

**Optimal Tracing and Replay for Debugging
Message-Passing Parallel Programs**

Robert H. B. Netzer and Barton P. Miller

Department of Computer Science
Brown University
Providence, Rhode Island 02912

CS-94-32
July 1994

Optimal Tracing and Replay for Debugging Message-Passing Parallel Programs

Robert H. B. Netzer
Dept. of Computer Science
Brown University
Box 1910
Providence, RI 02912
rn@cs.brown.edu

Barton P. Miller
Computer Sciences Dept.
University of Wisconsin–Madison
1210 W. Dayton St.
Madison, WI 53706
bart@cs.wisc.edu

Abstract

A common debugging strategy involves re-executing a program (on a given input) over and over, each time gaining more information about bugs. Such techniques can fail on message-passing parallel programs. Because of nondeterminacy, different runs on the given input may produce different results. This non-repeatability is a serious debugging problem, since an execution cannot always be reproduced to track down bugs. This paper presents a technique for tracing and replaying message-passing programs. By tracing the order in which messages are delivered, a reexecution can be forced to deliver messages in their original order, reproducing the original execution. To reduce the overhead of such a scheme, we show that the delivery order of only messages involved in *races* need be traced (and not every message). Our technique makes run-time decisions to detect and trace racing messages, and is usually *optimal* in the sense that the minimal number of racing messages is traced. Experiments indicate that only 1% of the messages are often traced, gaining two orders of magnitude reduction over traditional techniques which trace every message. These traces allow an execution to be reproduced any number of times for debugging. Our work is novel in that we adaptively decide what to trace, and trace only those messages that introduce nondeterminacy. With our strategy, large reductions in trace size allow long-running programs to be replayed that were previously unmanageable. In addition, the reduced tracing requirements alleviate tracing bottlenecks, allowing executions to be debugged with substantially lower execution-time overhead.

Keywords: debugging, deterministic re-execution, nondeterminism, tracing, message-passing parallel programs.

1. Introduction

Message-passing parallel programs can be nondeterministic. Variations in scheduling and message latencies can cause two executions of the same program (on the same input) to produce different results. Such nondeterminacy may be intended, but it can cause serious problems while debugging: subsequent executions of the program may not reproduce the original bug. For example, an execution that core dumps may not be reproducible even after 1000 subsequent runs on the same input. This non-repeatability makes it difficult to use traditional sequential debugging techniques that require repeated execution. In this paper we present a mechanism for tracing the message delivery order so an execution can be repeatedly replayed. Replay is achieved by forcing each process during reexecution to receive messages in the same order as in the original execution. A critical cost in such a mechanism is the cost of tracing the delivery order. To reduce this cost, our scheme *adaptively* makes tracing decisions at run-time to trace only those messages that introduce nondeterminacy (and whose deliveries must be enforced to achieve replay), instead of tracing every message. Experiments show that only 1% of the messages are usually traced, improving by up to two orders of magnitude earlier techniques that trace every message. With such a reduction, long-running programs can now be replayed that could not have been previously replayed.

In a trace-and-replay scheme, the order in which messages are delivered (but not their contents) is first traced during execution. These traces are then used during re-execution to force each message to be delivered to the same operation as during the traced execution. Tracing the original execution is necessary because some messages may *race* with others, introducing nondeterminacy into the execution. Two messages race if they are simultaneously in transit and either could arrive first at some receive operation. If the original order in which racing messages are delivered is not recorded, their order cannot always be reproduced during replay. However, by tracing the original message deliveries and then forcing them to occur during replay, the computation and all its messages will be exactly reproduced¹. An erroneous execution can then be repeatedly replayed to analyze the execution more carefully and gain information about bugs.

Our main result is an adaptive tracing algorithm based on a proof that only racing messages need be traced to support replay. Our algorithm detects racing messages on-the-fly, and is *optimal* in most cases in the sense that only one message out of each pair of racing messages is traced (if any fewer messages were traced, some race would remain untraced, and insufficient information would exist to force a deterministic replay). Instead of tracing every message (as earlier schemes propose), our technique checks each message to determine if it races with another, and traces one of the racing messages. When a message is received, a race check is performed by analyzing the execution order between the previous receive operation in the same process and the message sender. The ordering information necessary for this check is maintained during execution by appending vector timestamps onto

¹ Interactions with the external environment must also be reproduced (such as return values from system calls). However, these interactions must be reproduced to replay sequential programs as well.

user messages. When a race is detected, the logical time of the message receipt is traced. To achieve replay, the traced logical times are appended onto messages during reexecution and each process is allowed to receive a message only at a time specified by the trace. This strategy is effective because the racing messages are exactly those that introduce nondeterminacy into the execution.

Our work is novel in that only the racing messages are traced. In contrast, earlier trace and replay schemes for message-passing programs require tracing every message[6, 2, 12, 10, 7, 4]. Replay was first introduced by Curtis and Wittie in the BugNet system for debugging distributed C programs[2]. LeBlanc and Mellor-Crummey[6] also addressed replay but considered both shared-memory and message-passing parallel programs. They trace only the order in which messages are delivered (and not their contents). By reproducing only the order of message deliveries, their contents (and hence the original computation) will also be reproduced. However, both of these schemes require emitting some type of trace for every message. Tracing every message can require huge amounts of storage for long-running programs, making them impossible to debug. In addition, as processors become faster and parallel machines become larger, tracing becomes an increasing bottleneck.

2. Example

To contrast traditional trace and replay schemes with our tracing strategy, we present an example message-passing program. This example shows that tracing every message sent during execution is sufficient to provide replay but is not necessary. Instead, tracing only the racing messages is sufficient to provide correct replay of all messages (even those that do not race).

Figure 1a shows a three-process message-passing program. Processes P_1 and P_3 send `Msg1` and `Msg2` to process P_2 . Process P_2 issues two `RECV` operations that will accept messages from any process. Figure 1b illustrates one possible execution of this program in which P_2 first received `Msg1` sent by P_1 , then received `Msg2` from P_3 . However, because these two messages *race*, they are not guaranteed to be delivered in this order. Intuitively, two messages race if either could be received first (due to the unpredictability of schedulers and message delays). For example, if `Msg1` were delayed (because of variations in message latencies), `Msg2` could instead be received first by P_2 , as shown in Figure 1c. This nondeterminacy causes a problem when debugging, since re-executing the program (on the same input) is not guaranteed to reproduce the original execution.

To replay the execution for debugging, we must first trace the order in which the messages are delivered, and then use this trace to force a re-execution to exhibit the same message deliveries. Earlier trace and replay schemes propose tracing the order in which *all* messages are delivered[6, 2, 12, 10, 7, 4]. For example, they would record that `Msg1` was delivered to the first `RECV` in P_2 and that `Msg2` was delivered to the second `RECV`. During replay, the receive operations are modified to accept only the appropriate messages. However, in this example, it suffices to trace only one of the two messages. If only the delivery of `Msg1` to the first `RECV` in P_2 is recorded, sufficient information still exists for replay. By forcing *only* this message to be delivered to the appropriate receive (the first `RECV` in P_2), the other message will automatically be delivered to the correct operation — it has no

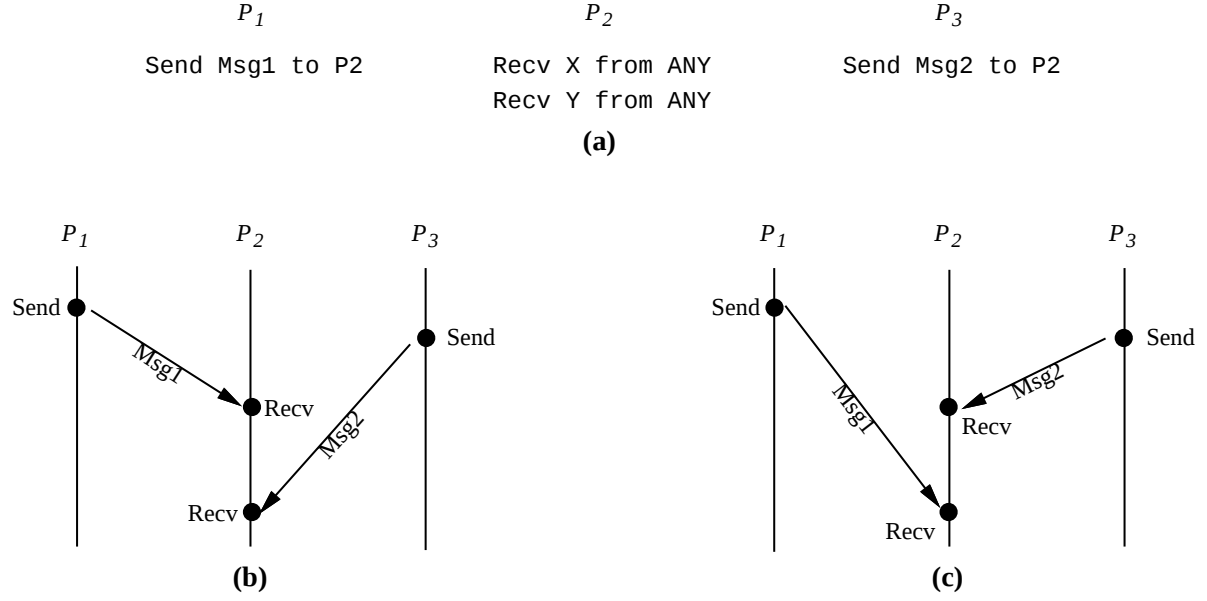


Figure 1. (a) example message-passing program, and (b),(c) two possible executions

where else to go. One of our results is a proof that only racing messages need be traced. Non-racing messages cannot introduce nondeterminacy and thus their deliveries need not be enforced during replay.

In Section 3 we formally define these races. In Section 4 we show how to detect and trace them on-the-fly and provide replay from the traces. We also prove that in the common case our strategy traces the minimal number of racing messages whose order must be reproduced for replay. In Section 5 we present experimental results indicating that this strategy is effective in practice, even in the non-optimal case.

3. Formal Definition of Race

To formally define a race, we first present a model for reasoning about executions of a message-passing parallel program. One part of this model is a notation to represent the *actual* behavior exhibited by the execution. The other part of the model characterizes *potential* behavior (alternative executions that could have occurred instead). We then use the model to characterize what we mean by a race, and later to prove that our algorithm traces enough information about the actual execution to ensure deterministic replay.

3.1. Representing the Actual Behavior

An *actual program execution* represents one execution of a message-passing program, and is a pair, $P = \langle E, \xrightarrow{HB} \rangle$, where E is a finite set of *events* and $\xrightarrow{HB}$ is the *happened-before relation* defined over E [5].

We assume that an execution consists of a fixed number of processes, each of which performs a sequence of events. We distinguish between two types of events, *computation* and *synchronization*. A computation event simply represents all computation performed in a process between synchronization operations. A synchronization event is an execution instance of a send or receive operation (we assume send and receive are the only types of synchronization operations). We model message passing as occurring over *logical channels*, and assume that each send or receive event e specifies a set of logical channels (denoted $SEND(e)$ or $RECEIVE(e)$, respectively) over which it operates. For a send event e , identical copies of the message are sent over each channel in $SEND(e)$. For a receive event e , a single message is received over any channel in $RECEIVE(e)$ (and this channel is nondeterministically chosen if more than one channel has a message available). Without loss of generality, we assume that any message sent over a channel is received by exactly one receive event (all messages are eventually received). Modeling message passing with logical channels is very general; any message-passing scheme (such as ports, mailboxes, or links) can be represented.

The happened-before relation, $\xrightarrow{HB}$, shows the relative order in which events execute and how they potentially affect one another[5]. This relation is defined as the irreflexive transitive closure of the union of two other relations: $\xrightarrow{HB} = (\xrightarrow{XO} \cup \xrightarrow{M})^+$. The $\xrightarrow{XO}$ relation shows the order in which events in the same process execute. The i^{th} event in any process p (denoted $e_{p,i}$) always executes before the $i + 1^{\text{st}}$ event: $e_{p,i} \xrightarrow{XO} e_{p,i+1}$. The $\xrightarrow{M}$ relation shows the order in which messages are delivered: $a \xrightarrow{M} b$ means that a sent a message that b received (we also write $a \xrightarrow{M} b$ to denote the message a sent). An event a is said to happen before an event b iff a could affect b because they belong to the same process or because a sequence of messages was sent from a (or a following event) to b (or a preceding event).

3.2. Representing the Potential Behavior

Races cause nondeterminacy, meaning that a program execution different from the actual execution *could have* occurred. The second part of our model characterizes a set of such alternative executions. We call an execution of the program that had the potential of occurring a *feasible* program execution, denoted $P' = \langle E', \xrightarrow{HB'} \rangle$ (and when discussing feasible executions we implicitly assume that the input is fixed). Below we define a set of feasible executions appropriate for defining what we mean by a race.

First, we consider where a race might have allowed the execution to differ from what actually occurred. The point at which the actual execution could have differed is always a receive event at which either of two (or more) messages *could have* arrived (a *racing* receive event). As a mechanism to determine if a given event, r , is a racing receive, we consider whether it is possible to construct a *frontier* across the actual program execution P in the fol-

lowing way. A frontier divides the events E into two sets: those before the frontier and those after the frontier (Figure 2 shows an example frontier just above events a , r , and c). A receive event r is a racing receive iff we can construct a frontier where

- (1) the receive event r is just after the frontier (such as in Figure 2),
- (2) two send events are just after the frontier (e.g., events a and c in Figure 2) and they each send over at least one channel in $RECEIVE(r)$, and
- (3) all receive events before the frontier also have their senders before the frontier:

for all receive events y before the frontier, $x \xrightarrow{M} y \Rightarrow x$ is also before the frontier.

If such a frontier exists, then r could have accepted either of the messages in condition (2). Condition (2) states that the senders must send on a channel over which the receive is accepting messages. Condition (3) is required for the frontier to be a *consistent cut*[1], which means that it represents a state at which all processes could have simultaneously arrived. Being able to draw a frontier means that the send events *could have* executed at the same time (thus sending messages that are simultaneously in the network), and that the unpredictability of message latencies could have allowed either message to arrive at r first. For example, the frontier in Figure 2 means that the messages sent by a and c could have been in flight simultaneously, so the message sent by c could have arrived at r first (Figure 2b).

Next, we define a set F of feasible executions by considering all the frontiers that exist in the actual program execution P . For each possible frontier f , a program execution P'_f is in this set if it is identical to P up to the frontier, and then executes the send and receive events just after the frontier. Including all such executions in the set shows all the possible ways in which different messages could have been delivered to the receive event just after each frontier. Intuitively, this set shows us where in P nondeterminacy is introduced, and which alternative message deliveries could have occurred at those points.

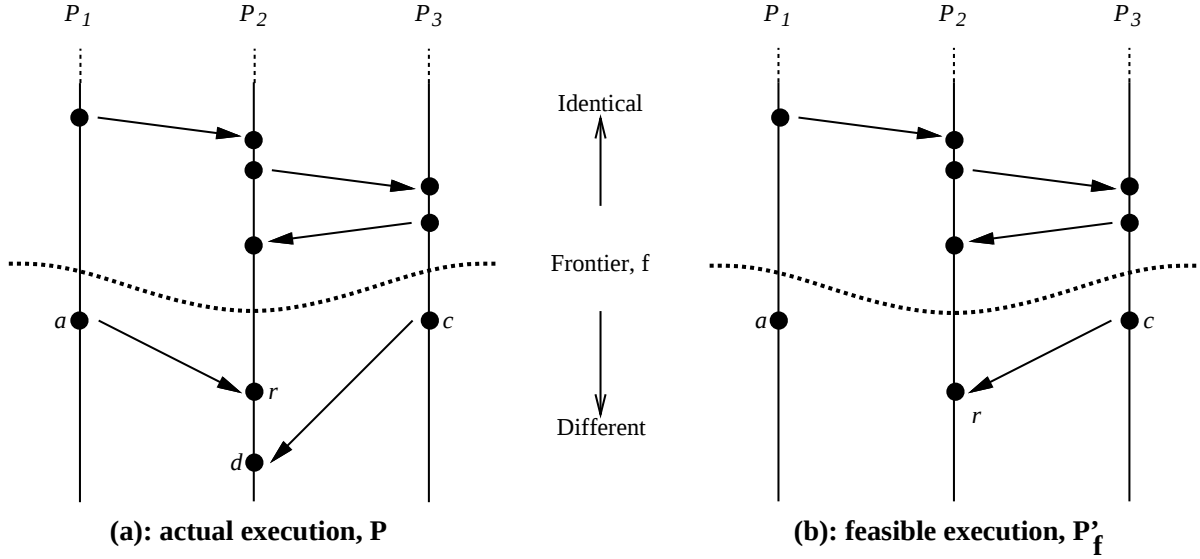


Figure 2. (a) a frontier f in actual program execution, and (b) a feasible execution P'_f in F . They are identical up to the frontier, beyond which each delivers a different message to r .

Definition 3.1

Let $P = \langle E, \xrightarrow{\text{HB}} \rangle$ be the actual program execution. Let f be any frontier that can be drawn across P as described above, and let r be the receive event, and a, b be the send events, just after the frontier. Then $P'_f = \langle E', \xrightarrow{\text{HB}'} \rangle$ is a program execution where

- (1) P'_f represents an execution the program could actually perform (i.e., a feasible execution),
- (2) P'_f is identical to P up to the frontier f :
 - (a) each process in P'_f performs the same events as the corresponding process in P before the frontier,
 - (b) P'_f exhibits the same message deliveries as P before the frontier:

for all $a, b \in E'$ where a and b are before the frontier, $a \xrightarrow{M'} b \Leftrightarrow a \xrightarrow{M} b$, and
- (3) after the frontier P'_f contains only a, b , and r ; and r receives the message sent by either a or b .

We define F to be the set of all such feasible executions P'_f for all possible frontiers f .

Figure 2b illustrates an example P'_f that belongs to F . Note that the message sent by a is not received by an event included in P'_f (since the only events we include after the frontier are a, r , and c); P'_f is only meant to show that

the message sent by c could have been received by r .

3.3. Definition of Race

If we look at a particular frontier, we will find a receive event that, during the actual execution P , could have received any of several messages. It is these messages that we wish to define as racing, and we say that they are involved in a *frontier race*. By considering all the possible frontiers, we define a binary relation over the messages in P to denote all the frontier races that exist in P .

Definition 3.2

Let P be the actual program execution. We say that $a \xrightarrow{M} r$ *RacesWith* $c \xrightarrow{M} d$ iff a program execution $P'_f = \langle E', \xrightarrow{HB'} \rangle$ exists in F such that $a, c, r \in E'$ and $c \xrightarrow{M'} r$ ($c \neq a$).

We are interested in frontier races because they show the messages that must be traced to provide a deterministic replay that reproduces P . If we trace enough information about the frontier races in P so that during a replay P' we can force these races to be resolved in the same order as in P , then P' will be a deterministic execution, and therefore must be identical to P (since determinacy means P' has no way to differ from P).

Theorem 1 (Replay Theorem)

A program execution that has no frontier races is deterministic.

Proofs of theorems appear in the appendix.

4. Message Tracing Using On-the-fly Race Detection

We now present our trace and replay strategy. Our approach is to locate the frontier races on-the-fly and to trace the second message involved in each detected race, instead of tracing every message. In this section we first show how the frontier races can be detected and traced on-the-fly (we discuss an implementation in the next section), and then discuss how to provide replay from the traces. We also prove that our algorithm is optimal in most cases in the sense that it traces only one message in each race.

4.1. On-the-fly Race Detection and Tracing

We detect frontier races on-the-fly by performing a race check after each receive. By analyzing the execution order between the sender and a previous receive in the same process, we can determine whether the received message races with another, and trace only a racing message. For simplicity, we assume that the receiving ends of logical channels are associated with a *single* process; e.g., messages to ports (but not mailboxes). Two messages can then race only if they are received by the same process, simplifying the tracing algorithm. Below we discuss handling more general (mailbox) communication.

Receive *Msg* from Channels:

```

1:   Send = event that sent Msg;
2:   PrevRecv = previous event (in the same process) such that RECEIVE(PrevRecv) contains
      the channel over which Msg was sent;
3:   PrevSend = event that sent a message to PrevRecv;
4:   if (PrevRecv  $\xrightarrow{HB}$  Send)
5:     /* frontier race detected */
6:     trace that a message was delivered from Send to Recv;
7:   /* else no frontier race yet detected */

```

(a)

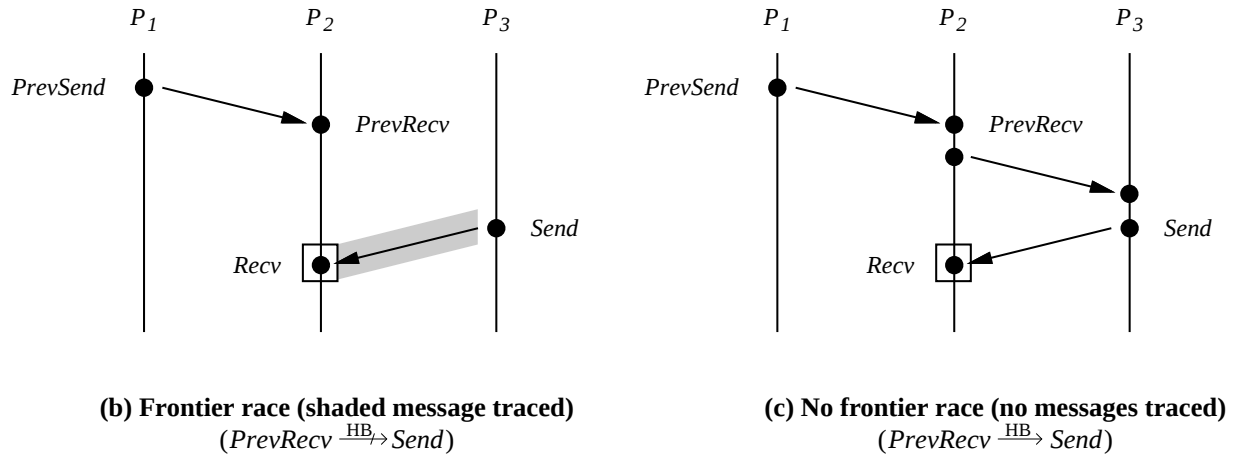


Figure 3. (a) tracing algorithm, (b),(c) example race checks performed at boxed receive

Figure 3a shows our on-the-fly race detection and tracing algorithm, which is invoked after each receive. Recall that a race exists when either of two messages could have arrived first at some receive. After a message is received, this algorithm determines whether the message could have instead been received by a previous event in the same process. To identify these situations, an earlier receive event is located (line 2) that specified a logical channel over which the current message was sent. Both the message accepted by this earlier receive and the current message are race candidates. As shown in Figure 3b, if *PrevRecv* did not happen before the sender of the current message (*Send*), then a frontier race exists — a frontier can be drawn just before *PrevSend*, *Send*, and *PrevRecv*. Both the previous and current messages could have been simultaneously in transit and either could have arrived first at *PrevRecv*. In this case the algorithm traces the second racing message. If instead *PrevRecv* happened before *Send* (as shown in Figure 3c), then no race exists: no frontier can be drawn as above (the two messages could not have been simultaneously in transit), and the algorithm emits no trace. We prove in the appendix (Theorem 2) that this algorithm traces *at least* one message in each frontier race.

The traces only need identify the sending and receiving events of the traced message. These events can be identified by maintaining in each process a local counter (incremented after every synchronization operation) that is used to assign serial numbers to events[6]. It suffices to trace the event serial numbers of the sender and receiver and the process number of the sender. If one trace file is maintained for each process in the program execution, the process number of the receiver is implicit and need not be recorded.

Because we assume that the receiving end of each logical channel is associated with a single process, to find races it suffices (in line 2) to locate the previous event in the *same* process that could have accepted the incoming message. In mailbox communication, a mailbox might have multiple simultaneous owners in different processes. Two messages (to the same mailbox) can race even if they are received by different processes. Detecting these races requires locating earlier events in any process that could have accepted the incoming message sent to the mailbox. Such events can be located by modifying the mailbox mechanism to store the last event in each process that received a message from the mailbox. Line 4 of the tracing algorithm can be modified to check the execution order of each of these events against the sender.

4.2. Replay

During replay each frontier race must be resolved in the same way as during the original execution. Theorem 1 proved that a replay free of frontier races is deterministic. To achieve a deterministic replay, only the delivery of the traced messages need be specially enforced; untraced messages will automatically arrive at the correct receive event and need no special treatment. We provide replay by tagging and buffering the racing messages so they can be accepted by receive events in the proper order. If necessary, the traced messages can be forced (at additional cost) to also arrive in their original order, alleviating the need for buffering and ensuring that no buffer overflows occur during replay.

To effect replay, the trace files must first be collated; a racing message is traced during execution when it is *received*, but tagging the message during replay requires a special action when it is *sent*. To produce the collated trace file for process p , all traces of messages sent from p must be collected (from the uncollated trace files) into a single file and sorted by sender serial number. During replay, as in the original execution, serial numbers must be assigned to events by maintaining a local counter that is incremented after every synchronization operation. This counter is used to ensure that racing messages are accepted by the intended receive. Before each send, the serial number of the next racing message (read from the trace file) is compared to the current value of the local counter. Because the trace file is sorted, these numbers will match if the message about to be sent was originally involved in a race. If the message to be sent originally raced, it is tagged with the serial number of its intended receive. If the message did not race, it is not specially tagged. Receives are modified to accept tagged messages only if their serial numbers match those on the tagged messages. Tagged messages with serial numbers that do not match are buffered so they can be accepted by later receives. Such buffering is often normally performed by message-passing systems that accept asynchronous messages. Untagged messages are accepted as usual.

The above strategy ensures that each message is *received* by the correct event, but does not guarantee that messages *arrive* at a process in the same order as during the original execution. Racing messages can still arrive in any order, and must be buffered so that they can be received in the correct order. Buffering is normally not a problem unless buffer space is limited. Buffer overflows may occur during replay that did not originally occur. To guarantee that no overflows occur, we can reproduce the original message arrival order by introducing *control* messages during replay. When a message $a \xrightarrow{M} b$ races with a message $c \xrightarrow{M} d$, a control message from a new send (added by the replay system) just after b to a new receive just before c ensures that $c \xrightarrow{M} d$ is not sent until $a \xrightarrow{M} b$ has been received. Figure 4 shows an example ($x \xrightarrow{M} y$ is the control message).

The control messages can be passed over any logical channel that exists between the two processes. They introduce orderings that remove all frontier races between user (i.e., non-control) messages. As a result, user messages will be delivered to each process in exactly the same order as during the original execution. In addition, at most one user message will ever be in transit over a channel², so each channel only need buffer at most one user message. The control messages may race among themselves, but such races are benign since control messages have no content and can be received in any order. In addition, channels need not be FIFO for this strategy to work, as the control messages effectively force FIFO delivery (since at most one message is ever in transit over a chan-

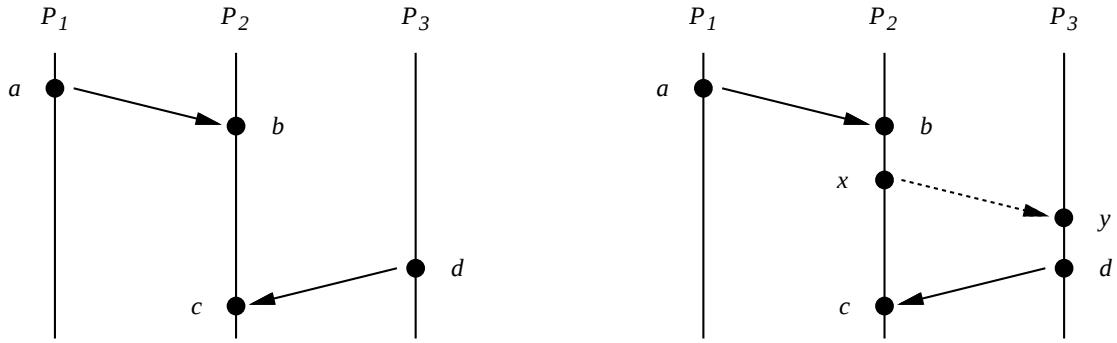


Figure 4. (a) a race between $a \xrightarrow{M} b$ and $c \xrightarrow{M} d$, and (b) the control message $x \xrightarrow{M} y$ added for replay

² A simple argument shows this: For a contradiction, assume that two messages $a \xrightarrow{M} r_1$ and $b \xrightarrow{M} r_2$ are both simultaneously in transit over some channel. The control messages that prevent these messages from racing cause either $r_1 \xrightarrow{M} b$ or $r_2 \xrightarrow{M} a$. But this means that one of the messages is received *before* the other is sent. Therefore they cannot be simultaneously in transit.

nel). Thus, at most one control message can ever be in transit from a process at any time, and thus at most $P - 1$ control messages (where P is the number of processes) will ever need to be buffered over any channel. If enough buffer space is reserved for this many control messages (which are small), buffer overflows will never occur during replay. However, because control messages introduce additional orderings that were not present during the original execution, they can reduce the amount of parallelism achievable during replay, and should only be used if buffer overflow is otherwise a problem.

4.3. Optimal Message Tracing

We now show one common case where our tracing algorithm is *optimal*. Since something must be traced about each racing message, we consider an optimal trace to be one where *only* one message in each race is traced (if any fewer messages were traced, there would be some untraced race that would cause replay to be nondeterministic). We characterize when the traces are optimal and present example executions for which optimal and non-optimal traces are generated. As shown later, even when non-optimal traces are recorded, they are usually small.

Our tracing algorithm traces only one message of each race if each message is either involved in only one race, or if messages participate in multiple races and the races are *transitive*. Transitive races often occur when receive operations specify either a single channel over which to accept a message, or *all* channels, instead of a subset of channels. Theorem 4 (in the appendix) proves that on executions for which the *RacesWith* relation is transitive (see Definition 3.2), our tracing algorithm is optimal. Figure 5a illustrates such an execution. Because all three receive events could have accepted messages from any channel, all three messages, $a \xrightarrow{M} b$, $c \xrightarrow{M} d$, and $e \xrightarrow{M} f$, race with each other. The optimal trace consists of recording any two of the three messages (such as the two shaded messages): without tracing these messages, replay cannot ensure that all three messages are delivered to the correct event. In this case, our algorithm traces the last two messages, which is an optimal trace.

Figure 5b shows an execution with non-transitive races; the first receive can accept messages only from processes 1 and 3, the second receive from processes 3 and 4, and the last receive from process 4. Because the second message could have been accepted by the first receive, the first two messages race. Similarly, because the third message could have been accepted by the second receive, the last two messages race. However, because the last message could not have been received by the first receive (it only accepts messages from processes 1 and 3), the first and last messages do *not* race. The optimal trace thus consists of recording *only* the second message: if replay ensures that the second message is delivered to the second receive operation, the other two messages will automatically arrive at the correct events (no other receives will accept them). Our tracing algorithm would trace the last two messages, which is not optimal.

5. Implementation and Experimental Results

We now discuss experiences with our trace and replay strategy. We first discuss an implementation of our tracing algorithm based on appending *vector timestamps* onto user messages. These timestamps provide informa-

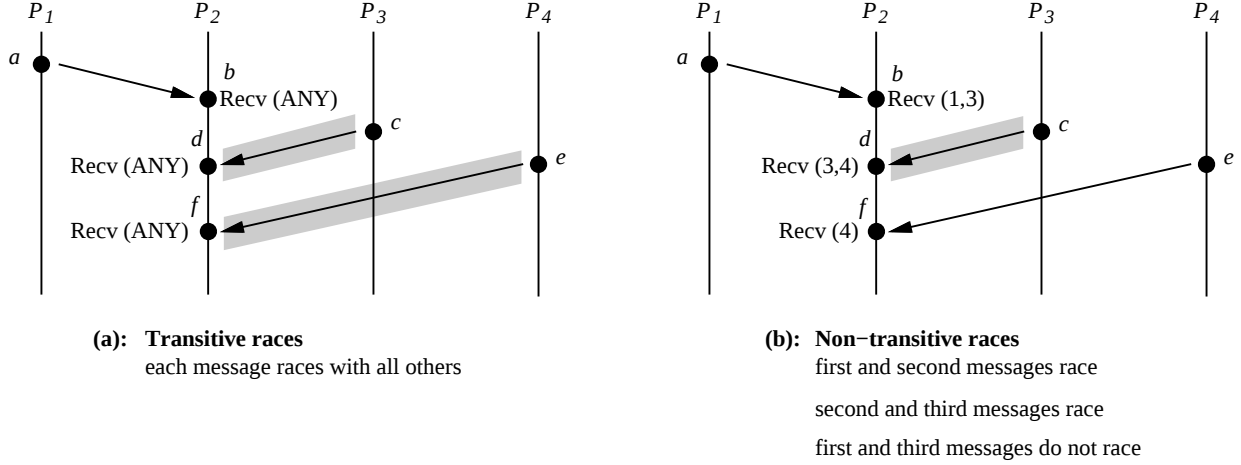


Figure 5. (a) Transitive races, (b) non-transitive races: shaded messages show the optimal trace (tracing algorithm would trace the last two messages)

tion about the $\xrightarrow{HB}$ relation needed to perform on-the-fly race checks. We then discuss experiments performed on a collection of message-passing programs on a 64-node Thinking Machines CM-5 and a 32-node Intel iPSC/2 hypercube. These experiments show that only 0 – 19% of the messages in these programs were traced, and in all but one case the optimal trace was generated. In addition, the small traces completely alleviated tracing bottlenecks that plague traditional schemes which trace every message. These results suggest that our trace and replay technique is very effective in practice, producing small traces with low execution-time overhead, providing a new technique to debug even long-running programs. We end by discussing how tracing overhead can perturb the execution, possibly causing behavior different than an untraced execution.

5.1. Implementation

The tracing algorithm presented in Section 4 detects races by determining the $\xrightarrow{HB}$ relation between the sender of a message and a previous receive. Our implementation of this algorithm uses a standard method of maintaining the $\xrightarrow{HB}$ relation during execution by keeping a *vector timestamp* in each process. A vector timestamp is a vector of length p (the number of processes) containing event serial numbers[3]. These timestamps are maintained by appending them onto user messages and updating them after each receive operation. The tracing algorithm detects races by comparing timestamp values and event serial numbers to determine whether the previous receive happened before the sender of the current message.

In our implementation, each process maintains both a local virtual clock, *Clock*, and a vector timestamp, *Timestamp*. The local clock is used to assign serial numbers to events: events are numbered sequentially within a process beginning with the number 1, and the clock is incremented after each operation. The timestamps are maintained so that at any point during execution, the i^{th} slot of the vector timestamp for process p (i.e., *Timestamp*[i]) equals the serial number of the last event in process i that happened before the most recent event in process p . By definition, the p^{th} slot equals the current value of p 's local clock. To maintain these timestamps, each process appends the current value of its timestamp onto the end of each message it sends. Upon receiving a message, it updates its timestamp by computing the component-wise maximum with the timestamp appended to the incoming message. Using vector timestamps to track the $\xrightarrow{\text{HB}}$ relation has the advantage that they work even if channels are not FIFO, and they do not require processes to synchronize their clocks[11].

The race check in line 4 of our tracing algorithm (Figure 4) is performed easily using the timestamps. The sender's timestamp (which is appended to the incoming message) is compared to the serial number of the previous receive to determine if the receive happened before the sender. The value of the p^{th} slot of the sender's timestamp equals the serial number of the most recent event in process p that happened before the sender. If the serial number of the previous receive is greater than this value, then the previous receive did not happen before the sender, and a frontier race exists.

5.2. Experimental Results

We implemented our tracing algorithm on two message-passing parallel machines: a 64-node Thinking Machines CM-5 and a 32-node Intel iPSC/2 hypercube. On each machine, two instrumented versions of the message-passing library were created. One version uses the traditional approach of tracing every message sent during execution, and the other version uses our tracing algorithm to trace only racing messages. We analyzed a collection of message-passing programs obtained from colleagues and measured two quantities. First, the percentage of messages that race was recorded. This percentage shows the trace size reduction obtained by our race-based tracing strategy. Second, the increase in execution time of both the traditional approach of tracing every message and our approach of tracing only racing messages was measured. These overheads show whether the cost of performing race checks outweighs the time savings obtained by not tracing non-racing messages. We found that often only 0 – 2% of the total messages were traced, and in cases where the execution-time overhead of tracing every message is high, race-based tracing is an order of magnitude faster.

Table 1 shows the results of our experiments on six programs. *det* computes the determinant of a matrix, and was run on a randomly generated 100×100 matrix. *line* computes the intersections of a collection of line segments, and was run on 1000 randomly generated segments. *mesh* computes finite differences over a grid to solve a differential equation, and was run on a 300×300 grid. *mult* multiplies two matrices, and was run on two randomly generated 100×100 matrices. *sys* uses Gaussian elimination to solve a linear system of equations, and was run on a

Program	Messages Sent	Messages Traced	% of Optimal Trace	Run-time Overhead	
				for tracing all msgs	for tracing racing msgs
det‡	4713	63 (1%)	optimal	568%	8%
line‡	31	0 (0%)	optimal	0.3%	8%
mesh‡	10210	1392 (14%)	optimal	28%	14%
mult‡	1120	0 (0%)	optimal	15%	0.1%
sys‡	9424	332 (2%)	optimal	561%	14%
tycho†	1791	412 (19%)	within 46%	—	—

† 64-node Thinking Machines CM-5

‡ 32-node Intel iPSC/2

Table 1. Results of message tracing

system of 300 randomly generated (linearly independent) equations. *tycho* is a cache simulator, and was run on a 10 MByte address trace.

Our first experimental result pertains to trace sizes: In two programs (*det* and *sys*), only 1 – 2% of the messages were traced — a two order of magnitude reduction over tracing every message. In two programs (*line*, *mult*), *none* of the messages raced, and no traces at all were generated. Executions of these programs (on the given input) are guaranteed to be reproduced automatically; nothing special need be done during replay. In two programs (*mesh*, *tycho*), 14 – 19% of the messages were traced. These cases represent programs that were designed to be highly internally nondeterministic (although their final results are deterministic); they use some form of a first-come first-served worker paradigm. Even in such cases, the number of racing messages was rather low. In five of the six programs the optimal trace was generated (because the races were transitive, as discussed in Section 4.3). In the other program (*tycho*), the trace size was within 46% (or less) of optimal³. These results suggest that our tracing strategy is effective, tracing as few as 0 – 2% of the messages, and no more than 19% of the messages

³ We derived this quantity by computing a lower bound on the optimal trace size. As shown in the proof of Theorem 3 (in Appendix A), computing an optimal trace is equivalent to computing a minimal vertex cover of a graph. To determine how the recorded trace size compares to the optimal trace, we used graph matching to estimate the optimal size to within a factor of two. For *tycho*, the recorded trace size was *no more* than 46% larger than optimal. Determining whether it was actually optimal would require computing a minimal vertex cover, which is an intractable problem in general.

even in programs that are highly nondeterministic. In addition, our tracing algorithm generated optimal traces for most of the test programs, and even when non-optimal traces were generated, they were small.

Our second result pertains to the execution-time overhead incurred by our tracing strategy. To assess this overhead, we analyzed three versions of each program: the original (uninstrumented) program, an instrumented version that traces every message, and an instrumented version that traces only racing messages. Each version of every program was executed 10 times and the average execution times were computed. The last two columns of Table 1 show the execution-time overhead incurred by tracing all messages and tracing only racing messages⁴.

Two programs (*det* and *sys*) suffered substantial slowdown (almost 600%) when using the traditional approach of tracing every message. These programs exhibited a high frequency of message operations. Tracing this high traffic introduced a bottleneck, since many trace records needed to be written in a short time. In contrast, our strategy of tracing only the racing messages reduced the tracing requirements to the point where the bottleneck was completely alleviated (resulting in a slowdown of only 8 – 14%). In the programs that had no racing messages (*line*, *mult*), the overhead of race-based tracing indicates the inherent cost of performing the on-the-fly race checks. This overhead shows that the cost of maintaining the vector timestamps and checking for races is low (0.1-8%). Because *line* had low message-passing traffic, tracing every message was cheaper than maintaining the timestamps. In *mult*, longer messages were passed between nodes, making the incremental cost of appending timestamps low. In both programs, tracing every message did not introduce a bottleneck, and both tracing strategies had reasonable overheads.

5.3. Perturbations Caused by Tracing

Although the overhead of our tracing strategy is low, it is non-zero. If the execution is nondeterministic (contains races), this overhead can potentially perturb the actual execution enough for it to exhibit behavior different than it would without tracing. Such *probe effect* is a drawback shared among all schemes that trace or monitor nondeterministic programs. In the worst case it appears impossible to completely eliminate any probe effect. Our test programs were very sensitive to variations in timing as they exhibited a different execution almost every time they were run, even when not instrumented. For such programs, the probe effect is probably great, although it is unclear if executions result that are *completely* different than what would otherwise normally occur. In practice, keeping the tracing overhead as low as possible can reduce the chance that such a perturbation will occur. Our adaptive tracing strategy contributes to this goal. In addition, current work is exploring additional analyses that can detect where in the execution the instrumentation may have changed the behavior. Dealing with the probe ef-

⁴ Because our CM-5 had not yet been equipped with I/O processors, overhead measurements for this machine were not reported. Our implementation performed trace I/O by sending messages over the diagnostic network. The unusually high cost of this I/O makes the overhead of tracing all messages several orders of magnitude higher than race-based tracing, making comparisons unrealistic.

fect in nondeterministic systems remains an issue of current research[9, 8].

6. Conclusions

In this paper we presented a trace and replay strategy for message-passing parallel programs that substantially reduces tracing overhead over past schemes. The small traces and low overhead produced allow even long-running programs to be replayed that previously could not have been replayed in practice. This work provides a new foundation on which efficient parallel-program debugging techniques can be built. We achieve these benefits by making tracing decisions at run-time, instead of tracing every message (as earlier work proposes). Race checks are performed after each receive operation to locate and trace those messages that introduce nondeterminacy. We prove that our tracing strategy is optimal (in the sense that a minimal number of racing messages is traced) for many programs which exhibit simple message patterns. Even when non-optimal traces are generated, experiments show that the traces are kept small, and are one to two orders of magnitude smaller than traces of every message.

Although the primary application of this work is replay, the on-the-fly tracing algorithm is also useful for directing the programmer to the locations of races. When the programmer intends the execution to be completely deterministic, but non-empty traces are generated, the traces can be perused to determine which messages (erroneously) raced. The traces can then be used to effect a replay where more detailed information about the races can be collected.

Future work includes a more precise characterization of when our algorithm is optimal (e.g., transitive races are sufficient for an optimal trace but not necessary). Better tracing strategies may be possible. By buffering more information about recent messages, more informed tracing decisions might trace fewer races. By employing optimizations for maintaining timestamps, the overhead of passing timestamps and performing race checks might also be reduced.

References

- [1] Ozalp Babaoglu and Keith Marzullo, “Consistent Global States of Distributed Systems: Fundamental Concepts and Mechanisms,” *Technical Report UBLCS-93-1*, University of Bologna, (January 1993).
- [2] R. Curtis and L. Wittie, “BugNet: A Debugging System for Parallel Programming Environments,” *Proc. of the 3rd Intl. Conf. on Dist. Computing Systems*, pp. 394-399 (1982).
- [3] C. J. Fidge, “Partial Orders for Parallel Debugging,” *SIGPLAN/SIGOPS Workshop on Parallel and Distributed Debugging*, pp. 183-194 Madison, WI, (May 1988). Also appears in *SIGPLAN Notices* **24**(1) (January 1989).
- [4] Arthur P. Goldberg, Ajei Gopal, Andy Lowry, and Rob Strom, “Restoring Consistent Global States of Distributed Computations,” *ACM/ONR Workshop on Parallel and Distributed Debugging*, pp. 144-154 Santa Cruz, CA, (May 1991).
- [5] Leslie Lamport, “Time, Clocks, and the Ordering of Events in a Distributed System,” *CACM* **21**(7) pp. 558-565 (July 1978).

- [6] Thomas J. LeBlanc and John M. Mellor-Crummey, “Debugging Parallel Programs with Instant Replay,” *IEEE Trans. on Computers* **C-36**(4) pp. 471-482 (April 1987).
- [7] Eric Leu, Andre Schiper, and Abdelwahab Zramdini, “Efficient Execution Replay Technique for Distributed Memory Architectures,” *2nd European Distributed Memory Computing Conference*, LNCS 487, Springer-Verlag, Munich, (1991).
- [8] James E. Lumpp, Jr., Julie A. Gannon, Mark S. Andersland, and Thomas L. Casavant, “A Technique for Recovering from Software Instrumentation Intrusion in Message-Passing Systems,” *Technical Report TR-ECE-920817*, University of Iowa Department of Electrical and Computer Engineering, (August 1992).
- [9] Allen D. Malony and Daniel A. Reed, “Models for Performance Perturbation Analysis,” *ACM/ONR Workshop on Parallel and Distributed Debugging*, pp. 15-25 Santa Cruz, CA, (May 1991).
- [10] Barton P. Miller and Jong-Deok Choi, “A Mechanism for Efficient Debugging of Parallel Programs,” *SIGPLAN Conf. on Programming Language Design and Implementation*, pp. 135-144 Atlanta, GA, (June 1988).
- [11] Reinhard Schwarz and Friedemann Mattern, “Detecting Causal Relationships in Distributed Computations: In Search of the Holy Grail,” *Technical Report SFB124-15/92, Dept. of Computer Science, Univ. of Kaiserslautern, Kaiserslautern, Germany*, (December 1992).
- [12] Kuo-Chung Tai and Sanjiv Ahuja, “Reproducible Testing of Communication Software,” *IEEE COMPSAC '87*, pp. 331-337 (1987).

Appendix. Proofs of Theorems

Theorem 1 (Replay Theorem).

A program execution that has no frontier races is deterministic.

Proof. To establish a contradiction, assume that some program execution $P = \langle E, \xrightarrow{HB} \rangle$ is nondeterministic but free of frontier races. Since P is nondeterministic, another execution of the program on the same input could produce a different execution, $P' = \langle E', \xrightarrow{HB'} \rangle$. P and P' exhibit the same events and message deliveries up to some point after which they differ. Let r be a receive event where they first differ. That is, $x \xrightarrow{HB} y \Leftrightarrow x \xrightarrow{HB'} y$ for all events x, y where $x \xrightarrow{HB} r$ and $y \xrightarrow{HB} r$. Also let s_1 and s_2 be operations that send messages to r in P and P' : $s_1 \xrightarrow{M} r$ and $s_2 \xrightarrow{M'} r$. The messages sent by s_1 and s_2 in P must race because P' meets the conditions in Definition 3.1 and thus belongs to F . The events x and y are all the events before the frontier, and s_1 , s_2 , and r are after the frontier. But P containing a frontier race contradicts the assumption. Thus, P' cannot be different than P , implying that P is deterministic. QED

Theorem 2 (Tracing Theorem).

The tracing algorithm (Figure 3) traces *at least* one message in each frontier race.

Proof. We prove below that a message is traced by the algorithm when some predicate is true (Lemma 1), and then prove that this predicate is true when a frontier race exists (Lemma 2). At least one message in each frontier race is thus traced. QED

Lemma 1.

If two messages, $Send \xrightarrow{M} Recv$ and $PSend \xrightarrow{M} PRecv$, exist such that $PRecv \xrightarrow{HB} Send \wedge PRecv \xrightarrow{XO} Recv \wedge SEND(Send) \cap RECEIVE(PRecv) \neq \emptyset$, then the tracing algorithm traces $Send \xrightarrow{M} Recv$.

Proof. To establish a contradiction, assume that the above conditions hold but the algorithm does not trace $Send \xrightarrow{M} Recv$. This message is not traced only if the algorithm finds that the previous receive, $PrevRecv$ (located in line 2), happened before the sender: $PrevRecv \xrightarrow{HB} Send$. We must also have $PRecv \xrightarrow{XO} PrevRecv$, otherwise the algorithm would find $PRecv$ as the previous receive. These orderings imply that $PRecv \xrightarrow{HB} Send$, which contradicts the assumption that $PRecv \xrightarrow{HB} Send$. QED

Lemma 2.

If $PSend \xrightarrow{M} PRecv$ RacesWith $Send \xrightarrow{M} Recv$, then $PRecv \xrightarrow{HB} Send \wedge PRecv \xrightarrow{XO} Recv \wedge SEND(Send) \cap RECEIVE(PRecv) \neq \emptyset$.

Proof. Because $PSend \xrightarrow{M} PRecv$ RacesWith $Send \xrightarrow{M} Recv$, there exists a $P' = \langle E', \xrightarrow{HB'} \rangle \in F$ such that $Send \xrightarrow{M'} PRecv$ (by Definition 3.2). We consider each term in the conjunct $PRecv \xrightarrow{HB} Send \wedge PRecv \xrightarrow{XO} Recv \wedge SEND(Send) \cap RECEIVE(PRecv) \neq \emptyset$.

- (1) To establish a contradiction, assume that $PRecv \xrightarrow{HB} Send$. In addition, by the definition of F (part (3) of Definition 3.1), $Send \xrightarrow{M'} PRecv$ implies that $\forall x \in E, x \xrightarrow{HB} Send \Leftrightarrow x \xrightarrow{HB'} Send$. Thus, if $PRecv \xrightarrow{HB} Send$, we must have $PRecv \xrightarrow{HB'} Send$, which contradicts the assumption that $Send \xrightarrow{M'} PRecv$.
- (2) To establish a contradiction, assume that $Recv \xrightarrow{XO} PRecv$. Then, P' cannot belong to F , since by its definition $Recv \xrightarrow{XO} PRecv$ implies that $Recv$ is before the frontier and thus $Send \xrightarrow{M'} Recv$ (since $Send \xrightarrow{M} Recv$), which contradicts the assumption that $Send \xrightarrow{M'} PRecv$.
- (3) Since $Send \xrightarrow{M'} PRecv$, we clearly have $SEND(Send) \cap RECEIVE(PRecv) \neq \emptyset$. QED

Theorem 3 (Tracing Complexity Theorem).

Given a program execution, $P = \langle E, \xrightarrow{HB} \rangle$, determining whether replay can be implemented by tracing k or fewer messages is an NP-hard problem.

Proof. We use a reduction from the vertex cover problem, known to be NP-complete: given an undirected graph, $G = (V, E)$, does G have a vertex cover with k or fewer vertices? A vertex cover is a subset V' of the vertices such that every edge is connected to some vertex in V' . Given a graph, G , we reduce the problem of determining whether it has a vertex cover with k or fewer vertices to the problem of determining whether a program execution, P , can be replayed from a trace of k or fewer messages.

From the graph G we construct P as follows. P contains two processes between which a message is sent for each of the n vertices in G . Process 1 in P contains n send operations, and process 2 contains n receive operations. The i^{th} send operation sends a null message over logical channel i , and the i^{th} receive operation specifies

that it will receive over logical channel i . Additional channels are specified by the receive operations so that two messages race iff an edge connects their corresponding nodes in G . For an edge from vertex i to vertex j , the i^{th} receive operation also specifies that it will receive over logical channel j . Because messages sent from process 1 may be delivered out of order, the i^{th} and j^{th} messages in P race iff an edge exists from vertex i to vertex j in G .

G has a vertex cover with k or fewer vertices iff P can be replayed from a trace of k or fewer messages. Assume that G has a vertex cover V' with k vertices. Each vertex in V' corresponds to one of the messages sent by P . P can be replayed from a trace of exactly these messages. Since two messages race iff an edge connects their corresponding vertices, a vertex cover ensures that at least one message in each race is traced. By the Replay Theorem (Theorem 2), a trace of these messages suffices for replay. Conversely, assume that P can be replayed from a trace T of k messages. T must contain at least one message in each frontier race, or else replay will not be frontier-race-free. Since two messages race in P iff an edge connects the corresponding vertices in G , the vertices corresponding to the messages in T are a vertex cover. QED

Theorem 4 (Optimality Theorem).

For any program execution, $P = \langle E, \xrightarrow{\text{HB}} \rangle$, for which the *RacesWith* relation is transitive, the tracing algorithm (Figure 3) traces a minimal number of racing messages required to implement replay.

Proof. As in Theorem 3, we can view the tracing problem as equivalent to computing a vertex cover of a graph. The program execution P defines a graph G : the messages in P define its vertices, and when two messages race an edge is drawn between the corresponding vertices. As discussed in Theorem 3, any trace sufficient for replay must cover the vertices of G . When the *RacesWith* relation is transitive, G becomes a forest of completely connected graphs. In this case, a minimal vertex cover is easily computed. For each completely connected component of n vertices, a minimal vertex cover consists of any $n - 1$ vertices. When races are transitive, our tracing algorithm traces all but one of the mutually racing messages (the first racing message is not traced), which corresponds to such a minimal vertex cover. Thus, the minimal number of racing message is traced. QED