha_{\text{lib}}}_{pRC11} M$ implies that $M|_{\bar{F}} = M'|_{\bar{F}}$, and so we also have $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$. To show that $\langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle \xrightarrow{\alpha_{\text{lib}}}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$, consider each case:
 - $\alpha_{\text{cl}} \in \text{PLab}$ (i.e., memory internal step): In this case $\bar{p}'_{\text{lib}} = \bar{p}_{\text{lib}}$, and the claim easily follows.
 - Otherwise, we have $\bar{p}'_{\text{lib}} \xrightarrow{\alpha_{\text{lib}}}_{Pr'[L]} \bar{p}_{\text{lib}}$. By [Propositions A.1, A.6](#) and [A.11](#) we have $\bar{p}_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\alpha_{\text{lib}}}_{Pr[L]} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$, from which the claim follows.

(3) $t_{\text{cl}} = t'_{\text{cl}} \cdot \alpha_{\text{cl}}$ for $\alpha_{\text{cl}} \in \text{HLab}_F$ and $t_{\text{lib}} = t'_{\text{lib}} \cdot \alpha_{\text{lib}}$ for $\alpha_{\text{lib}} \in \text{HLab}_F$: In this case we have $\alpha_{\text{cl}} = \alpha_{\text{lib}}$. Let $\alpha = \alpha_{\text{cl}}$, and let $\langle \bar{p}'_{\text{cl}}, M'_{\text{cl}} \rangle$ and $\langle \bar{p}'_{\text{lib}}, M'_{\text{lib}} \rangle$ be states for which the following hold:

$$\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'_{\text{cl}}}_{Pr[L'] \bowtie pRC11} \langle \bar{p}'_{\text{cl}}, M'_{\text{cl}} \rangle \xrightarrow{\alpha}_{Pr[L'] \bowtie pRC11} \langle \bar{p}_{\text{cl}}, M_{\text{cl}} \rangle$$

$$\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'_{\text{lib}}}_{Pr'[L] \bowtie pRC11} \langle \bar{p}'_{\text{lib}}, M'_{\text{lib}} \rangle \xrightarrow{\alpha}_{Pr'[L] \bowtie pRC11} \langle \bar{p}_{\text{lib}}, M_{\text{lib}} \rangle$$

By the induction hypothesis, $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'}_{Pr[L] \bowtie pRC11} \langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle$ for some t' and M' , such that $t'|_{\bar{F}} = t'_{\text{cl}}|_{\bar{F}}$, $t'|_F = t'_{\text{lib}}|_F$, $M'|_{\bar{F}} = M'_{\text{cl}}|_{\bar{F}}$, and $M'|_F = M'_{\text{lib}}|_F$. Let $t \triangleq t' \cdot \alpha$. It is easy to see that $t|_{\bar{F}} = t_{\text{cl}}|_{\bar{F}}$ and $t|_F = t_{\text{lib}}|_F$. Since $\alpha \in \text{HLab}_F$, we have $\text{lab}(\alpha) \neq \epsilon$ and so $M'_{\text{cl}} \xrightarrow{\alpha}_{pRC11} M_{\text{cl}}$ and $M'_{\text{lib}} \xrightarrow{\alpha}_{pRC11} M_{\text{lib}}$. Now, since $M'_{\text{cl}} \xrightarrow{\alpha}_{pRC11} M_{\text{cl}}$, using [Lemma 6.13](#), we have $M'_{\text{cl}}|_{\bar{F}} \xrightarrow{\alpha}_{pRC11} M_{\text{cl}}|_{\bar{F}}$, and so $M'|_{\bar{F}} \xrightarrow{\alpha}_{pRC11} M_{\text{cl}}|_{\bar{F}}$. Similarly, since $M'_{\text{lib}} \xrightarrow{\alpha}_{pRC11} M_{\text{lib}}$, using [Lemma 6.13](#), we have $M'_{\text{lib}}|_F \xrightarrow{\alpha}_{pRC11} M_{\text{lib}}|_F$, and so $M'|_F \xrightarrow{\alpha}_{pRC11} M_{\text{lib}}|_F$. Therefore, using [Prop. A.16](#), we obtain that $M' \xrightarrow{\alpha}_{pRC11} M$ for some M such that $M|_{\bar{F}} = M_{\text{cl}}|_{\bar{F}}$ and $M|_F = M_{\text{lib}}|_F$. To show that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$, it remains to show that $\langle \bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}], M' \rangle \xrightarrow{\alpha}_{Pr[L] \bowtie pRC11} \langle \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}], M \rangle$. If $\alpha \in \text{PLab}$, then $\bar{p}'_{\text{lib}} = \bar{p}_{\text{lib}}$ and $\bar{p}'_{\text{cl}} = \bar{p}_{\text{cl}}$, and so $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] = \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$ and we are done using the fact that $M' \xrightarrow{\alpha}_{pRC11} M$. Otherwise, we have $\bar{p}'_{\text{cl}} \xrightarrow{\alpha}_{Pr[L']} \bar{p}_{\text{cl}}$ and

$\bar{p}'_{\text{lib}} \xrightarrow{\alpha}_{Pr'[L]} \bar{p}_{\text{lib}}$. By [Propositions A.1, A.2, A.7](#) and [A.8](#), we obtain that $\bar{p}'_{\text{cl}}[F \mapsto \bar{p}'_{\text{lib}}] \xrightarrow{\alpha}_{Pr[L]} \bar{p}_{\text{cl}}[F \mapsto \bar{p}_{\text{lib}}]$, from which the claim follows. $\square$

LEMMA 7.6. *Under the conditions of [Thm. 7.3](#), $H_F(Pr[L]) \subseteq H_F(MGC[L^\#])$.*

PROOF. Assume otherwise, and let h be a shortest history in $H_F(Pr[L]) \setminus H_F(MGC[L^\#])$. Let t be a shortest trace of $Pr[L] \bowtie \text{pRC11}$ with $H_F(t) = h$. Since ε (the empty history) is clearly in $H_F(MGC[L^\#])$, we know that t is non-empty. Consider the last transition label α in t , and let t' such that $t = t' \cdot \alpha$. The minimality of t ensures that α must be an element of HLab_F (i.e., call, return, call propagation or return propagation of some method in F). The minimality of h further ensures that $H_F(t') \in H_F(MGC[L^\#])$. Let $t'_\#$ and $\langle \bar{p}'_\#, M'_\# \rangle$ such that $H_F(t'_\#) = H_F(t')$ and $\langle \bar{p}_{\text{init}}, M_{\text{init}} \rangle \xrightarrow{t'_\#}_{MGC[L^\#] \bowtie \text{pRC11}} \langle \bar{p}'_\#, M'_\# \rangle$. Let $\langle \bar{p}', M' \rangle$ be such that $\langle \bar{p}_{\text{init}}, M_{\text{init}} \rangle \xrightarrow{t'}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}', M' \rangle \xrightarrow{\alpha}_{Pr[L] \bowtie \text{pRC11}} \langle \bar{p}, M \rangle$. We consider the following cases:

- (1) $\alpha \in \text{Call}_F \cup \text{CP}_F$: Using [Lemma 7.5](#) (applied with $L := L^\#$, $L' := L$, $Pr := Pr$, and $Pr' := MGC$),

there exist $t''_\#$ and $M''_\#$ such that $\langle \bar{p}_{\text{init}}, M_{\text{init}} \rangle \xrightarrow{t''_\#}_{Pr[L^\#] \bowtie \text{pRC11}} \langle \bar{p}'[F \mapsto \bar{p}'_\#], M''_\# \rangle$, $H_F(t''_\#) = H_F(t')$, and $M''_\#|_F = M'|_F$.

Claim A.16.1: α is enabled in $M''_\#$.

Proof. If $\alpha \in \text{Call}_F$, then the claim is trivial. Otherwise, since α is enabled in M' , by [Lemma 6.13](#), α is also enabled in $M'|_F = M''_\#|_F$. By [Lemma 6.15 \(Item 2\)](#), α is also enabled in $M''_\#$. $\triangleleft$

Claim A.16.2: If $\alpha \in \text{Call}_F$, then α is enabled in $\bar{p}'[F \mapsto \bar{p}'_\#]$ (in the LTS $Pr[L^\#]$).

Proof. Since $\bar{p}' \xrightarrow{\alpha}_{Pr[L]} \bar{p}$, the claim follows from [Propositions A.2](#) and [A.9](#). $\triangleleft$

It follows that $h = H_F(t') \cdot \alpha = H_F(t''_\#) \cdot \alpha \in H_F(Pr[L^\#])$. Since $Pr \sqsubseteq_{L^\#} MGC$, we obtain $h \in H_F(MGC[L^\#])$, which contradicts our assumption.

- (2) $\alpha \in \text{Ret}_F \cup \text{RP}_F$: Using [Lemma 7.5](#) (applied with $L := L$, $L' := L^\#$, $Pr := MGC$, and $Pr' := Pr$), there

exist t'' and M'' such that $\langle \bar{p}_{\text{init}}, M_{\text{init}} \rangle \xrightarrow{t''}_{MGC[L] \bowtie \text{pRC11}} \langle \bar{p}'[F \mapsto \bar{p}'], M'' \rangle$, $H_F(t'') = H_F(t')$, and $M''|_F = M'|_F$.

Claim A.16.3: α is enabled in M'' .

Proof. If $\alpha \in \text{Ret}_F$, then the claim is trivial. Otherwise, since α is enabled in M' , by [Lemma 6.13](#), α is also enabled in $M'|_F = M''|_F$. By [Lemma 6.15 \(Item 1\)](#), α is also enabled in M'' . $\triangleleft$

Claim A.16.4: If $\alpha \in \text{Ret}_F$, then α is enabled in $\bar{p}'_\#[F \mapsto \bar{p}']$ (in the LTS $MGC[L]$).

Proof. Let $\tau = \text{Tid}(\alpha)$. Since $\alpha \in \text{Ret}_F$, we have $\text{method}(\alpha) \in F$, and since $\bar{p}' \xrightarrow{\alpha}_{Pr[L]} \bar{p}$, we also have $\bar{p}'(\tau).f = \text{method}(\alpha)$. Since $H_F(t') = H_F(t''_\#)$, by [Prop. A.11](#), we have $\bar{p}'_\#(\tau).f = \text{method}(\alpha)$ as well. Thus, since $\bar{p}' \xrightarrow{\alpha}_{Pr[L]} \bar{p}$, the claim follows from [Propositions A.1](#) and [A.10](#). $\triangleleft$

It follows that $h = H_F(t') \cdot \alpha = H_F(t'') \cdot \alpha \in H_F(MGC[L])$. Since $L \sqsubseteq_{MGC} L^\#$, we obtain $h \in H_F(MGC[L^\#])$, which contradicts our assumption. $\square$

LEMMA 7.7. *Suppose that $M_{\text{init}} \xrightarrow{t}_{\text{pRC11}} \langle G, K \rangle$. Then, there exist a trace t' and a knowledge mapping K' for G such that the following hold:*

- $M_{\text{init}} \xrightarrow{t'}_{\text{pRC11}} \langle G, K' \rangle$.
- $t'|_{\text{ProgLab}} = t|_{\text{ProgLab}}$.
- $K'(\tau, X) \cap (W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}) \subseteq \text{dom}(G.\text{hb}_X^?; [E^\tau])$ for every $\tau \in \text{Tid}$ and $X \in \text{Space}$.

PROOF (SKETCH). Let T be the set of all traces t_0 such that (1) $M_{\text{init}} \xrightarrow{t_0}_{\text{pRC11}} \langle G, K_0 \rangle$ for some knowledge mapping K_0 for G ; and (2) $t_0|_{\text{ProgLab}} = t|_{\text{ProgLab}}$. Then, $t \in T$, and so T is not empty. Let t' be a trace in T of minimal length, and let K' such that $M_{\text{init}} \xrightarrow{t'}_{\text{pRC11}} \langle G, K' \rangle$. Suppose (by way of contradiction) that $K'(\tau, X) \cap (W_X^{\text{prel}} \cup W^{\text{rel}} \cup \text{CR}) \not\subseteq \text{dom}(G.\text{hb}_X^?; [E^\tau])$ for some $\tau \in \text{Tid}$ and

$X \in \text{Space}$. Let t'' be the trace obtained from t' by omitting the last transition label in t' of the form $\text{EP}(e, \tau, X)$ with $e \in W_X^{\text{prel}} \cup W_X^{\text{rel}} \cup \text{CR}$ and $e \notin \text{dom}(G.\text{hb}_X^?; [E^\tau])$. Then, it can be shown that $t'' \in T$, which contradicts the minimality of t' . $\square$

THEOREM 7.3 (LIBRARY ABSTRACTION). *Let L and $L^\#$ be libraries implementing the same set F of methods. Let MGC and Pr be programs, such that both L and $L^\#$ are safe for both MGC and Pr . Suppose that $L \sqsubseteq_{\text{MGC}} L^\#$ and $\text{Pr} \sqsubseteq_{L^\#} \text{MGC}$. Then, the following hold:*

- If $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t}_{\text{Pr}[L] \bowtie \text{pRC11}} \langle \bar{p}, M \rangle$, then there exist $t^\#$ and $\langle \bar{p}^\#, M^\# \rangle$ such that the following hold:
 - (1) $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t^\#}_{\text{Pr}[L^\#] \bowtie \text{pRC11}} \langle \bar{p}^\#, M^\# \rangle$; (2) $t^\#|_{\bar{F}} = t|_{\bar{F}}$; (3) $\bar{p}^\# = \bar{p}[F \mapsto \bar{p}_{\text{lib}}]$ for some $\bar{p}_{\text{lib}}$ (in particular, $\bar{p}^\#(\tau) = \bar{p}(\tau)$ whenever $\bar{p}(\tau).f \notin F$); and (4) $M^\#|_{\bar{F}} = M|_{\bar{F}}$.
- $\text{Pr}[L]$ is not racy under pRC11 .

PROOF. By Lemma 7.6, we have $H_F(\text{Pr}[L]) \subseteq H_F(\text{MGC}[L^\#])$. Then, the first part of the claim directly follows using Lemma 7.5 by letting $L := L^\#$, $L' := L$, $\text{Pr} := \text{Pr}$, and $\text{Pr}' := \text{MGC}$.

For the second part, suppose that $\text{Pr}[L]$ is racy under pRC11 . Let $\langle \bar{p}, \langle G, K \rangle \rangle$ be a reachable state in $\text{Pr}[L] \bowtie \text{pRC11}$ such that G is racy. Let e_1, e_2 be two events that form a race in G . Let $E = \text{dom}((G.\text{po} \cup G.\text{rf})^*; [e_1, e_2])$ and $G' = G|_E$. Clearly, e_1 and e_2 form a race in G' . It is also easy to see that there exist $\bar{p}'$ and K' such that $\langle \bar{p}', \langle G', K' \rangle \rangle$ is reachable in $\text{Pr}[L] \bowtie \text{pRC11}$. By Lemma 7.7, there exists a knowledge mapping K'' for G' such that $\langle \bar{p}', \langle G', K'' \rangle \rangle$ is reachable in $\text{Pr}[L] \bowtie \text{pRC11}$ as well and $K''(\tau, X) \cap \text{CR} \subseteq \text{dom}(G'.\text{hb}_X^?; [E^\tau])$ for every $\tau \in \text{Tid}$ and $X \in \text{Space}$. Let $M = \langle G', K'' \rangle$ and $X = \text{sp}(e_1) = \text{sp}(e_2)$. By Prop. 6.6, F is encapsulated in M . Thus, since $\text{sp}(e_1) = \text{sp}(e_2)$, it follows that one of the following holds:

- (1) $e_1 \in E_{\bar{F}}$ and $e_2 \in E_{\bar{F}}$: Since $H_F(\text{Pr}[L]) \subseteq H_F(\text{MGC}[L^\#])$, by applying Lemma 7.5 with $L := L^\#$, $L' := L$, $\text{Pr} := \text{Pr}$, and $\text{Pr}' := \text{MGC}$, we obtain a program state $\bar{p}_\#$ and a memory state $M_\# = \langle G_\#, K_\# \rangle$ such that $\langle \bar{p}'[F \mapsto \bar{p}_\#], M_\# \rangle$ is reachable in $\text{Pr}[L^\#] \bowtie \text{pRC11}$ and $M_\#|_{\bar{F}} = M|_{\bar{F}}$. Since $\text{Pr}[L^\#]$ is not racy under pRC11 , we have $\langle e_1, e_2 \rangle \in G_\#.\text{hb}_X \cup G_\#.\text{hb}_X^{-1}$. Suppose w.l.o.g. that $\langle e_1, e_2 \rangle \in G_\#.\text{hb}_X$. Let $\tau = \text{tid}(e_2)$.

Claim A.16.1: $\langle e_1, e_2 \rangle \in G_\#|_{E_{\bar{F}}}.\text{hb}_X; [\text{CR}_F]; G_\#.\text{hb}_X$.

Proof. By Prop. A.14, we have $\langle e_1, e_2 \rangle \in G_\#|_{E_{\bar{F}}}.\text{hb}_X; ([\text{CR}_F]; G_\#.\text{hb}_X)^?$. Since $G_\#|_{E_{\bar{F}}}.\text{hb}_X = G'|_{E_{\bar{F}}}.\text{hb}_X \subseteq G'.\text{hb}_X$, and e_1, e_2 form a race in G' , we cannot have $\langle e_1, e_2 \rangle \in G_\#|_{E_{\bar{F}}}.\text{hb}_X$. $\triangleleft$

Let $f \in \text{CR}_F$ such that $\langle e_1, f \rangle \in G_\#|_{E_{\bar{F}}}.\text{hb}_X$ and $\langle f, e_2 \rangle \in G_\#.\text{hb}_X$.

Claim A.16.2: $\langle e_1, f \rangle \in G'.\text{hb}_X$.

Proof. Immediately follows using the fact that $G_\#|_{E_{\bar{F}}}.\text{hb}_X = G'|_{E_{\bar{F}}}.\text{hb}_X \subseteq G'.\text{hb}_X$. $\triangleleft$

Claim A.16.3: $f \in \text{dom}(G'.\text{hb}_X^?; [E^\tau])$.

Proof. Using Prop. 5.8, $\langle f, e_2 \rangle \in G_\#.\text{hb}_X$ implies that $f \in K_\#(\tau, X)$. Since $M_\#|_{\bar{F}} = M|_{\bar{F}}$, we have $f \in K''(\tau, X)$ as well. Since $K''(\tau, X) \cap \text{CR} \subseteq \text{dom}(G'.\text{hb}_X^?; [E^\tau])$, it follows that $f \in \text{dom}(G'.\text{hb}_X^?; [E^\tau])$. $\triangleleft$

From the two last claims, we obtain that $e_1 \in \text{dom}(G'.\text{hb}_X; [E^\tau])$.

Claim A.16.4: e_2 is G' .po-maximal.

Proof. Suppose otherwise, and let e be a G' .po-maximal event with $\text{tid}(e) = \tau$. Then, $e_1 \in \text{dom}(G'.\text{hb}_X; [E^\tau])$ implies that $\langle e_1, e \rangle \in G'.\text{hb}_X$. It follows that $\langle e_1, e \rangle \in (G'.\text{po} \cup G'.\text{rf})^+$. However, since $e \in G'.E$, we either have $\langle e, e_1 \rangle \in (G'.\text{po} \cup G'.\text{rf})^+$ or $\langle e, e_2 \rangle \in (G'.\text{po} \cup G'.\text{rf})^+$, and both contradict the fact that G' is dRC11-consistent by Lemma 5.9 (as they violate **PO-RF**). $\triangleleft$

Using the last claim, the fact that $e_1 \in \text{dom}(G'.\text{hb}_X; [E^\tau])$ implies that $\langle e_1, e_2 \rangle \in G'.\text{hb}_X$, which contradicts the fact that e_1 and e_2 form a race in G' .

- (2) $e_1 \in E_F$ and $e_2 \in E_F$: Since $H_F(\text{Pr}[L]) \subseteq H_F(\text{MGC}[L^\#])$, by applying [Lemma 7.5](#) with $L := L$, $L' := L^\#$, $\text{Pr} := \text{MGC}$, and $\text{Pr}' := \text{Pr}$, we obtain a program state $\bar{p}^*$ and a memory state $M^* = \langle G^*, K^* \rangle$ such that $\langle \bar{p}^*[F \mapsto \bar{p}^*], M^* \rangle$ is reachable in $\text{MGC}[L] \bowtie \text{pRC11}$ and $M^*|_F = M|_F$. Since $\text{MGC}[L]$ is not racy under pRC11 we have $\langle e_1, e_2 \rangle \in G'.\text{hb}_X \cup G'.\text{hb}_X^{-1}$. Suppose w.l.o.g. that $\langle e_1, e_2 \rangle \in G'.\text{hb}_X$. Let $\tau = \text{Tid}(e_2)$.

Claim A.16.5: $\langle e_1, e_2 \rangle \in G^*|_{E_F}.\text{hb}_X ; [\text{CR}_F] ; G^*.\text{hb}_X$.

Proof. By [Prop. A.14](#), we have $\langle e_1, e_2 \rangle \in G^*|_{E_F}.\text{hb}_X ; ([\text{CR}_F] ; G^*.\text{hb}_X)^\tau$. Since $G^*|_{E_F}.\text{hb}_X = G'|_{E_F}.\text{hb}_X \subseteq G'.\text{hb}_X$, and e_1, e_2 form a race in G' , we cannot have $\langle e_1, e_2 \rangle \in G^*|_{E_F}.\text{hb}_X$. $\triangleleft$
Let $f \in \text{CR}_F$ such that $\langle e_1, f \rangle \in G^*|_{E_F}.\text{hb}_X$ and $\langle f, e_2 \rangle \in G^*.\text{hb}_X$.

Claim A.16.6: $\langle e_1, f \rangle \in G'.\text{hb}_X$.

Proof. Immediately follows using the fact that $G^*|_{E_F}.\text{hb}_X = G'|_{E_F}.\text{hb}_X \subseteq G'.\text{hb}_X$. $\triangleleft$

Claim A.16.7: $f \in \text{dom}(G'.\text{hb}_X^\tau ; [E^\tau])$.

Proof. Using [Prop. 5.8](#), $\langle f, e_2 \rangle \in G^*.\text{hb}_X$ implies that $f \in K^*(\tau, X)$. Since $M^*|_F = M|_F$, we have $f \in K''(\tau, X)$ as well. Since $K''(\tau, X) \cap \text{CR} \subseteq \text{dom}(G'.\text{hb}_X^\tau ; [E^\tau])$, it follows that $f \in \text{dom}(G'.\text{hb}_X^\tau ; [E^\tau])$. $\triangleleft$

From the two last claims, we obtain that $e_1 \in \text{dom}(G'.\text{hb}_X ; [E^\tau])$.

Claim A.16.8: e_2 is G' .po-maximal.

Proof. Similar to the proof of [Claim A.16.4](#). $\triangleleft$

Using the last claim, the fact that $e_1 \in \text{dom}(G'.\text{hb}_X ; [E^\tau])$ implies that $\langle e_1, e_2 \rangle \in G'.\text{hb}_X$, which contradicts the fact that e_1 and e_2 form a race in G' . $\square$

COROLLARY 7.4 (COMPOSITIONALITY). *Let $L_1, \dots, L_n$ be libraries implementing pairwise disjoint sets of methods, such that $\text{Space}(L_1), \dots, \text{Space}(L_n), \text{Space}(L_1^\#), \dots, \text{Space}(L_n^\#)$, and $\text{Space}(\text{MGC} \setminus \text{dom}(L_1 \uplus \dots \uplus L_n))$ are pairwise disjoint. Suppose that for every $1 \leq i \leq n$, we have $L_i \sqsubseteq_{\text{MGC}_i} L_i^\#$ for $\text{MGC}_i = \text{MGC}[L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_n^\#]$. Then, $L_1 \uplus \dots \uplus L_n \sqsubseteq_{\text{MGC}} L_1^\# \uplus \dots \uplus L_n^\#$.*

PROOF. We prove the claim by induction on n . For $n = 1$, the claim trivially holds. For the induction step, let $L_1, \dots, L_n, L_1^\#, \dots, L_n^\#$ be libraries, and let MGC be a program satisfying the required conditions. For $\text{MGC}' = \text{MGC}[L_n^\#]$, we have that $\text{Space}(L_1), \dots, \text{Space}(L_{n-1}), \text{Space}(L_1^\#), \dots, \text{Space}(L_{n-1}^\#)$, and $\text{Space}(\text{MGC}' \setminus \text{dom}(L_1 \uplus \dots \uplus L_{n-1}))$ are pairwise disjoint. In addition, for every $1 \leq i \leq n-1$,

$$\begin{aligned} \text{MGC}'[L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_{n-1}^\#] &= \text{MGC}[L_n^\#][L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_{n-1}^\#] = \\ &= \text{MGC}[L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_n^\#]. \end{aligned}$$

Hence, for every $1 \leq i \leq n-1$, we have $L_i \sqsubseteq_{\text{MGC}_i} L_i^\#$ for $\text{MGC}_i = \text{MGC}'[L_1^\# \uplus \dots \uplus L_{i-1}^\# \uplus L_{i+1}^\# \uplus \dots \uplus L_{n-1}^\#]$. By the induction hypothesis, it follows that $L \sqsubseteq_{\text{MGC}'} L^\#$ for $L = L_1 \uplus \dots \uplus L_{n-1}$ and $L^\# = L_1^\# \uplus \dots \uplus L_{n-1}^\#$, which implies that $\text{MGC}[L] \sqsubseteq_{L_n^\#} \text{MGC}[L^\#]$. In addition, by assumption we have $L_n \sqsubseteq_{\text{MGC}_n} L_n^\#$ for $\text{MGC}_n = \text{MGC}[L^\#]$. Hence, the abstraction theorem ([Thm. 7.3](#)) ensures that $L_n \sqsubseteq_{\text{MGC}[L]} L_n^\#$. This implies that $L \uplus L_n \sqsubseteq_{\text{MGC}} L \uplus L_n^\#$. In addition, $L \sqsubseteq_{\text{MGC}'} L^\#$ also implies that $L \uplus L_n^\# \sqsubseteq_{\text{MGC}} L^\# \uplus L_n^\#$. Finally, $L \uplus L_n \sqsubseteq_{\text{MGC}} L^\# \uplus L_n^\#$ follows from $L \uplus L_n \sqsubseteq_{\text{MGC}} L \uplus L_n^\#$ and $L \uplus L_n^\# \sqsubseteq_{\text{MGC}} L^\# \uplus L_n^\#$. $\square$

A.5 Proof Sketch of Queue Library Correctness (§8.2)

THEOREM A.17. *For the queue implementation L_{queue} and specification $L_{\text{queue}}^\#$ as given in §8.2, we have $L_{\text{queue}} \sqsubseteq_{\text{MGCfree}} L_{\text{queue}}^\#$.*

PROOF. To show refinement we consider histories of $\text{MGCfree}[L_{\text{queue}}]$ and $\text{MGCfree}[L_{\text{queue}}^\#]$ which consists of labels corresponding to call, return, and the propagation of the call and return of the enqueue and dequeue methods. Since the local registers are not used for passing values to (and from) the functions, in the following discussion we ignore the bookkeeping of the register variables in the history labels. We observe the following,

Remark 5. For a given trace $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11})$ of the queue implementation, no two different W s appearing on the trace t can write to the same location in the queue. It is true because a write to the location $q[n]$ happens only after n is obtained by the most recent preceding RMW in trace t with $\text{loc}(\text{RMW}) = T$. Since, due the release acquire semantics, no two RMWs with $\text{loc}(\text{RMW}) = T$ can read from the same write, we will always get a different values for n . Therefore, any location in the queue can be written at most once. Let $v[n]$ represents the value written at the location $q[n]$ by some write on the trace t . Also, by looking at the structure of enqueue method (in specification program), it easy to see that no two different writes W s appearing on a given trace $t^\# \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$ can write to the same location in the queue. Let $v^\#[n]$ represents the value written at the location $q[n]$ by some write on the trace $t^\#$.

Remark 6. For a given trace $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11})$ of the implementation program, the value read by an RMW, with $\text{loc}(\text{RMW}) = H$ appearing in the trace t is the value written by the most recent previous RMW with $\text{loc}(\text{RMW}) = H$. On the other hand for a given trace $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$ of the specification program it is easy to see that the value read by a R with $\text{loc}(R) = H$ is same as the value written by the most recent previous write W with $\text{loc}(W) = H$.

Remark 7. A read R by thread τ , with $\text{loc}(R) = q[n]$, appearing on a given trace $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11})$ can either return $v[n]$ or $\perp$, if the prefix of t , containing this R as the last label, does not contain a propagation level of the form $\text{EP}(e, \tau, X_L)$ where e is an event corresponding to return of an enqueue method. Otherwise, such a read can only return $v[n]$. Similarly, a read R with $\text{loc}(R) = q[n]$ appearing on a given trace $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$ can either return $v^\#[n]$ or $\perp$ when the prefix of t , containing this R as the last label, does not contain a propagation level of the form $\text{EP}(e, \tau, X_L)$ where e is an event corresponding to return of an enqueue method. Otherwise, such a read can only return $v^\#[n]$.

Next we define the notion of *similar traces* to state and prove our main result. Let $L_W := \{ \langle \tau, l \rangle \mid \tau \in \text{Tid} \wedge l \in \{W(X_L, T, \cdot, \cdot), W(X_L, H, \cdot, \cdot)\} \}$ be the collection of labels corresponding to the execution of instructions `store`($X_L, T, \text{tp}+1, \text{rlx}$) and `store`($X_L, H, \text{hp}+1, \text{rlx}$) from the specification program of the queue object. Similarly, let $L_U := \{ \langle \tau, l \rangle \mid \tau \in \text{Tid} \wedge l \in \{ \text{RMW}(X_L, T, \cdot, \cdot), \text{RMW}(X_L, H, \cdot, \cdot) \} \}$ be the collection of RMWs labels corresponding to the successful execution of `CAS` instruction inside the enqueue and dequeue methods of the queue implementation. For $l_U \in L_U$ and $l_W \in L_W$, we say $l_U \approx l_W$ iff $\text{tid}(l_U) = \text{tid}(l_W) \wedge \text{loc}(l_U) = \text{loc}(l_W) \wedge \text{val}_W(l_U) = \text{val}_W(l_W)$. For the traces $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11})$ and $t^\# \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$, we say that t is *similar to* $t^\#$, denoted as $t \simeq t^\#$, if $\forall i, t|_{\text{HLab}_F \cup L_U}[i] \in \text{HLab}_F \implies t|_{\text{HLab}_F \cup L_U}[i] = t^\#|_{\text{HLab}_F \cup L_U}[i]$, $\forall i, t|_{\text{HLab}_F \cup L_U}[i] \in L_U \implies t|_{\text{HLab}_F \cup L_U}[i] \approx t^\#|_{\text{HLab}_F \cup L_U}[i]$, and $\forall n, v[n] \neq \perp \implies v[n] = v^\#[n]$.

Intuitively, if trace t is similar to trace $t^\#$, then their entries corresponding to history labels are exactly same and the entries from L_U in t is replaced with corresponding entries from L_W in $t^\#$. It is easy to see that $H_F(t) = H_F(t^\#)$ whenever $t \simeq t^\#$. Therefore, our main result follows from the following claim,

- For every $t \in \text{traces}(\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11})$ where $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t}_{\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11}} \langle \bar{p}, M \rangle$, there exists $t^\# \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$ such that $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t^\#}_{\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11}} \langle \bar{p}^\#, M^\# \rangle$, $t \simeq t^\#$, and locks are open in $M^\#$.

We prove the above claim using induction on the length of trace t . Consider the last transition label α in t , and let t' such that $t = t' \cdot \alpha$. Let $\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t'}_{\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11}} \langle \bar{p}', M' \rangle \xrightarrow{\alpha}_{\text{MGCfree}[L_{\text{queue}}] \bowtie \text{pRC11}} \langle \bar{p}, M \rangle$. By induction hypothesis, there exists $t'' \in \text{traces}(\text{MGCfree}[L_{\text{queue}}^\#] \bowtie \text{pRC11})$ such that

$\langle \bar{p}_{\text{Init}}, M_{\text{Init}} \rangle \xrightarrow{t''} \text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11} \langle \bar{p}'', M'' \rangle$ and $t' \simeq t''$. We consider the following four cases and show the existence of $t^{\#}$ such that $t \simeq t^{\#}$ in each of these cases. For all other possible values of α we let $t^{\#} := t''$ and it is easy to see that $t \simeq t''$.

- (1) $\alpha = \text{EP}(e, \tau, X_L)$ and $\text{typ}(e) \in \{\text{CALL}, \text{RET}\}$: Let $\text{tid}(\alpha) = \tau$, then α represents the propagation of call (or return) of an enqueue (or dequeue) method invoked by thread τ . If $\text{typ}(e) = \text{CALL}$, then it is easy to see that M'' can also take this step. On the other hand if $\text{typ}(e) = \text{RET}$, then $\text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11}$ can take the α step after propagating all the events on thread τ between the call and return of the method whose return is currently being propagated. Let $t^{\#}$ be the resulting trace of $\text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11}$, it is easy to see that $t \simeq t^{\#}$ and the locks are still open in $M^{\#}$.
- (2) $\text{typ}(\alpha) = \text{CALL}$: Let $\text{tid}(\alpha) = \tau$, then α signifies that the last executed instruction is either a call to enqueue or dequeue by thread τ . It is easy to see that thread τ of $\text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11}$ can also take an α step and $t \simeq t'' \cdot \alpha$ while the locks are still open in $M^{\#}$.
- (3) $\text{typ}(\alpha) = \text{RMW}$: Let $\text{tid}(\alpha) = \tau$, then α signifies that the last executed instruction is either a $\text{tp}' := \text{CAS}(X_L, T, \text{tp}, \text{tp} + 1, \text{rlx}, \text{rlx}, \text{rlx})$ or a $\text{hp}' := \text{CAS}(X_L, H, \text{hp}, \text{hp} + 1, \text{rlx}, \text{rlx}, \text{rlx})$ by thread τ . We consider the $\text{tp}' := \text{CAS}(X_L, T, \text{tp}, \text{tp} + 1, \text{rlx}, \text{rlx}, \text{rlx})$ step taken by the enqueue implementation. Note that due to induction hypothesis t'' already has a label corresponding to the call of the enqueue method by thread τ in $\text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11}$. We obtain $t^{\#}$ from t'' by appending in sequential order the labels corresponding to instruction between `pacquire`(LT) and `prelease`(LT) (including the labels for the `pacquire` and `prelease` instructions). It is easy to see that $t \simeq t^{\#}$ and the locks remain open in $M^{\#}$. We can give a similar argument when α corresponds to the $\text{hp}' := \text{CAS}(X_L, H, \text{hp}, \text{hp} + 1, \text{rlx}, \text{rlx}, \text{rlx})$ step taken by the dequeue method of the implementation program.
- (4) $\text{typ}(\alpha) = \text{RET}$: Let $\text{tid}(\alpha) = \tau$, then α signifies that the last executed instruction is a return of enqueue (or of dequeue) by thread τ . It is easy to see that thread τ of $\text{MGCfree}[L_{\text{queue}}^{\#}] \bowtie \text{pRC11}$ can also take this α step (i.e., $t \simeq t'' \cdot \alpha$), and the locks are open in $M^{\#}$. If α is a dequeue we claim that the value returned will be same. If the value returned, by the dequeue of the implemented library L_{queue} , is $\perp$ then due to [Remark 7](#) and induction hypothesis, the dequeue of the specification $L_{\text{queue}}^{\#}$ can return $\perp$. On the other hand, if the value returned is $v[n] \neq \perp$, then n is the value returned by the most recent RMW on the location H . In this case, the dequeue of implementation would have executed the instructions $\text{hp} := \text{load}(X_L, H, \text{rlx})$ and $\text{hc} := \text{load}(X_L, q[\text{hp}], \text{rlx})$. Therefore, due to [Remark 5](#), [Remark 6](#), [Remark 7](#), and the induction hypothesis the dequeue of specification can return the value $v^{\#}[n]$. In the present case, when $v[n] \neq \perp$, by induction hypothesis we know that $v[n] = v^{\#}[n]$ which completes our proof. $\square$

Received 2022-07-07; accepted 2022-11-07