

Runtime Instrumentation for Reactive Components

Luca Aceto  

Reykjavik University, Reykjavik, Iceland
Gran Sasso Science Institute, L'Aquila, Italy

Duncan Paul Attard  

University of Glasgow, Glasgow, UK

Adrian Francalanza  

University of Malta, Msida, Malta

Anna Ingólfssdóttir  

Reykjavik University, Reykjavik, Iceland

Abstract

Reactive software calls for instrumentation methods that uphold the reactive attributes of systems. Runtime verification imposes another demand on the instrumentation, namely that the trace event sequences it reports to monitors are *sound*—that is, they reflect actual executions of the system under scrutiny. This paper presents RIARC, a novel decentralised instrumentation algorithm for outline monitors meeting these two demands. Asynchrony in reactive software complicates the instrumentation due to potential trace event loss or reordering. RIARC overcomes these challenges using a next-hop IP routing approach to rearrange and report events soundly to monitors.

RIARC is validated in two ways. We subject its corresponding implementation to rigorous systematic testing to confirm its correctness. In addition, we assess this implementation via extensive empirical experiments, subjecting it to large realistic workloads to ascertain its reactivity. Our results show that RIARC optimises its memory and scheduler usage to maintain latency feasible for soft real-time applications. We also compare RIARC to inline and centralised monitoring, revealing that it induces comparable latency to inline monitoring in moderate concurrency settings where software performs long-running, computationally-intensive tasks, such as in Big Data stream processing.

2012 ACM Subject Classification Software and its engineering → Software verification and validation

Keywords and phrases Runtime instrumentation, decentralised monitoring, reactive systems

Digital Object Identifier [10.4230/LIPIcs.ECOOP.2024.16](https://doi.org/10.4230/LIPIcs.ECOOP.2024.16)

Related Version (*Extended*): <https://arxiv.org/abs/2406.19904> [8]

Supplementary Material *Software (Source code)*: <https://doi.org/10.5281/zenodo.10634182>

Funding This work is supported by the Reykjavik University Research Fund, the Doctoral Student Grant (No: 207055) and the MoVeMnt project (No: 217987) under the IRF, and the STARDUST project (No: EP/T014628/1) under the EPSRC.

Acknowledgements We thank our reviewers and the Artefact Evaluation Committee for their feedback. Thanks also to Keith Bugeja, Simon Fowler, Simon Gay, and Phil Trinder for their input.

1 Introduction

Modern software is generally built in terms of concurrent components that execute without relying on a global clock or shared state [86]. Instead, components interact via non-blocking messaging, creating a loosely-coupled architecture known as a *reactive system* [9, 93], which

- responds in a timely manner (is *responsive*),
- remains available in the face of failure (is *resilient*),
- reacts to inputs from users or their environment (is *message-driven*), and
- grows and shrinks to accommodate varying computational loads (is *elastic*).



© Luca Aceto, Duncan Paul Attard, Adrian Francalanza, and Anna Ingólfssdóttir;
licensed under Creative Commons License CC-BY 4.0

38th European Conference on Object-Oriented Programming (ECOOP 2024).

Editors: Jonathan Aldrich and Guido Salvaneschi; Article No. 16; pp. 16:1–16:33



Leibniz International Proceedings in Informatics

LIPICS Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

The real-world behaviour of reactive systems is hard to understand statically, and *monitoring* is used to inspect their operation at *runtime*, *e.g.* for debugging [110], security checking [61], profiling [75], resource usage analysis [35], *etc.* This paper considers runtime verification (RV), an application of monitoring used to detect whether the *current* execution of a system under scrutiny (SuS) deviates from its correct behaviour [16, 70, 21]. A RV monitor is a *sequence recogniser* [122, 100]: a state machine that incrementally analyses a *finite* fragment of the runtime information exhibited by a SuS to reach an *irrevocable* verdict (see [6, 5] for details).

Instrumentation lies at the core of runtime monitoring [69, 21, 63]. It is the mechanism by which runtime information from a SuS is extracted and reported to monitors as a stream of system events called a *trace*. Software is typically instrumented in one of two ways. Inline instrumentation, or *inlining*, modifies the SuS by injecting tracing instructions at specific joinpoints, *e.g.* using AspectJ [89] or BCEL [53]. Outline instrumentation, or *outlining*, uses an external tracing infrastructure to gather events, *e.g.* LTTng [55] or OpenJ9 [57], thereby treating the SuS as a *black box*. A key requirement setting RV apart from monitoring, *e.g.*, telemetry [84] or profiling [120, 25], is that the instrumentation must report *sound traces*.

► **Definition 1** (Sound traces). *A finite trace T is sound w.r.t. a system component P iff it is*

1. Complete. *T contains all the events exhibited by P so far, and*
2. Consistent. *The event sequence in T reflects the order the events occur locally at P .* ◀

Traces violating this soundness invariant are unfit for RV, as omitted, spurious, or out-of-sequence events incorrectly characterise the system behaviour, *nullifying* the verdicts that monitors flag [21, 51]. Reactive software imposes another requirement: that the instrumentation *safeguards* the responsive, resilient, message-driven, and elastic attributes of the SuS. This necessitates an instrumentation method which is itself *reactive*, in order to

1. not hamper the SuS by inducing unfeasible runtime overhead (is responsive),
2. permit monitors to fail independently of SuS components (is resilient),
3. react to trace events without blocking the SuS (is message-driven), and
4. grow and shrink in proportion to the size of the SuS (is elastic).

Limitations of current RV instrumentation methods State-of-the-art RV tools use instrumentation methods that do not satisfy *all* of the conditions 1–4 above. This renders them inapplicable to reactive software; see [63, tables 3 and 4] for details. Many approaches, including [23, 30, 48, 74, 109, 121, 126, 18], assume systems with a *fixed* architecture where the number of components remains constant at runtime, failing to meet condition 4. Works foregoing the assumption of a fixed system size, such as [44, 90, 59, 58, 24, 30, 67, 3], inline the SuS with monitors *statically*. Inlining monitors pre-deployment inherently accommodates systems that grow and shrink (condition 4) as a by-product of the embedded monitor code that executes on the same thread of system components; see fig. 1a. This scheme, however, has shortcomings that make it less suited to reactive software. Recent studies [21, 51] observe that the lock-step execution of the SuS and monitors can impair the operation of the instrumented system, *e.g.* slow runtime analyses manifest as high latencies [36], and faulty monitors may break the system [68], which do not meet conditions 1 and 2 (*e.g.* M_Q in fig. 1a). Other works [45, 15] argue that errors, such as deadlocks or component crashes, are hard to detect since the monitoring logic shares the runtime thread of the affected component. Entwining the execution of the SuS and monitors may also diminish the scalability, performance, and resource usage efficiency of the monitored system because inlined monitor code cannot be run on separate threads [12]. Lastly, inlining is *incompatible* with unmodifiable software, such as closed-source components (*e.g.* R in figs. 1a–1c), making outlining the only alternative.

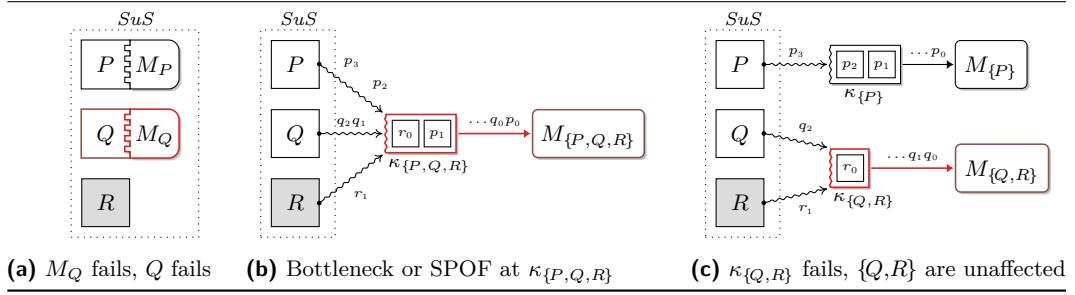
Outline instrumentation *can* address the limitations of inlining by isolating the SuS and its monitors (works [44, 36, 37] that view externalised monitors as ‘outline’ embed tracing code to extract events from the SuS, subjecting them to the cons of inlining). The latest survey on decentralised RV [70, tables 1 and 2] establishes that outlining-based tools, *e.g.* [49, 17, 18, 71, 36, 37, 124, 64], are variations on *centralised* instrumentation. In this set-up, events exhibited by SuS components are funnelled through a *global* trace buffer (*e.g.* $\kappa_{\{P,Q,R\}}$ in fig. 1b) that a singleton monitor can analyse asynchronously, meeting condition 3. Yet, the central buffer introduces contention and sacrifices the scalability of the SuS [11], violating condition 4. Centralised architectures are prone to single point of failures (SPOFs) [93, 92] (violating condition 2), which is not ideal for monitoring medium-scale reactive systems.

Contribution We propose RIARC, a *decentralised* instrumentation algorithm for outline monitors that overcomes the above shortcomings, fulfilling conditions 1–4. Outline monitors minimise latency effects due to slow trace event analyses associated with inlining (meeting condition 1). While RIARC does not handle monitor failure explicitly, it intrinsically enjoys a degree of fault tolerance by isolating the SuS and its decentralised monitor components (meeting condition 2); *e.g.* monitors $M_{\{P\}}$ and $M_{\{Q,R\}}$ in fig. 1c. RIARC uses a tracing infrastructure to obtain system events passively without modifying the SuS (meeting condition 3). The algorithm equips each isolated monitor with a *local* trace buffer, using it to report events based on the SuS components a monitor is tasked to analyse (*e.g.* buffers $\kappa_{\{P\}}$ and $\kappa_{\{Q,R\}}$ in fig. 1c). RIARC reorganises its instrumentation set-up to reflect dynamic changes in the SuS. It reacts to specific events in traces to instrument monitors for new SuS components and to remove redundant monitors when it detects graceful or abnormal component terminations. This enables RIARC to grow and shrink the verification set-up on demand (meeting condition 4). Given the challenges of fulfilling the conditions 1–4, we scope our work to settings where communication is reliable (*i.e.*, no message corruption, duplication, and loss) [56] and Byzantine failures do not arise [95].

To the best of our knowledge, the approach RIARC advocates is novel. One reason why outlining has never been adopted for decentralising monitors are the onerous conditions 1–4 imposed by reactive software. Utilising non-invasive tracing makes our set-up necessarily *asynchronous*. At the same time, this complicates the instrumentation, which must ensure trace soundness (def. 1), notwithstanding the inherent phenomena arising from the concurrent execution of the SuS and monitors, *e.g.* trace event reordering and process crashes. Consequently, the second reason is that the overhead incurred to uphold this invariant—in addition to scaling the verification set-up as the SuS executes—is perceived as prohibitive when compared to inlining. This opinion is often reinforced when the viability of outline instrumentation is predicated on empirical criteria tied to monolithic, batch-style programs, that *may not* apply to reactive software (*e.g.* percentage slowdown); *e.g.* see [96, 113, 112, 46, 45, 118, 29, 97].

This paper shows how instrumenting outline monitors under conditions 1–4 can be achieved using a decentralised approach that guarantees def. 1, while *also* exhibiting overheads considered feasible for typical soft real-time reactive systems. Concretely, we

- (i) recall the benefits of the actor model [81, 10] for building reactive systems and argue how our model of processes and tracers readily maps to that setting, sec. 2;
- (ii) give a decentralised instrumentation algorithm for outline monitors, detailing how the reactive characteristics of the SuS can be preserved whilst ensuring def. 1, sec. 3;
- (iii) show the implementability of our algorithm in an actor language and systematically validate the correctness of its corresponding implementation w.r.t. def. 1 by exhaustively inducing interleaved executions for a selection of instrumented systems, sec. 4;



■ **Figure 1** P, Q, R instrumented in inline (left), centralised (middle) and decentralised (right) modes

- (iv) back up the feasibility of the implemented algorithm via a comprehensive empirical study that uses various workload configurations surpassing the state of the art, showing that the induced overhead minimally impacts the reactive attributes of the SuS, sec. 5.

The extended version [8] contains the full details about RIARC and further discussion of our experiments and results. That material is ancillary to the one presented in this paper.

2 A computational model for reactive systems

The actor model [81, 10] emerged as *the* paradigm to design and build reactive systems [32]. *Actors*—the units of decomposition in this model—are abstractions of concurrent entities that share no mutable memory with other actors. Instead, actors interact through asynchronous message passing and alter their internal state based on the messages they consume. Asynchronous communication decouples actors spatially and temporally, which fully isolates system components and establishes the foundation for resiliency and elasticity [31, 93]. Each actor is equipped with an incoming message buffer called the *mailbox*, from which messages deposited by other actors can be selectively read. Besides sending and receiving messages, actors can *spawn* other actors. Actors in a system are addressable by their unique process identifier (PID), which they use to engage in directed, *point-to-point* communication. This idea of addressability is central to the actor model: it enables reasoning about decentralised computation, as the identity of components or messages can be propagated through a system and used in handling complex tasks, such as process registration and failure recovery [32]. As is often the case in decentralised computations, we assume that messages exchanged between pairs of processes are always received in the order in which they have been sent [42].

Frameworks, notably Erlang [12], Elixir [87], Akka [1] for Scala [116], along with others [117, 129], instantiate the actor model. We adopt Erlang since its ecosystem is specifically engineered for highly-concurrent, soft real-time reactive systems [130, 13, 43]. The Erlang virtual machine (EVM) implements actors as lightweight processes. It employs *per process* garbage collection that, unlike the JVM, does not subject the virtual machine to global unpredictable pauses [85, 115]. This factor minimises the impact on the soft real-time properties of a system *and* is also crucial to the empirical evaluation of sec. 5 since it stabilises the variance in our experiments. The EVM exposes a flexible *process tracing* API aimed at reactive software [41]. Erlang provides other components, *e.g.* supervision trees, message queues, *etc.*, for building fault-tolerant distributed applications. While we scope our work to fault-free settings (see sec. 1), adopting Erlang gives us the foundation upon which our work can be naturally extended to address these aspects. Henceforth, we follow the established convention in Erlang literature and use the terms *actor*, *process*, and *component* synonymously.

2.1 Process tracing and trace partitioning

Processes in a concurrent system form a *tree*, starting at the *root* process that spawns *child* processes, and so forth¹. Concurrency induces inherent *partitions* to the execution of the SuS in the form of isolated traces that reflect the *local* behaviour at each process [18]. RIARC exploits this aspect to attain several benefits. First, one can *selectively* specify the SuS processes to be instrumented. The upshot is that fewer trace events need to be gathered, improving *efficiency*. Another benefit of partitioned traces is that each process can be dynamically instrumented, free from assumptions about the number of processes the SuS is expected to have. This makes the RV set-up *elastic*. Lastly, the instrumentation set-up can *partially fail*, as faulty SuS or monitor processes do not imperil the execution of one another.

► **Example 2** (Trace partitions). Trace partitions enable RIARC to instrument a system in various arrangements. Fig. 2a depicts an interaction sequence for the execution of the SuS from sec. 1. In this interaction, the root process, P , spawns Q and communicates with it, at which point Q spawns process R ; P and Q eventually terminate. We denote the process *spawning* and *termination* trace events by $\diamond$ and $\star$, and use $!$ and $?$ for *send* and *receive* events respectively. The *sound* trace partitions for the processes in fig. 2a are ' $\diamond_P.!.P.\star_P$ ' for P , ' $?_Q.\diamond_Q.\star_Q$ ' for Q , and the empty trace for R . ◀

A centralised set-up such as that of fig. 1b can be obtained by instrumenting $\{P, Q, R\}$ with one monitor, $M_{\{P, Q, R\}}$, whereas instrumenting the components $\{P\}$ and $\{Q, R\}$ with monitors $M_{\{P\}}$ and $M_{\{Q, R\}}$ gives the decentralised arrangement of fig. 1c. Each of these instrumentation arrangements generates different executions.

► **Example 3** (Sound traces). For the case of fig. 1b, RIARC can report to $M_{\{P, Q, R\}}$ one of four possible traces ' $\diamond_P.!.P.\star_P.?_Q.\diamond_Q.\star_Q$ ', ' $\diamond_P.!.P.?_Q.\star_P.\diamond_Q.\star_Q$ ', ' $\diamond_P.!.P.?_Q.\diamond_Q.\star_P.\star_Q$ ', or ' $\diamond_P.!.P.?_Q.\diamond_Q.\star_Q.\star_P$ '. These *sound* traces result from the interleaved execution of processes P , Q . Any other trace, e.g. ' $\diamond_P.\star_P.?_Q.\diamond_Q.\star_Q$ ' or ' $\diamond_P.!.P.\star_P.?_Q.\star_Q.\diamond_Q$ ', is *unsound* since it contradicts the local behaviour at processes P and Q of fig. 2a. The former trace omits the request $!.P$ that P makes to Q (it is *incomplete* w.r.t. P), and the latter trace inverts $\diamond_Q$ and $\star_Q$, suggesting that Q spawns R after Q terminates (it is *inconsistent* w.r.t. Q). ◀

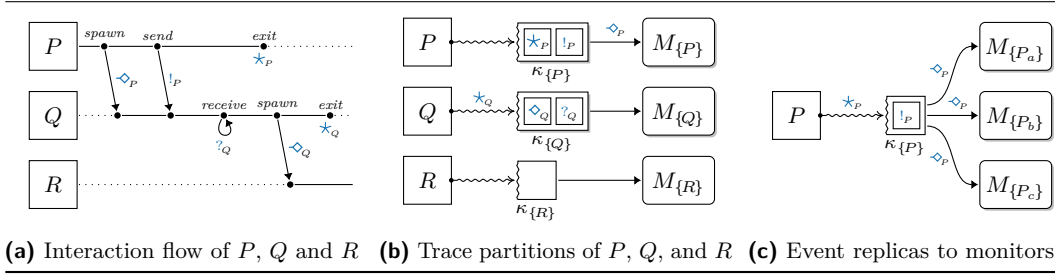
► **Example 4** (Separate instrumentation). Fig. 2b shows another decentralised set-up, where P , Q , and R are instrumented separately. In this case, the instrumentation should report to $M_{\{P\}}$, $M_{\{Q\}}$ and $M_{\{R\}}$ the events observed *locally* at each process, as stated in ex. 2. ◀

RIARC makes two assumptions about process tracing in order to support the instrumentation arrangements shown in figs. 1b, 1c, and 2b:

- A₁ Tracing processes sets.** Tracing can gather events for *sets* of SuS processes, e.g. $\kappa_{\{P, Q, R\}}$ in fig. 1b gathers the events of $\{P, Q, R\}$, and $\kappa_{\{Q, R\}}$ in fig. 1c gathers the events of $\{Q, R\}$.
- A₂ Tracing inheritance.** Tracing gathers the events of a SuS process *and* the children it spawns by default to eliminate the risk that trace events from child processes are missed.

We opt for tracing inheritance since it follows established centralised RV monitoring tools, including [17, 40, 49, 109]. In fact, tracing assumptions **A₁** and **A₂** mean that centralised set-ups like that of fig. 1b can be obtained just by tracing the root process P . Tracing inheritance requires the instrumentation to *intervene* if it needs to channel the events of a child process into a *new* trace partition that is *independent* from that of its parent, e.g. as in

¹ For example, using `spawn()` in Erlang [41] and Elixir [87], `ActorContext.spawn()` in Akka [1], `Actor.createActor()` in Thespian [117], `CreateProcess()` in Windows [107], etc.



■ **Figure 2** SuS with processes P , Q , and R instrumented with independent monitors

fig. 1c. In such cases, the instrumentation must first stop tracing the child process, allocate a fresh trace buffer, and resume tracing the child process. The out-of-sync execution of the SuS and instrumentation complicates the creation of these new trace partitions because it can lead to reordered or missed events. This, in turn, would violate trace soundness, def. 1.

We supplement A_1 and A_2 with the following to keep our exposition in sec. 3 manageable:

- A_3 *Single-process tracing.* Any SuS process can be traced *at most* once at any point in time.
- A_4 *Causally-ordered spawn events.* Tracing gathers the spawn trace event of a parent process before *all* the events of the child process spawned by that parent, *e.g.* if P spawns Q , and Q receives, as in fig. 2a, the reported sequence is ' $\star_P.\star_Q$ ' rather than ' $?_Q.\star_P$ '.

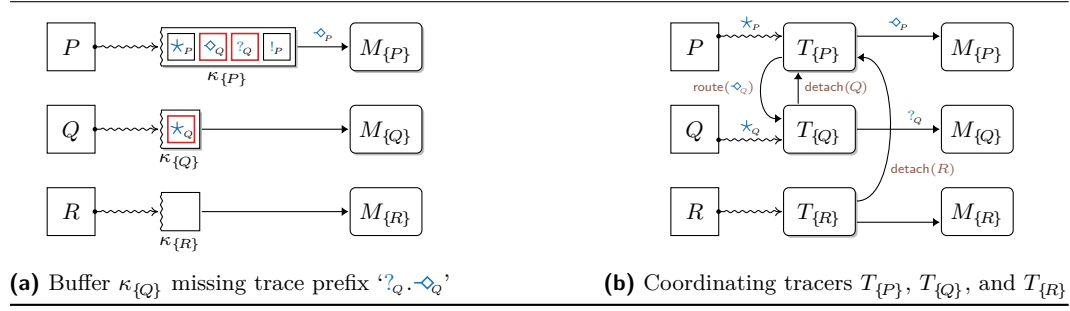
The constraint of tracing assumption A_3 is easily overcome by replicating trace events for a process and reporting them to different monitors (*e.g.* the events in the trace partition of process P are replicated to monitors $M_{\{P_a\}}$, $M_{\{P_b\}}$, $M_{\{P_c\}}$ in fig. 2c). Tracing assumption A_4 requires trace buffers to reorder $\star$ events using the spawner and spawned process information carried by each event before reporting them to monitors. Sec. 3.3 gives more details.

► **Example 5 (Unsound traces).** Fig. 3a shows one possible configuration that can be reached by our three-process system introduced in fig. 2a, where the trace buffer $\kappa_{\{P\}}$ contains the events for both P and Q . The trace in buffer $\kappa_{\{Q\}}$ is unsound, as it inaccurately characterises the local behaviour of process Q (the sound trace for Q should be ' $?_Q.\star_Q$ ', not ' $\star_Q$ '). ◀

RIARC programs trace buffers to coordinate with one another to ensure that sound traces are invariably reported to monitors. We refer to a trace buffer and the coordination logic it encapsulates as a *tracer*. RIARC employs an approach based on *next-hop routing* in IP networks [79, 103] to counteract the effects of trace event reordering and loss by rearranging and forwarding events to different tracers. Fig. 3b conveys our organisation of tracers (refer to [8, fig. 10 in app. A] for legend). Sec. 3 details how RIARC dynamically reorganises the tracer choreography and performs next-hop routing.

2.2 Modelling decentralised instrumentation

Since RV monitors are passive verdict-flagging machines (refer to sec. 1), they are orthogonal to our instrumentation. We, thus, focus our narrative on tracers and omit monitors, except when relevant in the surrounding context. The model assumes a set of SuS process, $P, Q, R \in \text{PRC}$, and tracer names, $T \in \text{TRC}$, together with a countable set of PID values to reference processes. We distinguish between SuS and tracer PIDs, which we denote respectively by the sets, $p_s, q_s \in \text{PID}_S$ and $p_T, q_T \in \text{PID}_T$. The variables i_s and j_s , and i_T and j_T range over PIDs from the corresponding sets PID_S and PID_T . We also assume the function signature sets, $f_s \in \text{SIG}_S$, $f_T \in \text{SIG}_T$, and, $f_M \in \text{SIG}_M$, to denote SuS, tracer, and RV monitor functions, together with the variables ς_s , ς_T , and ς_M that range over each signature set. New SuS processes are created



■ **Figure 3** Choreographed tracers coordinating to ensure sound traces

via the function $\text{spwn}(\varsigma_s)$ that accepts the function signature ς_s to be spawned, and returns a fresh PID, ι_s . We overload spwn to spawn tracer signatures ς_T equivalently, returning corresponding PIDs, ι_T . The function self obtains the PID of the process invoking it. We write P as shorthand for a singleton process set $\{P\}$ to simplify notation.

RIARC uses three message types, $\tau \in \{\text{evt}, \text{dte}, \text{rtd}\}$. These determine when to *create* or *terminate* tracer processes, and what trace events to *route* between tracers:

- **evt** are *trace events* gathered via process tracing,
- **dte** are *detach* requests that tracers exchange to reorganise the tracer choreography, and
- **rtd** are *routing* packets that transport **evt** or **dte** messages forwarded between tracers.

We encode messages m as tuples. Trace event messages, $\langle \text{evt}, \ell, \iota_s, j_s, \varsigma_s \rangle$, comprise the event label ℓ that ranges over the SuS events $\diamond$ (*spawn*), $\star$ (*exit*), $!$ (*send*), and $?$ (*receive*). The PID value ι_s identifies the SuS process exhibiting the trace event, and is defined for *all* events. The SuS PID j_s and function signature ς_s depend on the type of the event. Tbl. 1a catalogues the values defined for each event. We write trace events in their shorthand form, omitting undefined values (denoted by $\perp$), e.g. $\langle \text{evt}, \star, \iota_s \rangle$ instead of $\langle \text{evt}, \star, \iota_s, \perp, \perp \rangle$.

Detach request messages have the form $\langle \text{dte}, \iota_T, \iota_s \rangle$. A tracer with the PID ι_T uses **dte** to request that another tracer *stop* tracing the SuS PID ι_s . Routing packet messages, $\langle \text{rtd}, \iota_T, m \rangle$, move **evt** and **dte** messages between tracers. The PID ι_T identifies the tracer that embeds the

Label ℓ	Index	Description (ι_s and j_s are SuS PIDs)
$\diamond$	$e.\iota_s$	Parent PID spawning new child PID j_s
	$e.j_s$	Child PID spawned by parent PID ι_s
	$e.\varsigma_s$	Signature ς_s spawned by parent PID ι_s
$\star$	$e.\iota_s$	Terminated PID
	$e.j_s, e.\varsigma_s$	Undefined for exit events
$!$	$e.\iota_s$	Sending PID
	$e.j_s$	Recipient PID
	$e.\varsigma_s$	Undefined for send events
$?$	$e.\iota_s$	Recipient PID
	$e.j_s, e.\varsigma_s$	Undefined for receive events

(a) Messages encoding *spawn*, *exit*, *send*, and *receive* events

(b) Detach and routing messages

■ **Table 1** Trace event (**evt**), detach request (**dte**), and routing packet (**rtd**) message index names

Requirement	Approach
R ₁ Growing the set-up	Instrument tracers on-demand to create new trace partitions
R ₂ Ensuring complete traces	Route trace events to deliver them to the correct tracer
R ₃ Ensuring consistent traces	Prioritise routed trace events before others
R ₄ Isolating tracers	Detach tracers from others once all trace events are routed
R ₅ Minimising overhead	Target specific processes to instrument
R ₆ Shrinking the set-up	Garbage collect redundant tracers and monitors

■ **Table 2** RIARC approach to ensure trace soundness (def. 1) and reactive instrumentation (sec. 1)

message m into the routing packet and dispatches it to other tracers. Tbl. 1b summarises detach request and routing packet messages.

► **Note 6 (Notation).** We reserve the variables e , d , and r for the messages types `evt`, `dte`, and `rtd` respectively. Our model uses the suggestive dot notation ($\cdot$) to index message fields, *e.g.* $m.\tau$ reads the message type, $e.\ell$ reads the trace event label, *etc.* (see tbl. 1). For simplicity, we occasionally write the label ℓ in lieu of the full trace event form, *e.g.* we write $\star$ instead of $\langle \text{evt}, \star, \iota_s \rangle$, *etc.* ◀

3 Decentralised instrumentation

Our reason for encapsulating trace buffers and their coordination logic as tracers stems from the actor model. Trace buffers align with actor mailboxes, which localise the tracer state and enable tracers to run *independently*. The main logic replicated at each tracer is given in algs. 1–3. Tracers operate in two modes, *direct* ($\circ$) and *priority* ($\bullet$), to counteract the effects of trace event reordering. We organise our tracer logic in algs. 1 and 3 to reflect these modes, respectively. Algs. 1 and 3 use the function `ANALYSEEVENT`, which analyses events; see [8, app. C.5.2] for details. Auxiliary tracer logic referenced in this section is given in [8, app. A].

Every tracer maintains an internal state σ consisting of the following three maps:

- the *routing* map, Π , governing how events are routed to other tracers,
- the *instrumentation* map, Λ , that determines which SuS processes to instrument, and
- the *traced-processes* map, Γ , tracking the SuS process set that the tracer currently traces.

Tbl. 2 summarises the challenges that RIARC needs to overcome to attain the reactive characteristics stated in sec. 1. Requirements R_1 and R_6 in tbl. 2 oblige the instrumentation to reorganise dynamically while the SuS executes to preserve its *elasticity*. Requirement R_4 offers a modicum of *resiliency* between the SuS and tracer processes, whereas R_5 minimises the instrumentation overhead by gathering only the events monitors require. This keeps the overall set-up *responsive*. Since RIARC builds on the actor model, it fulfils the *message-driven* requirement intrinsically. *Trace soundness* is safeguarded by requirements R_2 and R_3 .

The operations `TRACE`, `CLEAR` and `PREEMPT` give access to the tracing infrastructure. `TRACE`(ι_s, ι_T) enables a tracer with PID ι_T to register its interest in receiving trace events of a SuS process with PID ι_s . This operation can be undone using `CLEAR`(ι_s, ι_T), which *blocks* the calling tracer ι_T and returns once all the trace event messages for the SuS process ι_s that are in transit to the tracer ι_T have been delivered to ι_T . It is worth remarking that this behaviour conforms to our proviso in sec. 1, *i.e.*, no communication faults. `PREEMPT`(ι_s, ι_T) combines `CLEAR` and `TRACE`. It enables the tracer pre-empting ι_T to take control of tracing the SuS process ι_s from another tracer ι'_T that is currently tracing ι_s . Tracers use `CLEAR` or `PREEMPT`

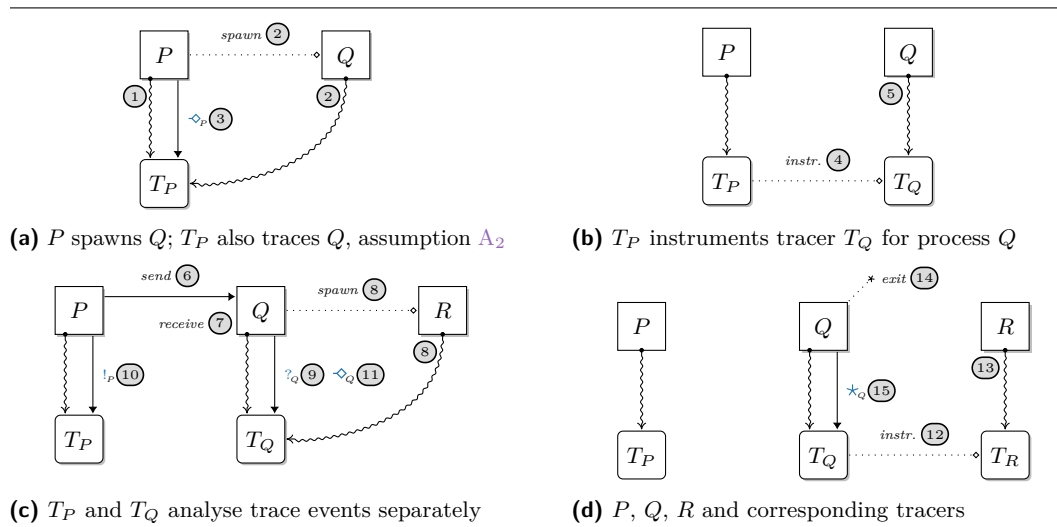
to modify the default process-tracing inheritance behaviour that tracing assumption A_2 describes. We refer readers to [8, alg. 5 in app. A] for the specifics of these operations.

We focus our presentation in secs. 3.1–3.6 of how RIARC addresses the challenges listed in tbl. 2 on the set-up of fig. 2b, where the processes P , Q and R , are instrumented separately. This specific case highlights two aspects. First, it *emphasises* the complications that RIARC overcomes to establish the desired set-up while ensuring trace soundness, def. 1. Second, fig. 2b *covers all* other possible instrumentation set-ups. Disjoint sets of SuS processes, including the one shown in fig. 1c, can be obtained when tracers do not act on certain $\diamond$ (*spawn*) events, as sec. 3.1 explains. Notably, *any* centralised set-up, *e.g.* the one in fig. 1b, emerges naturally when the root tracer disregards all $\diamond$ events exhibited by the SuS.

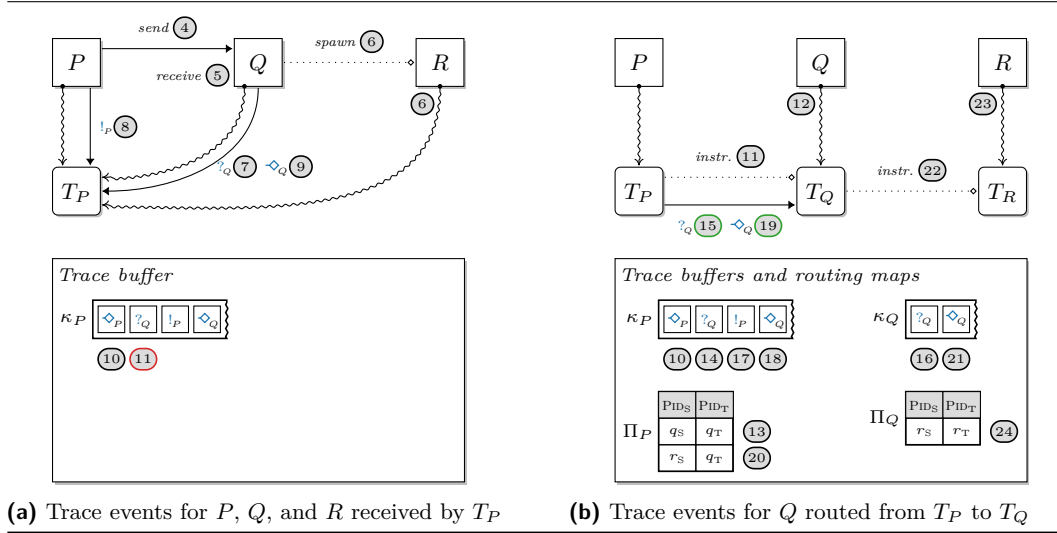
► **Note 7 (Naming conventions).** For clarity, we adopt the convention that a SuS process P is spawned from the signature f_{s_P} and is assigned the PID p_s . A tracer for P is named T_P (short for $T_{\{P\}}$) and has the PID p_T . Other processes are treated likewise, *e.g.* the SuS process Q has signature f_{s_Q} , PID q_s , while the tracer T_Q for Q has PID q_T , *etc.* ◀

3.1 Growing the set-up

Fig. 4 illustrates how the hierarchical creation sequence of SuS processes described in sec. 2.1 is exploited to instrument separate tracers. RIARC programs tracers to react to $\diamond$ (*spawn*) events in the trace. In fig. 4a, the root tracer T_P traces process P , step ①. When P spawns process Q , Q automatically inherits T_P (tracing assumption A_2 from sec. 2.1). Steps ② in fig. 4a emphasise that tracing inheritance is instantaneous. The event $e = \langle \text{evt}, \diamond, p_s, q_s, f_{s_Q} \rangle$ is generated by P when it spawns its child Q , step ③ in fig. 4a. The PID values of the parent and child processes carried by e , namely p_s and q_s , are accessed via the indexes $e.i_s$ and $e.j_s$ respectively (see tbl. 1a). Tracer T_P uses this PID information to instrument a new tracer T_Q for process Q in step ④ of fig. 4b. By invoking $\text{PREEMPT}(q_s, q_T)$, T_Q takes over tracing process Q from the former tracer T_P going forward. T_Q creates a new trace partition for process Q that is independent of the partition of P , step ⑤. Meanwhile, T_P receives the send event $\langle \text{evt}, !, p_s, q_s \rangle$ in step ⑩ after P messages Q in step ⑥ of fig. 4c. Subsequent $\diamond$ events that T_P or T_Q may gather are handled as described in steps ③–⑤. Figs. 4c and 4d show



■ **Figure 4** Growing the tracer instrumentation set-up for processes P , Q and R (monitors omitted)



■ **Figure 5** Next-hop trace event routing using tracer routing maps Π (monitors omitted)

how the final tracer T_R is instrumented in step 12 after Q spawns R in step 8. As before, T_Q traces R automatically in step 8. T_Q receives the event $\langle \text{evt}, \diamond, q_s, r_s, f_{s_R} \rangle$ generated by Q in step 11. T_R invokes $\text{PREEMPT}(r_s, r_T)$ to create the trace partition for R in step 13.

3.2 Ensuring complete traces

The asynchrony between the SuS and tracer processes can induce the interleaved execution shown in fig. 5, as an alternative execution to that shown in figs. 4b–4d. In fig. 5a, T_P is slow to handle $\diamond_P$ it receives in 3 of fig. 4a and fails to instrument T_Q promptly. Consequently, the events $?_Q$ and $\diamond_Q$ that Q exhibits are sent to T_P in steps 7 and 9 of fig. 5a. Step 11 shows the case where $\langle \text{evt}, ?, q_T \rangle$ is processed by T_P , rather than by the intended tracer T_Q that would have been instrumented by T_P . This error breaches the *completeness* property of trace soundness w.r.t. Q , as the events $?_Q$ and $\diamond_Q$ meant for Q reach the wrong tracer T_P .

To address this issue, RIARC uses a next-hop routing approach, where tracers *retain* the events they should handle and *forward* the rest to neighbouring tracers. We use the term *dispatch tracer* (*dispatcher* for short) to describe a tracer that receives trace events meant to be handled by another tracer. For instance, T_P in fig. 5a becomes the dispatch tracer for Q when it receives the events $?_Q$ and $\diamond_Q$ exhibited by Q , steps 7 and 9. We expect these events to be handled by T_Q once it is instrumented. Dispatchers are tasked with embedding trace event (*evt*) or detach requests (*dtc*) into routing packet messages (*rtid*) and transmitting them to the next *known* hop. In fig. 5b, T_P dispatches the events $?_Q$ and $\diamond_Q$ as follows. It first instruments T_Q with Q in step 11. Next, T_P prepares $\langle \text{evt}, ?, r_s \rangle$ and $\langle \text{evt}, \diamond, q_s, r_s, f_{s_R} \rangle$ for transmission by embedding each in *rtid* messages (steps 14 and 18). T_P forwards the resulting routing packets, $\langle \text{rtid}, p_T, \langle \text{evt}, ?, r_s \rangle \rangle$ and $\langle \text{rtid}, p_T, \langle \text{evt}, \diamond, q_s, r_s, f_{s_R} \rangle \rangle$, to its next-hop neighbour T_Q in steps 15 and 19. The trace event $\langle \text{evt}, !, p_s, q_s \rangle$, however, is not forwarded but handled by T_P , as step 17 shows. Concurrently, T_Q acts on the forwarded events $?_Q$ and $\diamond_Q$ in steps 16 and 21 and instruments T_R as a result, step 22.

Tracers determine the events to retain or forward using the routing map, $\Pi: \text{PID}_S \rightarrow \text{PID}_T$. Every tracer queries its private routing map for each message it receives on SuS PID $m.i_s$. A tracer forwards a message to its neighbouring tracer with PID i_T if a next-hop for that

■ **Algorithm 1** Logic handling $\circ$ trace events, detach request dispatching, and forwarding

```

1  def LOOP $\circ$ ( $\sigma, \varsigma_M$ )
2  forever do
3     $m \leftarrow$  next message from trace buffer  $\kappa$ 
4    match  $m.\tau$  do
5      case evt:  $\sigma \leftarrow$  HANDLEEVENT $\circ$ ( $\sigma, \varsigma_M, m$ )
6      case dtc:  $\sigma \leftarrow$  DISPATCHDTC( $\sigma, \varsigma_M, m$ )
7      case rtd:  $\sigma \leftarrow$  FORWDRTD $\circ$ ( $\sigma, \varsigma_M, m$ )
8  def HANDLEVT $\circ$ ( $\sigma, \varsigma_M, e$ )
9  match  $e.l$  do
10 case  $\diamond$ : return HANDLSPWN $\circ$ ( $\sigma, \varsigma_M, e$ )
11 case  $\star$ : return HANDLEEXIT $\circ$ ( $\sigma, \varsigma_M, e$ )
12 case  $!, ?$ : return HANDLCOMM $\circ$ ( $\sigma, \varsigma_M, e$ )
13 def HANDLSPWN $\circ$ ( $\sigma, \varsigma_M, e$ )
14 match  $\sigma.\Pi(e.i_s)$  do
15 case  $\perp$ : # No next-hop for  $e.i_s$ ; handle  $e$ 
16   ANALYSEEVT( $\varsigma_M, e$ )
17    $\sigma \leftarrow$  INSTRUMENT $\circ$ ( $\sigma, e, \text{self}()$ )
18 case  $j_T$ : # Next-hop for  $e.i_s$  exists via  $j_T$ 
19   DISPATCH( $e, j_T$ )
20   # Set next-hop of  $e.j_s$  to tracer of  $e.i_s$ 
21    $\sigma.\Pi \leftarrow \sigma.\Pi \cup \{e.j_s, j_T\}$ 
22   return  $\sigma$ 
23 def HANDLEEXIT $\circ$ ( $\sigma, \varsigma_M, e$ )
24 match  $\sigma.\Pi(e.i_s)$  do
25 case  $\perp$ : # No next-hop for  $e.i_s$ ; handle  $e$ 
26   ANALYSEEVT( $\varsigma_M, e$ )
27    $\sigma.\Gamma \leftarrow \sigma.\Gamma \setminus \{e.i_s, \circ\}$ 
28   TRYGC( $\sigma$ )
29 case  $j_T$ : DISPATCH( $e, j_T$ )
30   return  $\sigma$ 
31 def HANDLCOMM $\circ$ ( $\sigma, \varsigma_M, e$ )
32 match  $\sigma.\Pi(e.i_s)$  do
33 case  $\perp$ : ANALYSEEVT( $\varsigma_M, e$ )
34 case  $j_T$ : DISPATCH( $e, j_T$ )
35   return  $\sigma$ 
36 def DISPATCHDTC( $\sigma, d$ )
37 match  $\sigma.\Pi(d.i_s)$  do
38 case  $\perp$ : fail dtc next-hop must be defined
39 case  $j_T$ :
40   DISPATCH( $d, j_T$ )
41   # Next-hop for  $d.i_s$  no longer needed
42    $\sigma.\Pi \leftarrow \sigma.\Pi \setminus \{d.i_s, j_T\}$ 
43   TRYGC( $\sigma$ )
44   return  $\sigma$ 
45 def FORWDRTD $\circ$ ( $\sigma, r$ )
46  $m \leftarrow r.m$  # Read embedded message in  $r$ 
47 match  $m.\tau$  do
48 case dtc: return FORWDDTC( $\sigma, r$ )
49 case evt: return FORWDEVT( $\sigma, r$ )
50 def FORWDDTC( $\sigma, r$ )
51  $d \leftarrow r.m$ 
52 match  $\sigma.\Pi(d.i_s)$  do
53 case  $\perp$ : fail dtc next-hop must be defined
54 case  $j_T$ :
55   FORWD( $r, j_T$ )
56   # Next-hop for  $d.i_s$  no longer needed
57    $\sigma.\Pi \leftarrow \sigma.\Pi \setminus \{d.i_s, j_T\}$ 
58   TRYGC( $\sigma$ )
59   return  $\sigma$ 
60 def FORWDEVT( $\sigma, r$ )
61  $e \leftarrow r.m$ 
62 match  $\sigma.\Pi(e.i_s)$  do
63 case  $\perp$ : fail evt next-hop must be defined
64 case  $j_T$ :
65   FORWD( $r, j_T$ )
66   # For spawn events, tracer also sets a
67   # new next-hop for  $e.j_s$ 
68   # Next-hop of  $e.j_s$  to same tracer of  $e.i_s$ 
69   if ( $e.l = \diamond$ )
70      $\sigma.\Pi \leftarrow \sigma.\Pi \cup \{e.j_s, j_T\}$ 
71   return  $\sigma$ 

```

message exists, *i.e.*, $\Pi(m.i_s) = i_T$. When the next-hop is undefined, *i.e.*, $\Pi(m.i_s) = \perp$, m is handled by the tracer. HANDLSPWN, HANDLEEXIT and HANDLCOMM in alg. 1 implement this forwarding logic on lines 14, 23 and 31.

Dynamically populating the routing map is key to transmitting messages between tracers. A tracer adds the new mapping $e.j_s \mapsto j_T$ to its routing map Π in case 1 or 2 below whenever it processes spawn trace events $e = \langle \text{evt}, \diamond, i_s, j_s, \varsigma_s \rangle$. One of two cases is considered for $e.i_s$:

1. $\Pi(i_s) = \perp$. The next-hop for e is undefined, which cues the tracer to instrument the SuS process with PID j_s . When applicable, the tracer processes the event *and* instruments a separate tracer with PID j_T . It then adds the mapping $e.j_s \mapsto j_T$ to Π . The tracer leaves Π *unmodified* and handles the event itself if a separate tracer is not required. Opting for a separate tracer is determined by the instrumentation map Λ , as discussed in sec. 3.5.

■ **Algorithm 2** Tracer instrumentation operations for direct ($\circ$) and priority ($\bullet$) modes

Expect: $e = \langle \text{evt}, \diamond, \iota_s, j_s, \varsigma_s \rangle$	Expect: $e = \langle \text{evt}, \diamond, \iota_s, j_s, \varsigma_s \rangle$
<pre> 1 def INSTRUMENT$\circ$(σ, e, ι_T) 2 if $((\varsigma_M \leftarrow \sigma.\Lambda(e.\varsigma_s)) \neq \perp)$ 3 # New tracer j_T for new SuS process $e.j_s$ 4 $j_T \leftarrow \text{spwn}(\text{TRACER}(\sigma, \varsigma_M, e.j_s, \iota_T))$ 5 $\sigma.\Pi \leftarrow \sigma.\Pi \cup \{ \langle e.j_s, j_T \rangle \}$ 6 else 7 # In $\circ$ mode, this tracer has detached 8 # all processes from its dispatcher ι_T 9 # This tracer traces new SuS process $e.j_s$ 10 # by tracing inheritance assumption A_2 11 $\sigma.\Gamma \leftarrow \sigma.\Gamma \cup \{ \langle e.j_s, \circ \rangle \}$ 12 return σ </pre>	<pre> 8 def INSTRUMENT$\bullet$(σ, e, ι_T) 9 if $((\varsigma_M \leftarrow \sigma.\Lambda(e.\varsigma_s)) \neq \perp)$ 10 # New tracer j_T for new SuS process $e.j_s$ 11 $j_T \leftarrow \text{spwn}(\text{TRACER}(\sigma, \varsigma_M, e.j_s, \iota_T))$ 12 $\sigma.\Pi \leftarrow \sigma.\Pi \cup \{ \langle e.j_s, j_T \rangle \}$ 13 else 14 # In $\bullet$ mode, this tracer must detach 15 # SuS process $e.j_s$ from its dispatcher ι_T 16 DETACH($e.j_s, \iota_T$) 17 # This tracer traces new SuS process $e.j_s$ 18 $\sigma.\Gamma \leftarrow \sigma.\Gamma \cup \{ \langle e.j_s, \bullet \rangle \}$ 19 return σ </pre>

2. $\Pi(\iota_s) = j_T$. The next-hop for e is defined, and the tracer forwards the event to the neighbouring tracer j_T . The tracer also records a new next-hop by adding $e.j_s \mapsto j_T$ to Π . The addition of $e.j_s \mapsto j_T$ in cases 1 and 2 ensures that future events originating from j_s can always be forwarded via a next-hop to a neighbouring tracer j_T (see invariants on lines 37, 51, and 60). Fig. 5b shows the routing maps of the tracers T_P and T_Q . T_P adds $q_s \mapsto q_T$ in step 13 after processing $\langle \text{evt}, \diamond, p_s, q_s, f_{s_Q} \rangle$ from its trace buffer in 10. T_P then instruments Q with the tracer T_Q in step 11; an instance of case 1. The function INSTRUMENT in alg. 2 details this on line 4, where the mapping $e.j_s \mapsto j_T$ is added to Π following the creation of tracer j_T , line 3. Step 20 of fig. 5b is an instance of case 2. Here, T_P adds $r_s \mapsto q_T$ to Π_P after processing $\langle \text{evt}, \diamond, q_s, r_s, f_{s_R} \rangle$ for R in step 18 since $\Pi_P(q_s) = q_T$. Crucially, T_P does not instrument a new tracer, but delegates the task to T_Q by forwarding $\diamond_Q$. Lines 20 and 64 in alg. 1 (and later line 24 in alg. 3) are manifestations of this, where the mapping $e.j_s \mapsto j_T$ is added after the $\diamond$ event e is forwarded to the next-hop j_T . T_Q instruments the SuS process R in step 22 with T_R , which has the PID r_T . It then adds the mapping $r_s \mapsto r_T$ to Π_Q in step 24, as no next-hop is defined for q_s , i.e., $\Pi_Q(q_s) = \perp$. Henceforth, any events exhibited by R and received at T_P can be dispatched by the latter tracer through T_Q to T_R .

Note that every tracer is *only* aware of its neighbouring tracers. This means messages may pass through multiple tracers before reaching their intended destination. Next-hop routing keeps the logic inside RIARC straightforward since tracers forward messages based on local information in their routing map. This approach makes the instrumentation set-up adaptable to dynamic changes in the SuS and has been shown to induce lower latency when compared to general routing strategies [79, 103]. The DAG of interconnected tracers induced by next-hop routing ensures that every message is eventually delivered to the correct tracer if a path exists or handled by the tracer otherwise. Fig. 5b illustrates this concept, where the next-hop mappings inside Π_P point to T_Q , and the mappings in Π_Q point to T_R . Consequently, any events that R exhibits and that T_P receives are forwarded *twice* to reach the target tracer T_R : from tracer T_P to T_Q , and from T_Q to T_R . RIARC relies on the operations DISPATCH and FORWD to achieve next-hop routing (see [8, alg. 4 in app. A]). DISPATCH creates a routing packet, $\langle \text{rtd}, \iota_T, m \rangle$, and embeds the trace event or detach message m to be routed. Alg. 1 shows how tracers handle routing packets. For instance, FORWDEVT extracts the embedded message from the routing packet on line 58 and queries the routing map to determine the next-hop, line 59. If found, the packet is forwarded, as FORWD(r, j_T) on line 62 indicates. Crucially, the **fail** invariant on line 60 asserts that the next-hop for a routing packet is *always* defined. The cases for DISPATCHDTC and FORWDDTC in alg. 1 are analogous.

3.3 Ensuring consistent traces

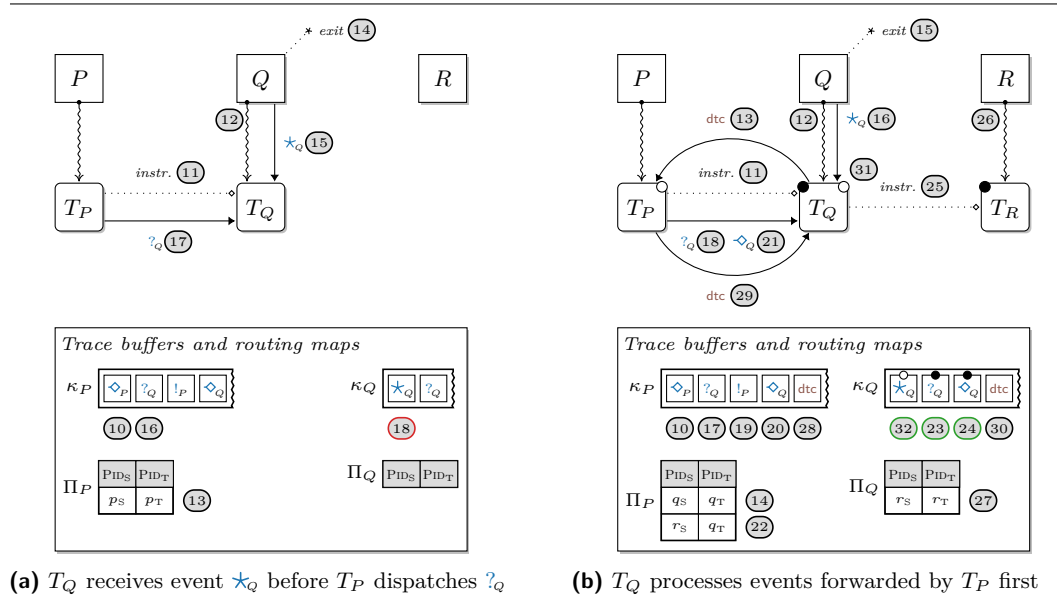
Next-hop routing alone does not guarantee trace consistency, *i.e.*, that the order of events in the trace reflects the one in which these occur locally at SuS processes, def. 1. Trace event reordering arises when a tracer gathers events of a SuS process (we call these *direct events*) and simultaneously receives *routed events* concerning said process from other tracers. Fig. 6a gives another interleaving to the one of fig. 5b to underscore the deleterious effect such a race condition provokes when events are reordered at T_Q . In step (12) T_Q takes over T_P to continue tracing process Q . T_Q collects the event $\star_Q$ in step (15), which happens before T_Q receives the routed event $?_Q$ concerning Q in step (17) of fig. 6a. If T_Q processes events from its trace buffer κ_Q in sequence, as in step (18), it violates trace consistency w.r.t. Q (the correct trace ordering should be ‘ $?_Q \cdot \diamond_Q \cdot \star_Q$ ’). Naïvely handling $\star$ before $?$ erroneously reflects that Q receives messages after it terminates.

RIARC tracers resolve this issue by prioritising the processing of routed trace events using selective message reception [41]. In doing so, tracers encode the invariant that ‘*routed* events temporally precede all others that are gathered *directly* by the tracer’. RIARC tracers operate in one of two modes, priority ($\bullet$) and direct ($\circ$), which adequately distinguishes past (*i.e.*, routed) and current (*i.e.*, direct) events from the perspective of the tracer receiving them.

Fig. 6b illustrates this concept. It shows that when in priority mode, T_Q dequeues the routed events $?_Q$ and $\diamond_Q$ labelled by $\bullet$ first. The event $?_Q$ is handled in step (23), whereas $\diamond_Q$ results in the instrumentation of tracer T_R in step (25) of fig. 6b. Meanwhile, T_Q can still receive events directly from Q while priority events are being handled. Yet, direct trace events from Q are considered only *after* T_Q transitions to direct mode. Newly-instrumented tracers default to $\bullet$ mode to implement the described logic; see [8, line 14 in alg. 4 of app. A].

LOOP $\bullet$ in alg. 3 shows the logic prioritising routed events, which are dequeued on line 3 and handled on line 6. HANDLSPWN, HANDLEXIT, and HANDLCOMM in LOOP $\circ$ and LOOP $\bullet$ handle events *differently*. A tracer in direct mode performs *one* of three actions (see alg. 1):

1. it *analyses* events for RV purposes via the function ANALYSEEVT(ζ_M, e), *e.g.* line 32,



■ **Figure 6** Trace event reordering using priority ($\bullet$) and direct ($\circ$) tracer modes (monitors omitted)

■ **Algorithm 3** Logic handling • trace events, detach request acknowledgements, and forwarding

<pre> 1 def LOOP_•(σ, ς_M) 2 forever do 3 $r \leftarrow$ next rtd message from trace buffer κ 4 $m \leftarrow r.m$ # Read embedded message in r 5 match $m.\tau$ do 6 case evt: $\sigma \leftarrow$ HANDLEVT_•(σ, ς_M, r) 7 case dtc: 8 # dtc ack relayed from dispatch tracer 9 $\sigma \leftarrow$ HANDLDTC(σ, ς_M, r) 9 def HANDLEVT_•(σ, ς_M, r) 10 $e \leftarrow r.m$ 11 match $e.l$ do 12 case $\diamond$: return HANDLSPWN_•(σ, ς_M, r) 13 case $\star$: return HANDLEEXIT_•(σ, ς_M, r) 14 case $!?$: return HANDLCOMM_•(σ, ς_M, r) 15 def HANDLSPWN_•(σ, ς_M, r) 16 $e \leftarrow r.m$ 17 match $\sigma.\Pi(e.i_s)$ do 18 case $\perp$: # No next-hop for $e.i_s$; handle e 19 ANALYSEVT(ς_M, e) 20 $\iota_T \leftarrow r.i_T$ # Read PID of dispatch tracer 21 $\sigma \leftarrow$ INSTRUMENT_•(σ, e, ι_T) 22 case j_T: # Next-hop for $e.i_s$ exists via j_T 23 FORWD(r, j_T) 24 # Set next-hop of $e.j_s$ to tracer of $e.i_s$ 25 $\sigma.\Pi \leftarrow \sigma.\Pi \cup \{e.j_s, j_T\}$ 26 return σ </pre>	<pre> 26 def HANDLEEXIT_•(σ, ς_M, r) 27 $e \leftarrow r.m$ 28 match $\sigma.\Pi(e.i_s)$ do 29 case $\perp$: # No next-hop for $e.i_s$; handle e 30 ANALYSEVT(ς_M, e) 31 $\sigma.\Gamma \leftarrow \sigma.\Gamma \setminus \{e.i_s, \bullet\}$ 32 TRYGC(σ) 33 case j_T: FORWD(r, j_T) 34 return σ 35 def HANDLCOMM_•(σ, ς_M, r) 36 $e \leftarrow r.m$ 37 match $\sigma.\Pi(e.i_s)$ do 38 case $\perp$: ANALYSEVT(ς_M, e) 39 case j_T: FORWD(r, j_T) 40 return σ 41 def HANDLDTC(σ, ς_M, r) 42 $d \leftarrow r.m$ 43 match $\sigma.\Pi(d.j_s)$ do 44 case $\perp$: 45 assert $d.i_T = \text{self}()$ unexpected dtc ack 46 $\sigma.\Gamma \leftarrow (\sigma.\Gamma \setminus \{d.j_s, \bullet\}) \cup \{d.j_s, \circ\}$ 47 if $(\{i_s, \gamma\} \mid \langle i_s, \gamma \rangle \in \sigma.\Gamma, \gamma = \bullet\} = \emptyset)$ 48 LOOP_o(σ, ς_M) # Put tracer in o mode 49 case j_T: 50 assert $d.i_T \neq \text{self}()$ dtc meant for ι_T 51 FORWD(r, j_T) 52 return σ </pre>
---	--

2. it *dispatches* events that it directly gathers using DISPATCH(e, j_T), when events ought to be handled by other tracers, *e.g.* line 33, or
3. it *forwards* routed events to the next-hop through FORWD(r, j_T), *e.g.* line 62.

Tracers in priority mode exclusively handle routed messages as points 1 and 3 describe, *e.g.* lines 38 and 39 in alg. 3. However, no event dispatching is performed.

3.4 Isolating tracers

A tracer in priority mode coordinates with the dispatch tracer of a particular SuS process it traces. This enables the tracer to determine when *all* of the events of that process have been routed to it by the dispatch tracer. The negotiation is effected using *dtc*, which the tracer sends to the relevant dispatch tracer. Each tracer records the set of processes it traces in the *traced-processes map*, $\Gamma: \text{PID}_s \rightarrow \{\circ, \bullet\}$. A SuS process mapping is added to Γ when a tracer starts gathering trace events for that process and removed once the process terminates. Lines 6 and 14 in alg. 2 add fresh mappings to Γ ; lines 26 in alg. 1 and 31 in alg. 3 purge mappings from Γ . A tracer in priority mode must issue a *dtc* request for *each* process it tracks in Γ before it can transition to direct mode and start operating on the trace events it gathers directly. The detach request, $d = \langle \text{dtc}, \iota_T, \iota_s \rangle$, contains the PIDs of the issuing tracer and the SuS process to be detached from the dispatch tracer. Once the tracer receives an acknowledgement to the *dtc* request for the SuS PID $d.i_s$ from the dispatch tracer, it updates

the corresponding entry $d.i_s \mapsto \bullet$ in Γ , marking it as detached, $d.i_s \mapsto \circ$. Alg. 3 shows this logic on line 46. A tracer transitions from priority to direct mode once *all* the processes in its Γ map are marked detached; line 47 in alg. 3 performs this check. Once in direct mode, tracers are isolated from others in the choreography.

Fig. 6b depicts the tracer T_Q in priority mode sending the detach request $\langle \text{dtc}, q_T, q_s \rangle$ for SuS PID q_s to the dispatch tracer. This happens in step (13), after T_Q starts tracing Q directly in step (12). Alg. 2 effects this transaction with the dispatch tracer by the operation DETACH on line 13; see [8, app. A] for definition of DETACH. The dtc request issued by T_Q is deposited in the trace buffer of the dispatch tracer T_P after the events $?_Q$ and $\neg_Q$. T_P processes the messages in its buffer sequentially in (10), (17), (19), (20) and (28), and forwards $?_Q$ and $\neg_Q$ to T_Q , steps (18) and (21). Crucially, T_P *acknowledges* the dtc request issued by T_Q : T_P dispatches dtc back to tracer T_Q , as step (29) indicates. T_Q first handles the events $?_Q$ and $\neg_Q$ (tagged with $\bullet$ in fig. 6b) in steps (23) and (24). Lastly, T_Q handles dtc in (30) and marks process Q as detached from its dispatch tracer T_P . The update on the traced-process map Γ is performed by HANDLDTC on line 46 in alg. 3. Tracer T_Q in fig. 6b transitions to direct mode in step (31), when the only process Q that it traces is detached. T_Q resumes handling $\star_Q$ in step (32), which is consistent w.r.t. the events exhibited locally at Q , i.e., $?_Q \cdot \neg_Q \cdot \star_Q$.

An acknowledgement to a detach request sent from a dispatch tracer, $\langle \text{dtc}, i_T, i_s \rangle$, is generally propagated through multiple next-hops before it reaches the tracer with PID i_T issuing the request. Since a dtc request informs the dispatch tracer that i_T is gathering trace events for the SuS PID i_s *directly*, the next-hop entries in the routing maps of tracers on the DAG path from the dispatch tracer to i_T are *stale*. Each tracer on this DAG path purges the next-hop entry for the SuS PID i_s in Γ once it forwards dtc to the neighbouring tracer. DISPATCHDTC and FORWDDTC in alg. 1 perform this clean-up. Fig. 6b does not illustrate the latter clean-up flow, which we summarise next. After receiving dtc , the dispatch tracer T_P removes from Π_P the next-hop mapping $q_s \mapsto q_T$ and calls DISPATCHDTC to acknowledge the detach request $\langle \text{dtc}, q_T, q_s \rangle$ it receives from T_Q . Similarly, T_P removes $r_s \mapsto q_T$ once it acknowledges the detach request $\langle \text{dtc}, r_T, r_s \rangle$ sent from T_R . Once T_Q receives the routing packet $\langle \text{rtd}, p_T, \langle \text{dtc}, r_T, r_s \rangle \rangle$ that embeds the detach acknowledgement T_P sends, it removes the next-hop mapping $r_s \mapsto r_T$ from Π_Q . T_Q then forwards this dtc acknowledgement to T_R .

RIARC ensures that all routing packets carrying dtc acknowledgements terminate at the tracers that issued these dtc requests. This requires *one* of two tracer conditions to hold:

1. either the tracer cannot forward the dtc acknowledgement to a next-hop, meaning that the tracer sent the dtc request, or
2. the tracer can forward the dtc acknowledgement via a next-hop, in which case the tracer did not issue the dtc request.

Alg. 3 enforces this invariant on lines 44 and 45 for case 1, and on lines 49 and 50 for case 2.

3.5 Minimising overhead

Instrumenting specific processes—in contrast to fully instrumenting the SuS—reduces the volume of gathered trace events and helps lower the runtime overhead induced. RIARC uses the instrumentation map, $\Lambda: \text{SIG}_S \rightarrow \text{SIG}_M$, to this end. Λ specifies the SuS function signatures to instrument and the corresponding RV monitor signatures tasked with the analysis via ANALYSEEVT. RIARC utilises the signature e_{cs} carried by spawn events $e = \langle \text{evt}, \neg, i_s, j_s, \text{cs} \rangle$ to determine whether the SuS process spawning e_{cs} requires a separate tracer. The INSTRUMENT operations in alg. 2 perform this check against Λ (lines 2 and 9). If a separate tracer is not required, e_{js} is instrumented using the tracer of its parent process, $e.i_s$; see tracing assumptions A_1 and A_2 . This logic caters for all the set-ups shown in figs. 1b, 1c, and 2b.

3.6 Shrinking the set-up

RIARC remains elastic by discarding unneeded tracers. Tracers in direct and priority mode purge SuS PID references from the traced-process map when handling $\star$ trace events. `HANDLEEXITo` and `HANDLEEXIT•` implement this logic in algs. 1 and 3 on lines 26 and 31. Tracer termination does *not* occur when the tracer has no processes left to trace, *i.e.*, when $\Gamma = \emptyset$, since the tracer may be required to forward trace events to neighbouring tracers. Instead, tracers perform a garbage collection check each time a mapping from Γ or Π is removed. A tracer terminates when $\Gamma = \Pi = \emptyset$, indicating that it has no SuS processes left to trace or any next-hop forwarding to perform. TRYGC used on lines 27, 41, and 55 in alg. 1, as well as on line 32 in alg. 3 encapsulates this check. Note that garbage collection never prematurely disrupts the RV analysis that tracers conduct, as invocations to `ANALYSE EVT` always precede TRYGC checks in our logic of algs. 1 and 3.

4 Correctness validation

We assess the validity of RIARC in two stages. First, we confirm its implementability by instantiating the core logic of algs. 1–3 to Erlang. Our implementation targets two RV scenarios: online and offline monitoring [63, 21]. Second, we subject the implementation to a series of systematic tests using a selection of instrumentation set-ups. These tests exhaustively emulate the interleaved execution of the SuS and tracer processes by generating all the *valid* permutations of events in a set of traces. This exercises the tracer choreography invariants mentioned in sec. 3, confirming the integrity of the tracer DAG topology under each interleaving. We also use specialised RV monitor signatures in `ANALYSE EVT` to assert the soundness (def. 1) of trace event sequences analysed by tracers; see algs. 1 and 3 in sec. 3.

4.1 Implementability

Our implementation of RIARC maps the tracer processes from sec. 3 to Erlang actors. The routing (Π), instrumentation (Λ), and traced-processes (Γ) maps constituting the tracer state σ are realised as Erlang maps for efficient access. Trace event buffers κ coincide with actor mailboxes, while the remaining logic in algs. 1–3 translates directly to Erlang code. This one-to-one mapping gives us confidence that our implementation reflects the algorithm logic.

In *online* RV, monitors analyse trace events while the SuS executes, whereas the *offline* setting defers this analysis until the system terminates; [8, fig. 11 in app. B.1] captures the distinction in process tracing between online and offline instrumentation in our setting (showing trace buffers only). The online instrumentation set-up employs the tracing infrastructure offered by the EVM, which deposits SuS trace event messages in tracer mailboxes. Erlang tracing complies with tracing assumption A_1 , enabling RIARC to instrument disjoint SuS processes sets. We configure the EVM with the `set_on_spawn` flag so that spawned processes automatically inherit the same tracer as their parent [41]. This tracer assignment is atomic, meeting tracing assumption A_2 . We also use the `procs`, `send`, and `receive` tracing flags, which constrain the events emitted by the EVM to $\diamond$, $\star$, $!$, and $?$. The EVM enforces single-process tracing, *i.e.*, tracing assumption A_3 , and guarantees that $\diamond$ events of descendant processes are causally-ordered [127], *i.e.*, tracing assumption A_4 .

The offline counterpart differs only in its tracing layer, where events are read as *recorded* runs of the SuS. Recorded runs can be acquired externally, *e.g.* using DTrace [35] or LTTng [55], making it possible to monitor systems that execute outside of the EVM. Our bespoke offline tracing engine of [8, fig. 11b in app. B.1] fulfils tracing assumptions A_1 – A_4 . This is crucial

since it permits the *same* implementation of RIARC to be used in online and offline settings. Sec. 4.2 leverages this aspect to validate RIARC exhaustively using trace permutations.

We develop two versions of the TRACE, CLEAR, and PREEMPT functions of [8, alg. 5 in app. A] to standardise tracing for online and offline use. The overloads for online use access the EVM tracing via the Erlang built-in primitive `trace` [41]. The second set of overloads wraps around our offline tracing engine to replay files containing specifically-formatted trace events. Offline tracing relaxes tracing assumption A_4 , as recorded runs do not generally guarantee that the $\diamond$ events of descendant SuS processes are causally ordered. Our offline tracing logic relies on the PID information carried by $\diamond$ events to rearrange them and recover the causal ordering per tracing assumption A_4 . $\text{TRACE}(\iota_s, \iota_T)$ registers a tracer ι_T with the offline tracing engine, which maintains an event buffer for ι_T , together with a set of SuS PIDs that ι_T traces. A tracer can use TRACE with multiple SuS PIDs to register to obtain events for a process set, *i.e.*, tracing assumption A_1 . The tracing engine accumulates the events it reads from file in each tracer buffer and delivers events to the corresponding tracer mailbox once the casual ordering between $\diamond$ events of descendant SuS processes is established. Our offline tracing engine implements tracing inheritance (tracing assumption A_2) and enforces single-process tracing (tracing assumption A_3); [8, ex. 7 in app. B.1] sketches how the tracing engine uses its internal tracer buffers to deliver events to tracers.

4.2 Correctness

Conventional testing does not guarantee the absence of concurrency errors due to the different interleaved executions that may be possible [104]. While subjecting the system under test to high loads raises the likelihood of obtaining more coverage, this still depends on external factors, such as scheduling, which dictate the executions induced in practice. Controlling the conditions for concurrency testing requires a *systematic exploration* of all the interleaved executions [73]. In fact, it is *not the size* of the testing load that matters, but the choice of interleaved executions that exhaust the space of possible system states [14]. *Concuerror* [47] is a tool for systematic Erlang code testing. Unfortunately, we could not use *Concuerror* to test our RIARC implementation, as we were unable to integrate it with Erlang tracing.

We, nevertheless, adopt the systematic scheme advocated by *Concuerror*. Our approach uses the offline tracing tool described in sec. 4.1 to induce specific interleaved sequences for instrumentation set-ups, such as those of figs. 1b, 1c, and 2a. We obtain these sequences by taking all the sound (def. 1) event permutations of traces produced by the SuS. These sequences are then replayed by the offline tracing engine to systematically induce interleaved SuS executions. Our final RIARC implementation embeds further invariants besides those mentioned in sec. 3, *e.g.* the **assert** and **fail** statements in algs. 1 and 3. Readers are referred to [8, app. B.2] for the full list. We ascertain *trace soundness* for each SuS interleaving that is emulated. This is accomplished via the function `ANALYSEVT`, which we preload with monitors that assert the event sequence expected at each tracer. We also use identical tests in our empirical evaluation of sec. 5 under high loads. It is worth mentioning that while we systematically drive the execution of the SuS, we do not control the execution of tracers. Yet, we indirectly induce various dynamic tracer arrangements in the monitor DAG topology under the different groupings of SuS process sets that tracers instrument. For example, we fully instrument system depicted in fig. 2a in all its configurations, *e.g.* $C_1 = [T_{\{P\}} \rightsquigarrow \{P\}, T_{\{Q\}} \rightsquigarrow \{Q\}, T_{\{R\}} \rightsquigarrow \{R\}]$, $C_2 = [T_{\{P,Q\}} \rightsquigarrow \{P,Q\}, T_{\{R\}} \rightsquigarrow \{R\}]$, $\dots$, $C_5 = [T_{\{P,Q,R\}} \rightsquigarrow \{P,Q,R\}]$, as well as instrument it partially, *e.g.* $C_6 = [T_{\{P\}} \rightsquigarrow \{P\}]$, $C_7 = [T_{\{P,Q\}} \rightsquigarrow \{P,Q\}]$, *etc.* Each of these configurations, when individually paired with every fabricated interleaved execution of the SuS, indicate that our RIARC implementation and corresponding logic of sec. 3 is correct.

5 Empirical evaluation

We assess the feasibility of our RIARC implementation, confirming it safeguards the *responsive*, *resilient*, *message-driven*, and *elastic* attributes of the SuS. Sec. 4 targets a small selection of instrumentation set-ups to induce interleaved execution sequences and validate correctness exhaustively. We now employ *stress testing* [108] to investigate how RIARC performs in terms of the *runtime overhead* it exhibits. Our study focusses on *online* monitoring, as its overhead requirement is far more stringent than offline monitoring [62, 63, 21, 70]. We evaluate RIARC against inline instrumentation since the latter is regarded as the most efficient instrumentation technique [61, 60, 21]. This comparison establishes a solid basis for our results to be generalised reliably. We also compare RIARC to centralised instrumentation to confirm that the latter approach does not scale under typical loads.

Our experiments are extensive. We use two hardware platforms to model edge-case scenarios based on limited hardware and general-case scenarios using commodity hardware. The evaluation subjects inline, centralised, and RIARC instrumentation to high loads that go beyond the state of the art and use realistic workload profiles. We gauge overhead under three performance metrics, the *response time*, *memory consumption*, and *scheduler utilisation*, which are crucial for reactive systems [7, 108]. Our results confirm that the overhead RIARC induces is adequate for applications such as soft real-time systems [41, 93], where the latency requirement is typically in the order of seconds [91]. We also show that RIARC yields overhead comparable to inlining in settings exhibiting moderate concurrency.

5.1 Benchmarking tool

Benchmarking is standard practice for gauging runtime overhead in software [99, 76, 34]. Frameworks, including DaCapo [27] and Savina [83], offer limited concurrency, making them inapplicable to our case; see [8, app. C.1] for detailed reasons. Industry-proven *synthetic* load testing benchmarking tools cater to reactive systems, *e.g.* Apache JMeter [66], Tsung [114], and Basho Bench [22]. Their general-purpose design, however, necessarily treats systems as a black box by gathering metrics externally, which may impact measurement *precision* [7]. Moreover, these load testers generate standard workloads, *e.g.* Poisson processes [78, 101, 88], but lack others, *e.g.* load bursts, that replicate typical operation or induce edge-case stress.

We adopt BenchCRV [7], another synthetic load testing tool specific to RV benchmarking for reactive systems. BenchCRV sets itself apart from the tools mentioned above because it does not require external software (*e.g.*, a web server) to drive tests. Instead, BenchCRV produces different SuS models that *closely emulate* real-world software behaviour. These models are based on the master-worker paradigm [119]: a pervasive architecture in distributed (*e.g.* Big Data stream processing frameworks, render farms) and concurrent systems [128, 72, 54, 131]. Like Tsung and Basho Bench, BenchCRV exploits the lightweight EVM process model to generate highly-concurrent synthetic workloads.

BenchCRV creates master-worker models and induces workloads derived from configurable parameters. In these models, the master process spawns a series of workers and allocates tasks. The volume of workers per benchmark run is set via the parameter n . Each worker task consists of a *batch* of requests that the worker receives, processes, and echoes back to the master process. The amount of requests batched in one task is given by the parameter w . Workers terminate when all of their allotted tasks are processed and acknowledged by the master. BenchCRV creates workers based on *workload profiles*. A profile dictates how the master spreads its creation of workers along the loading timeline, t , given in seconds. BenchCRV supports three workload profiles based on ones typical in practice:

Steady models the SuS under stable workload (Poisson process).

Pulse models the SuS under gradually rising and falling workload (Normal distribution).

Burst models the SuS under stress due to workload spikes (Log-normal distribution).

BenchCRV records three performance metrics to give a multi-faceted view of system overhead:

Mean response time in milliseconds (ms), gauging monitoring latency effects on the SuS.

Mean memory consumption in GB, gauging monitoring memory pressure on the SuS.

Mean scheduler utilisation as a percentage of the total processing capacity, showing how monitors maximise the scheduler use.

The prevalent use of the master-worker paradigm, the veracity with which BenchCRV models systems, the range of realistic workload profiles, and the choice of runtime metrics it gathers make this tool ideal for our experiments. We refer readers to [8, app. C.2] and [7] for details.

5.2 Benchmark configuration

The BenchCRV master-worker models we generate take the role of the SuS in our experiments. We consider *edge-case* and *general-case* hardware platform set-ups for the following reasons:

P_E Edge-case captures platforms with *limited* hardware. It uses an Intel Core i7 M620 64-bit CPU with 8GB of memory, running Ubuntu 18.04 LTS and Erlang/OTP 22.2.1.

P_G General-case captures platforms with *commodity* hardware. It uses an Intel Core i9 9880H 64-bit CPU with 16GB of memory, running macOS 12.3.1 and Erlang/OTP 25.0.3.

The EVMs on platforms **P_E** and **P_G** are set with 4 and 16 scheduling threads, respectively. These scheduler settings coincide with the processors available on each SMP [12] platform. We also use the **P_E** and **P_G** platforms with two concurrency scenarios for reactive systems:

C_H High concurrency scenarios perform short-lived tasks, *e.g.* web apps that fulfil thousands of HTTP client requests by fetching static content or executing back-end commands.

C_M Moderate concurrency scenarios engage in long-running, computationally-intensive tasks, *e.g.* Big Data stream processing frameworks.

Our benchmark workloads match the hardware capacity afforded by **P_E** and **P_G**:

High concurrency benchmarks on **P_E** set $n = 100k$ workers and $w = 100$ work requests per worker. These generate $\approx (n \times w \text{ requests} \times w \text{ responses}) = 20M$ message exchanges between the master and worker processes, totalling $\approx (20M \times ! \text{ events} \times ? \text{ events}) = 40M$ analysable trace events. Platform **P_G** sets $n = 500k$ workers batched with $w = 100$ requests to produce $\approx 100M$ messages and $\approx 200M$ trace events. The high concurrency model **C_H** is studied in sec. 5.4.

Moderate concurrency benchmarks on **P_G** set $n = 5k$ workers and $w = 10k$ work requests per worker. These settings yield roughly the same number of trace events as on **P_G** with concurrency scenario **C_H**. The moderate concurrency model **C_M** is studied in sec. 5.5.

All experiments in secs. 5.4 and 5.5 use a total loading time of $t = 100s$. Each experiment consists of *ten* benchmarks that apply Steady, Pulse, and Burst workloads. We repeat every experiment *thrice* to obtain *negligible variability* and ensure the accuracy of our results; see [8, app. C.4] for a summary of these workloads and [8, app. C.5] for the precautions we take.

The hardware, OS, and Erlang versions of platforms **P_E** and **P_G**, combined with the workloads of concurrency scenarios **C_H** and **C_M** provide generality to our conclusions.

5.3 Instrumentation configuration

One challenge in conducting our experiments is the lack of RV monitoring tools targeting the EVM. To the best of our knowledge [63, tables 3 and 4], detectEr [71, 17, 18, 16, 69, 39]

is the only RV tool for Erlang that implements centralised outline instrumentation². We are unaware of inline RV tools besides [37] and [3, 4]. Since the former tool is *unavailable*, we use the latter, more recent work³. In our experiments, we instrument the master *and each* worker process in the SuS models generated from sec. 5.2 to exert the highest possible load and capture *worst-case* scenarios. BenchCRV annotates work requests and responses with a unique sequence number to account for each message in benchmark runs. We leverage this numbering to write specialised monitor replicas that ascertain the *soundness* of trace event sequences reported to every RV monitor linked with the master and workers; see [8, app. C.5] for details. Equally crucial, this runtime checking introduces a degree of *realistic* RV analysis slowdown that is *uniform* across all monitors in the inline, centralised, and RIARC monitoring set-ups. We empirically estimate this slowdown at $\approx 5\mu\text{s}$ per analysed event.

5.4 High concurrency benchmarks

We study runtime overhead in the high concurrency scenario C_H with two aims. First, we show the effect overhead has on the SuS as it executes. Specifically, we consider how the memory consumption and scheduler utilisation impact the *latency* a client of the SuS experiences, *e.g.* end-user or application. We use the edge-case platform P_E for these experiments; analogous results obtained on P_G are detailed in [8, app. C]. Our second goal targets the general-case platform P_G to assess the *scalability* of the instrumentation methods through their optimal use of the *additional* memory and scheduler capacity afforded by P_G .

The charts in secs. 5.4.1–5.4.3 plot performance metrics, *e.g.* memory consumption (*y*-axis) against the number of concurrent worker processes or the execution duration (*x*-axis). Since inline instrumentation prevents us from delineating the SuS and monitoring-induced runtime overhead, we follow the standard RV literature practice and include the *baseline* plots, *e.g.* [18, 71, 45, 37, 98, 113, 111]. Baseline plots show the *unmonitored* SuS to compare the relative overhead between each evaluated instrumentation method.

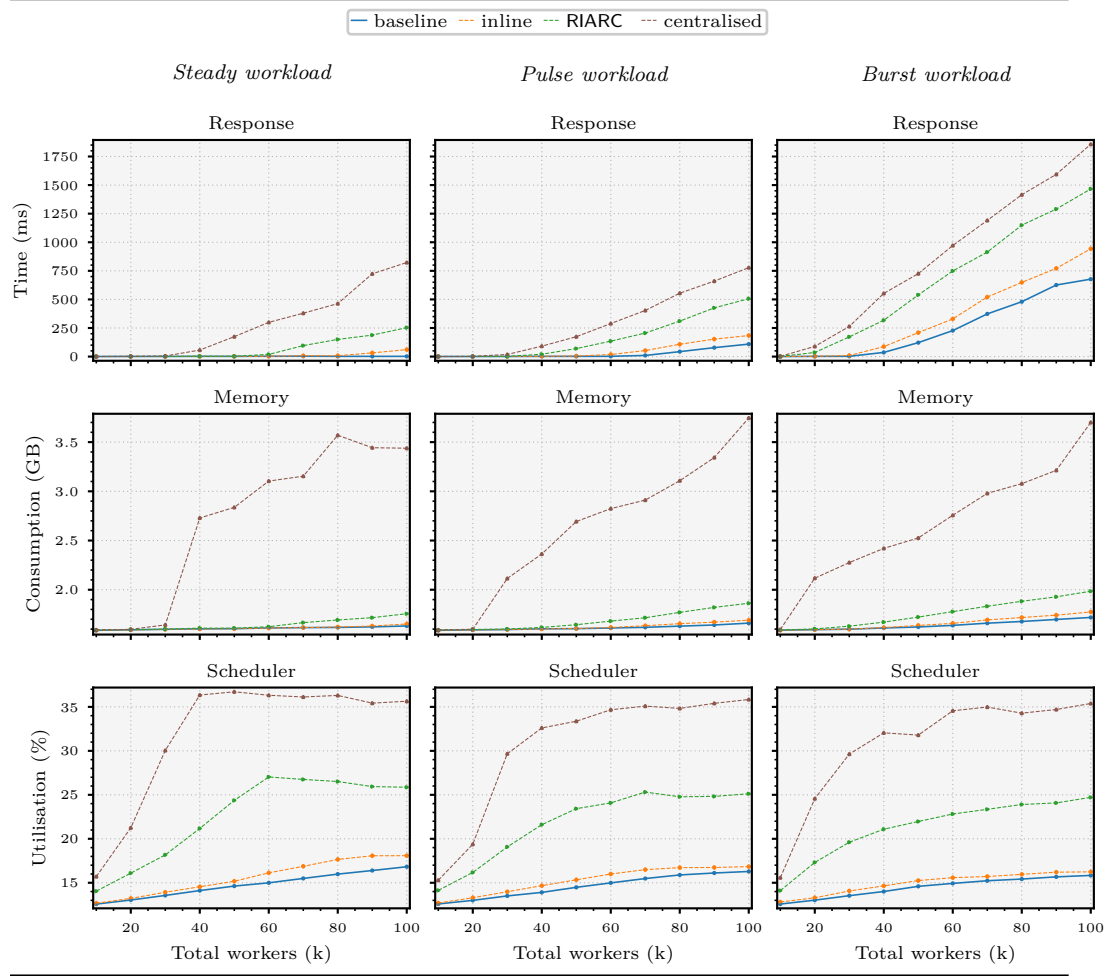
5.4.1 Instrumentation overhead

The first set of experiments isolates the instrumentation overhead induced on the SuS: this is the aggregated cost of tracing *and* reporting the traces soundly per def. 1 to RV monitors. Crucially, these experiments *omit monitors*, as we want to quantify the instrumentation overhead and understand its impact on the SuS. This enables us to focus on the differences between inlining—regarded as the most efficient instrumentation method [61, 60, 21]—and outlining. As far as we know [63, 70], outlining has *never* been used for decentralised RV in a *dynamic* setting such as ours. While we confirm that inline instrumentation uses less memory and scheduler capacity, RIARC dynamically scales and economises their use *without* adverse impact on the latency. In fact, the latency induced by RIARC is a mere 519ms higher than that of inline instrumentation at the peak stress-inducing loading point of 3.7k workers/s under Burst workloads. Our experiments indicate that centralised instrumentation manages resources poorly due to its inability to scale, increasing the chances of failure; see sec. 5.4.2.

Fig. 7 plots our results. Centralised instrumentation carries the largest overhead penalty. Regardless of the workload applied, it uses the most memory, $\approx 3.8\text{GB}$, highlighting its ineptitude to scale. This stems from the backlog of trace event messages that accumulate in the mailbox of the central tracer and is a manifestation of two aspects. First, the central

² <https://bitbucket.org/duncanatt/detector-lite>

³ <https://github.com/ScienceofComputerProgramming/SCIC0-D-22-00294>



■ **Figure 7** Isolated instrumentation overhead (*high* workload, 100k workers)

tracer does not consume events at the same rate worker processes produce them. Evidence of this *bottleneck* is visible as high scheduler utilisation in fig. 7 (bottom). This values settles at $\approx 36\%$ for the benchmarks with ≈ 40 k workers under the Steady workload and ≈ 60 k workers under Pulse and Burst workloads. Interpreting these $< 36\%$ scheduler usage values in isolation may suggest that centralised instrumentation has the potential to scale. However, its memory consumption plots in fig. 7 (middle) contradict this erroneous hypothesis.

By contrast, RIARC uses fewer resources to yield lower response times across the three workloads. The scheduler utilisation for RIARC slightly plateaus in the Steady (≈ 60 k workers) and Pulse (≈ 70 k workers) workload charts. This is not owed to scalability limitations of RIARC but to the intrinsic throttling instigated by the master process [119]. In fact, the plots for the baseline system and inline instrumentation in fig. 7 (middle) exhibit analogous signs of throttling. Even at a peak Burst workload of 3.7k workers/s, inline and RIARC instrumentation consume fairly similar amounts of memory, 1.7GB *vs.* 1.9GB, respectively.



■ **Figure 8** Instrumentation and RV monitoring overhead gap (*high* workload, 100k workers)

5.4.2 Monitoring overhead

Our second set of experiments extends the results of sec. 5.4.1 and quantifies the cost of RV monitoring. The *runtime monitoring* overhead combines the instrumentation and slowdown due to the RV analysis, established at $\approx 5\mu\text{s}$ per event in sec. 5.3 for our experiments. Fig. 8 plots the instrumentation (*instr.*) overhead from sec. 5.4.1 next to the runtime monitoring overhead (*mon.*). It shows that the RV analysis slowdown aggravates centralised monitoring to the point of crashing. Inline and RIARC monitoring are minimally affected. Our results also reveal that the instrumentation incurs the *major* overhead portion, not the RV analysis. Sec. 5.6 comments on this finding in the context of existing RV tools.

Fig. 8 plots our results under the Steady and Burst workloads; [8, fig. 14 in app. C.6.1] includes all three workloads. The charts for centralised monitoring exhibit a significant disparity between the instrumentation and runtime monitoring bar plots as the workload increases. This trend is consistent across both workloads in fig. 8. The lack of scalability of centralised monitoring in fig. 8 manifests as an increase in memory consumption but stabilised scheduler usage, as in fig. 7. Memory consumption and scheduler usage for centralised monitoring grow rapidly beyond $\approx 30\text{k}$ and $\approx 20\text{k}$ workers under the Steady and

Burst workloads, respectively. Bottlenecks led our experiments to crash (shown as missing bar plots in fig. 8). Crashes occur at $\approx 70k$ workers under the Steady and at $\approx 80k$ under Burst workload. By analysing the resulting dumps, we could attribute these crashes to memory exhaustion, which caused the EVM to fail. The dumps indicate severe memory pressure due to the vast backlog of trace event messages in the mailbox of the central tracer.

Inline and RIARC monitoring scale to accommodate the RV analysis slowdown. This is confirmed by cross-referencing the memory consumption and scheduler utilisation in fig. 8 for both monitoring methods. Each displays comparable overhead in their respective instrumentation and corresponding runtime monitoring bar plots. Fig. 8 (top) shows that inline and RIARC monitoring increase the latency, albeit for different reasons. The internal operation of RIARC enables us to deduce that its latency stems from message routing and dynamic tracer reconfiguration. Its scheduler utilisation plots support this observation. The latency due to inlining is a direct effect of RV analysis slowdown, provoked by the lock-step execution of monitors and the SuS. Other works, *e.g.* [45, 36], offer similar observations.

Dissecting our results uncovers further subtleties. The optimal scheduler utilisation of RIARC implies that its monitors are only active when triggered by trace events but remain idle otherwise. This inference is supported by the absence of sudden or continued memory growth for RIARC in fig. 8 (middle). The instrumentation and runtime monitoring latency bar plots for inline monitoring exhibit a growing pairwise gap that starts at $\approx 80k$ workers in fig. 8 (top right). The respective gap for RIARC at this mark is perceptibly lower. We credit this lower latency gap to outlining, which absorbs the slowdown effect of RV analyses. This leads us to conjecture that RIARC could accommodate monitors that perform richer RV analyses with minimal impact on the SuS. Our calculations from fig. 8 (top right) put the latency at 1093ms for inline monitoring *vs.* 1547ms for RIARC at a peak Burst workload of 3.7k workers/s: a 454ms difference, which is *lower* than the 519ms gap measured in sec. 5.4.1. Sec. 5.5 shows this gap is negligible in moderate concurrency scenarios.

5.4.3 Resource usage

We employ platform P_G with high concurrency C_H to confirm that our observations about inline and RIARC monitoring transfer to general cases. Secs. 5.4.1 and 5.4.2 deem centralised monitoring to be impractical. We, thus, omit it from the sequel; see [8, app. C.6.3] for results.

Our experiments now use 16 scheduling threads, $n=500k$ workers, and $w=100$ requests per worker, producing $\approx 100M$ messages and $\approx 200M$ trace events; [8, fig. 13 in app. C.4] render these Steady, Pulse, and Burst workload models. Secs. 5.4.1 and 5.4.2 bound the memory and scheduler metrics to the period the SuS executes to portray the *actual overhead* impact on the system. We refocus that view to assess the monitoring overhead in *its entirety*—from the point of SuS launch until monitors complete their RV analysis. Doing so reveals how inline and RIARC monitoring optimise the use of added memory and processing capacity. Results show that inline and RIARC monitoring are elastic and dynamically adapt to changes in the applied workloads; [8, app. C.6.3] confirms that centralised monitoring lacks this trait.

Fig. 9 gives a complete benchmark run under the Steady and Burst workloads. We relabel the x -axis with the benchmark duration and omit the response time plots since response time is inapplicable to these experiments (latency is an attribute of the SuS, not the monitors). In this run, the Steady workload generates a sustained load of $\approx 5k$ workers/s whereas Burst peaks at $\approx 17.8k$ workers/s under maximum load at $\approx 5s$; see [8, fig. 13 in app. C.4].

Fig. 9 (top) illustrates the memory consumption patterns for inline and RIARC monitoring, which exhibit *elasticity*. This elastic behaviour occurs at different points in the plots. Inline monitoring peaks at $\approx 3.7GB$ at $\approx 72s$ and RIARC at $\approx 5.7GB$ at $\approx 100s$ under the

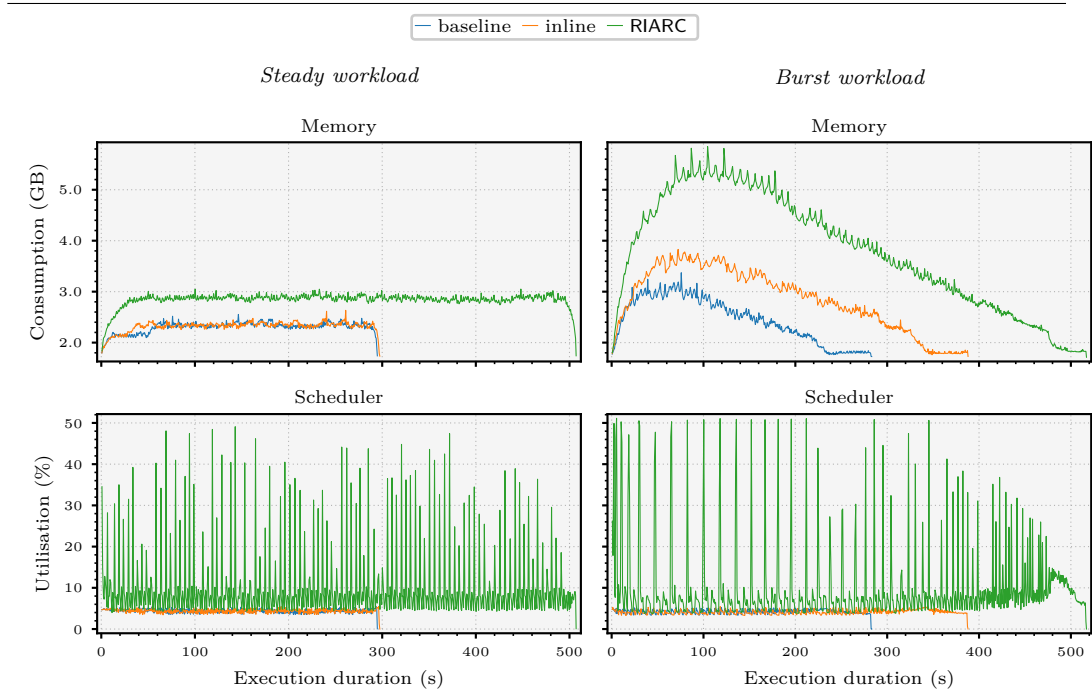
Burst workload. The memory consumption for both methods stabilises at around $\approx 36s$ under the Steady workload, with $\approx 2.3GB$ for inline and $\approx 2.7GB$ for RIARC monitoring. Elasticity in these methods is due to different reasons: it is intrinsic to inline monitoring (see sec. 1), whereas the RIARC spawns and garbage collects monitors on demand (secs. 3.1 and 3.6). These observations are certified by [8, fig. 16 in app. C.6.3] under the Pulse workload. Centralised monitoring is *insensitive* to the workload applied, as [8, figs. 17 and 18 in app. C.6.3] reconfirm.

The effect of dynamic message routing and tracer reconfiguration that RIARC performs is evident in the scheduler utilisation plots of fig. 9. Under the Steady and Burst workloads, scheduler utilisation oscillates continually due to the sustained influx of trace events. Oscillations corroborate our observation in sec. 5.4.2 about RIARC, namely, that monitors are activated by trace events but remain idle otherwise. Active monitor periods manifest as peaks in fig. 9. Idle periods, where monitors are placed in the EVM waiting queues, are reflected as regions with low and stable scheduler utilisation. These oscillations showcase the message-driven aspect of RIARC, which analyses events asynchronously. Inlining exhibits minimal scheduler utilisation oscillations due to its lock-step execution with the SuS.

5.5 Moderate concurrency benchmarks

Our last experiment studies moderate concurrency scenarios C_M . The general-case platform P_G sets $n = 5k$ workers and $w = 10k$ requests per worker, and uses 16 EVM schedulers. We show that under these loads, RIARC induces overhead on par with inline monitoring.

Moderate concurrency alters the execution of the master-worker model, compared to our benchmarks of secs. 5.4.1–5.4.3. In this set-up, the master creates most of its worker processes at the initial stage of benchmark runs and spends the remaining time allocating



■ **Figure 9** Inline and RIARC monitoring resource usage (*high* workload, 500k workers)

work requests. This change grows the request throughput, *e.g.* see [8, tbl. 5 in app. C.4]. One consequence is that centralised monitoring consistently crashes under the rapid accumulation of messages in its mailbox. We, thus, limit our study to inline and RIARC monitoring.

Tbl. 3 compares the results taken on platform P_G from sec. 5.4.3 with 500k workers (high concurrency, C_H) against the ones on P_G with 5k workers (moderate concurrency, C_M). The figures shown estimate the percentage overhead w.r.t. the baseline systems C_H and C_M at this *maximum* load. Our ensuing discussion is limited to the overhead under the Steady and Burst workloads since each respectively captures the SuS operation in *typical* and *severe* load conditions. Readers are referred to [8, fig. 20 in app. C.6.4] for the overhead comparison given in absolute metric values for the entirety of benchmark runs.

Tbl. 3 indicates that the memory consumption overhead due to inline monitoring is not affected under the Steady workload, which remains at 1% in both the high and moderate concurrency scenarios C_H and C_M . However, it decreases from 16% in C_H to 1% in C_M . We observe the opposite effect on the scheduler utilisation overhead for inline monitoring. For the moderate concurrency case C_M , the scheduler overhead under the Steady and Burst workloads increases to 3% and 4% respectively.

Tbl. 3 also shows that under the Steady workload, RIARC induces a 23% memory overhead in concurrency scenario C_H *vs.* 8% in concurrency scenario C_M , a decrease of 15%. Under the Burst workload, this overhead is reduced by 46%, from 56% in C_H to 10% in C_M . The scheduler utilisation overhead for RIARC from C_H to C_M also registers drops of $\approx 71\%$ under both Steady and Burst workloads. We attribute these overhead improvements to the lower number of worker processes the master creates in the moderate concurrency set-up, C_M . The long-running worker processes induce stability in the SuS. RIARC adapts to this change favourably by performing fewer trace event routing and tracer reconfigurations. The ramification of this adaptability is perceivable in the latency overhead discussed next.

RIARC inflates the latency overhead from 95% in C_H to 194% in C_M under the Steady workload (+99%), and from 97% in C_H to 190% in C_M under the Burst workload (+93%). However, RIARC induces *less latency* overhead than inline monitoring. Tbl. 3 reveals that the latency overhead for inline monitoring grows from 4% in the high concurrency set-up C_H to 246% in the moderate concurrency set-up C_M under the Steady workload (+242%). It also grows under the Burst workload, from 55% in C_H to 193% in C_M (+138%). In fact, our calculations confirm that the *absolute* response time for inline monitoring is slightly worse than that of RIARC in C_M : 116ms *vs.* 98ms under the Steady, and 182ms *vs.* 179ms under the Burst workloads respectively. This latency degradation for inline monitoring stems from the $\approx 5\mu s$ slowdown induced by the RV analysis, which results in frequent ‘pausing’ of worker processes. Monitors comprising richer analyses produce longer pauses in worker processes, which can degrade the response time further [45, 36, 68].

Concurrency	Workload	Response time %		Memory consumption %		Scheduler utilisation %	
		Inline	RIARC	Inline	RIARC	Inline	RIARC
C_H (500k)	Steady	4	95	1	23	0	123
	Burst	55	97	16	56	0	123
C_M (5k)	Steady	246	194	1	8	3	52
	Burst	193	190	1	10	4	50

Table 3 Percentage overhead on C_H (500k) and C_M (5k) w.r.t. baseline at *maximum* workload

5.6 Discussion

The RIARC scheduler utilisation in tbl. 3 is higher than the reported values for inline monitoring. This should not be construed as an inefficiency. From a reactive systems perspective, growth in the scheduler utilisation indicates *scalability*, as the low memory consumption in tbl. 3 affirms. RIARC benefits from the ample schedulers to improve the overall system response time *without* overtaxing the system. Indeed, [8, fig. 20 in app. C.6.4] demonstrates that the mean absolute scheduler utilisation in the benchmarks of sec. 5.5 is just $\approx 10\%$ under both the Steady and Burst workloads. Tbl. 3 shows that the reduction in latency makes RIARC comparable to inline monitoring in moderate concurrency scenarios.

Sec. 1 names *responsiveness* as a key reactive systems attribute [93]. RIARC prioritises responsiveness by isolating its monitors into asynchronous concurrent units. This design naturally exploits the available processing capacity of the host platform by maximising monitor *parallelism* when possible. Inline monitoring reaps fewer benefits in identical settings because its lock-step execution with the SuS robs it of potential parallelism gains.

Secs. 5.4.1–5.4.3 attest to the impracticality of centralised monitoring for reactive systems. Bottlenecks hinder its ability to scale, compelling it to consume inordinate amounts of memory, which can lead to failure, as sec. 5.4.2 shows. Despite these shortcomings, many RV tools in this setting use centralised monitoring, *e.g.* [49, 17, 125, 64, 80, 109, 71, 36, 40, 37, 2, 102].

6 Conclusion

Reactive software calls for instrumentation methods that uphold the responsive, resilient, message-driven, and elastic attributes of systems. This is attainable *only if* the instrumentation exhibits these qualities. Runtime verification imposes another demand on the instrumentation: the trace event sequences it reports to monitors must be *sound*, *i.e.*, traces do not omit events and preserve the ordering with which events occur locally at processes.

This paper presents RIARC, a novel decentralised instrumentation algorithm for outline monitors meeting these two demands. RIARC uses outline monitors to decouple the runtime analysis from system components, which minimises latency and promotes *responsiveness*. Outline monitors can fail independently of the system and each other to improve *resiliency*. RIARC gathers events non-invasively via a tracing infrastructure, making it *message-driven* and suited to cases where inlining is inapplicable. The algorithm is *elastic*: it reacts to specific events in the trace to instrument and garbage collect monitors on demand.

Our asynchronous setting complicates the instrumentation due to potential trace event loss or reordering. RIARC overcomes these challenges using a next-hop IP routing approach to rearrange and report events soundly to monitors. We validate RIARC by subjecting its corresponding Erlang implementation to rigorous systematic testing, confirming its correctness. This implementation is validated via extensive empirical experiments. These subject the implementation to large realistic workloads to ascertain its reactivity. Our experiments show that RIARC optimises its memory and scheduler usage to maintain latency feasible for soft real-time applications. We also compare RIARC to inline and centralised monitoring, revealing that it induces *comparable* latency to inlining under moderate concurrency.

Related work Other work on inlining besides that cited in sec. 1, *e.g.* [77, 24, 49, 48, 52], does not separate the instrumentation and runtime analysis. This view is commonplace in monolithic settings, where the instrumentation is often assumed to induce minimal runtime overhead. As a result, many inline approaches focus on the efficiency of the analysis but neglect the instrumentation cost (*e.g.* [62] attributes overhead solely to the analysis). These

arguments for monolithic systems are often ported to concurrent settings. For instance, [106, 125, 28, 45, 124, 65, 20] propose efficient runtime monitoring algorithms but do not account for, nor quantify, the overhead due to gathering trace events. Tools that measure the runtime overhead, such as [40, 36, 18, 33, 71, 132], coalesce the instrumentation and runtime analysis costs, making it difficult to gauge the source of inefficiencies. Some literature [38, 51] even extends the assumption about minimal instrumentation overhead to offline monitoring, stating that the instrumentation consists of ‘only’ capturing trace events. Sec. 5.4.1 shows this *not* to be the case. We are unaware of empirical studies such as ours that concretely distinguish between and quantify the instrumentation and runtime analysis overhead.

Sec. 5.6 remarks that centralised monitoring is used for concurrent runtime verification despite its evident limitations. One plausible reason for this is that the empirical scrutiny of such tools lacks proper benchmarking (*e.g.* [49, 17, 125, 64, 80]) or uses insufficient workloads that fail to expose the issues of centralised set-ups (*e.g.* [109, 71, 36, 40, 37, 2, 102]). Gathering inadequate metrics can also bias the interpretation of empirical data; see sec. 5.4.1. Works, such as [37, 18, 33, 123], consider the memory consumption and latency metrics. Our evaluation of inline, centralised, and RIARC monitoring uses (i) *combinations* of hardware and software, with (ii) two concurrency models that test *edge-case* and *general-case* scenarios, under (iii) *high* workloads that go beyond the state of the art, applying (iv) *realistic* workload profiles, interpreted against (v) *relevant* performance metrics that give a multi-faceted view of runtime overhead. To the best of our knowledge, this is generally not done in other studies, *e.g.* [113, 112, 46, 45, 118, 29, 105, 37, 40, 18, 49, 50, 52, 71, 58, 59, 26, 109, 96, 33].

Outline instrumentation decouples the execution of the SuS and monitor components in space (*i.e.*, isolated threads) and time (*i.e.*, asynchronous messaging). The tracing infrastructure outline instrumentation uses mirrors the publish-subscribe (Pub/Sub) pattern [128]. In this set-up, consumers subscribe to a *broker* that advertises events. Centralised instrumentation follows a Pub/Sub approach: the SuS produces trace events and deposits them into *one* global trace buffer that tracers receive from (see fig. 1b). Despite similarities, *e.g.* tracers register and deregister with the tracing infrastructure at runtime, RIARC differs from conventional Pub/Sub messaging in three fundamental aspects. Chiefly, Pub/Sub publishers are unaware of the subscribers interested in receiving messages because this bookkeeping task is appointed to the broker. By contrast, next-hop routing relies on knowing the *explicit* address of recipients to forward messages. Furthermore, in Pub/Sub messaging, subscribers do not communicate with publishers, whereas RIARC tracers exchange *direct* detach requests between one another to reorganise the choreography (refer to sec. 3.4). Lastly, Pub/Sub brokers are typically predefined and remain fixed, while trace partitioning *reconfigures* the tracing topology, creating and destroying brokers in reaction to dynamic changes in SuS.

One assumption we make about process tracing is A_4 , *i.e.*, tracing gathers the spawn events of parent processes before all the events of child processes. While A_4 induces a partial order over trace events, it is *weaker* than happened-before causality [94], as the events gathered from sets of child SuS processes need not be causally ordered. Demanding the latter condition would entail additional computation on the part of the tracing infrastructure and could increase runtime overhead. Maintaining minimal overhead is critical to our instrumentation because it preserves the responsiveness attribute of reactive systems. Tracing assumption A_4 and the RIARC logic detailed in sec. 3 guarantee trace soundness (def. 1), which suffices for RV monitoring. Since our work targets soft real-time systems [93, 91] scoped in a reliable messaging setting (see sec. 1), we do not tackle the problem of ensuring time-bounded causally-ordered message delivery [19] nor implement exactly-once delivery semantics [82]. We will address these challenges in future extensions of this work.

References

- 1 Francisco Lopez-Sancho Abraham. *Akka in Action*. Manning, 2023.
- 2 Luca Aceto, Antonis Achilleos, Elli Anastasiadi, and Adrian Francalanza. Monitoring Hyperproperties with Circuits. In *FORTE*, volume 13273 of *LNCS*, pages 1–10, 2022.
- 3 Luca Aceto, Antonis Achilleos, Duncan Paul Attard, Léo Exibard, Adrian Francalanza, and Anna Ingólfssdóttir. A Monitoring Tool for Linear-Time μ HML. In *COORDINATION*, volume 13271 of *LNCS*, pages 200–219, 2022.
- 4 Luca Aceto, Antonis Achilleos, Duncan Paul Attard, Léo Exibard, Adrian Francalanza, and Anna Ingólfssdóttir. A Monitoring Tool for Linear-time μ hml. *Sci. Comput. Program.*, 232:103031, 2024.
- 5 Luca Aceto, Antonis Achilleos, Adrian Francalanza, Anna Ingólfssdóttir, and Karoliina Lehtinen. Adventures in Monitorability: From Branching to Linear Time and Back Again. *PACMPL*, 3:52:1–52:29, 2019.
- 6 Luca Aceto, Antonis Achilleos, Adrian Francalanza, Anna Ingólfssdóttir, and Karoliina Lehtinen. An Operational Guide to Monitorability with Applications to Regular Properties. *Softw. Syst. Model.*, 20:335–361, 2021.
- 7 Luca Aceto, Duncan Paul Attard, Adrian Francalanza, and Anna Ingólfssdóttir. On Benchmarking for Concurrent Runtime Verification. In *FASE*, volume 12649 of *LNCS*, pages 3–23, 2021.
- 8 Luca Aceto, Duncan Paul Attard, Adrian Francalanza, and Anna Ingólfssdóttir. Runtime Instrumentation for Reactive Components. *CoRR*, abs/2406.19904, 2024.
- 9 Luca Aceto, Anna Ingólfssdóttir, Kim Guldstrand Larsen, and Jiří Srba. *Reactive Systems: Modelling, Specification and Verification*. Cambridge University Press, 2007.
- 10 Gul Agha, Ian A. Mason, Scott F. Smith, and Carolyn L. Talcott. A Foundation for Actor Computation. *JFP*, 7:1–72, 1997.
- 11 Gene M. Amdahl. Validity of the Single Processor Approach to Achieving Large Scale Computing Capabilities. In *AFIPS Spring Joint Computing Conference*, volume 30 of *AFIPS Conference Proceedings*, pages 483–485, 1967.
- 12 Joe Armstrong. *Programming Erlang: Software for a Concurrent World*. Pragmatic Bookshelf, 2007.
- 13 Joe Armstrong. Erlang. *Commun. ACM*, 53(9):68–75, 2010.
- 14 Stavros Aronis. *Effective Techniques for Stateless Model Checking*. PhD thesis, Uppsala University, Sweden, 2018.
- 15 Duncan Paul Attard, Luca Aceto, Antonis Achilleos, Adrian Francalanza, Anna Ingólfssdóttir, and Karoliina Lehtinen. Better Late than Never or: Verifying Asynchronous Components at Runtime. In *FORTE*, volume 12719 of *LNCS*, pages 207–225, 2021.
- 16 Duncan Paul Attard, Ian Cassar, Adrian Francalanza, Luca Aceto, and Anna Ingólfssdóttir. Introduction to Runtime Verification. In *Behavioural Types: from Theory to Tools*, Automation, Control and Robotics, pages 49–76. River, 2017.
- 17 Duncan Paul Attard and Adrian Francalanza. A Monitoring Tool for a Branching-Time Logic. In *RV*, volume 10012 of *LNCS*, pages 473–481, 2016.
- 18 Duncan Paul Attard and Adrian Francalanza. Trace Partitioning and Local Monitoring for Asynchronous Components. In *SEFM*, volume 10469 of *LNCS*, pages 219–235, 2017.
- 19 Roberto Baldoni, Achour Mostéfaoui, and Michel Raynal. Causal Delivery of Messages with Real-Time Data in Unreliable Networks. *Real Time Syst.*, 10(3):245–262, 1996.
- 20 Howard Barringer, Yliès Falcone, Klaus Havelund, Giles Reger, and David E. Rydeheard. Quantified Event Automata: Towards Expressive and Efficient Runtime Monitors. In *FM*, volume 7436 of *LNCS*, pages 68–84, 2012.
- 21 Ezio Bartocci, Yliès Falcone, Adrian Francalanza, and Giles Reger. Introduction to Runtime Verification. In *Lectures on Runtime Verification*, volume 10457 of *LNCS*, pages 1–33. Springer, 2018.
- 22 Basho. Bench, 2017. URL: https://github.com/basho/basho_bench.

- 23 David A. Basin, Felix Klaedtke, and Eugen Zalinescu. Failure-Aware Runtime Verification of Distributed Systems. In *FSTTCS*, volume 45 of *LIPIcs*, pages 590–603, 2015.
- 24 Andreas Bauer and Yliès Falcone. Decentralised LTL Monitoring. *FMSD*, 48:46–93, 2016.
- 25 André Bento, Jaime Correia, Ricardo Filipe, Filipe Araújo, and Jorge Cardoso. Automated Analysis of Distributed Tracing: Challenges and Research Directions. *J. Grid Comput.*, 19(1):9, 2021.
- 26 Shay Berkovich, Borzoo Bonakdarpour, and Sebastian Fischmeister. Runtime Verification with Minimal Intrusion through Parallelism. *FMSD*, 46:317–348, 2015.
- 27 Stephen M. Blackburn, Robin Garner, Chris Hoffmann, Asjad M. Khan, Kathryn S. McKinley, Rotem Bentzur, Amer Diwan, Daniel Feinberg, Daniel Frampton, Samuel Z. Guyer, Martin Hirzel, Antony L. Hosking, Maria Jump, Han Bok Lee, J. Eliot B. Moss, Aashish Phansalkar, Darko Stefanovic, Thomas VanDrunen, Daniel von Dincklage, and Ben Wiedermann. The DaCapo Benchmarks: Java Benchmarking Development and Analysis. In *OOPSLA*, pages 169–190, 2006.
- 28 Eric Bodden. The Design and Implementation of Formal Monitoring Techniques. In *OOPSLA Companion*, pages 939–940, 2007.
- 29 Eric Bodden, Laurie J. Hendren, Patrick Lam, Ondrej Lhoták, and Nomair A. Naeem. Collaborative Runtime Verification with Tracematches. *J. Log. Comput.*, 20:707–723, 2010.
- 30 Borzoo Bonakdarpour, Pierre Frgaignaud, Sergio Rajsbaum, David A. Rosenbluth, and Corentin Travers. Decentralized Asynchronous Crash-Resilient Runtime Verification. In *CONCUR*, volume 59 of *LIPIcs*, pages 16:1–16:15, 2016.
- 31 Jonas Bonér, Dave Farley, Roland Kuhn, and Martin Thompson. The Reactive Manifesto. Technical report, 2014.
- 32 Jonas Bonér and Viktor Klang. Reactive Programming *vs.* Reactive Systems. Technical report, Lightbend Inc., 2016.
- 33 Christian Bartolo Burlò, Adrian Francalanza, and Alceste Scalas. On the Monitorability of Session Types, in Theory and Practice. In *ECOOP*, volume 194 of *LIPIcs*, pages 20:1–20:30, 2021.
- 34 Rajkumar Buyya, James Broberg, and Andrzej M. Goscinski. *Cloud Computing: Principles and Paradigms*. Wiley-Blackwell, 2011.
- 35 Bryan Cantrill. Hidden in Plain Sight. *ACM Queue*, 4:26–36, 2006.
- 36 Ian Cassar and Adrian Francalanza. On Synchronous and Asynchronous Monitor Instrumentation for Actor-based Systems. In *FOCLASA*, volume 175 of *EPTCS*, pages 54–68, 2014.
- 37 Ian Cassar and Adrian Francalanza. On Implementing a Monitor-Oriented Programming Framework for Actor Systems. In *IFM*, volume 9681 of *LNCS*, pages 176–192, 2016.
- 38 Ian Cassar, Adrian Francalanza, Luca Aceto, and Anna Ingólfssdóttir. A Survey of Runtime Monitoring Instrumentation Techniques. In *PrePostiFM*, volume 254 of *EPTCS*, pages 15–28, 2017.
- 39 Ian Cassar, Adrian Francalanza, Duncan Paul Attard, Luca Aceto, and Anna Ingólfssdóttir. A Suite of Monitoring Tools for Erlang. In *RV-CuBES*, volume 3 of *Kalpa Publications in Computing*, pages 41–47, 2017.
- 40 Ian Cassar, Adrian Francalanza, and Simon Said. Improving Runtime Overheads for detectEr. In *FESCA*, volume 178 of *EPTCS*, pages 1–8, 2015.
- 41 Francesco Cesarini and Simon Thompson. *Erlang Programming: A Concurrent Approach to Software Development*. O’Reilly Media, 2009.
- 42 Bernadette Charron-Bost, Friedemann Mattern, and Gerard Tel. Synchronous, Asynchronous, and Causally Ordered Communication. *Distributed Comput.*, 9(4):173–191, 1996.
- 43 Natalia Chechina, Kenneth MacKenzie, Simon J. Thompson, Phil Trinder, Olivier Boudeville, Viktoria Fordós, Csaba Hoch, Amir Ghaffari, and Mario Moro Hernandez. Evaluating Scalable Distributed Erlang for Scalability and Reliability. *IEEE Trans. Parallel Distributed Syst.*, 28(8):2244–2257, 2017.

- 44 Feng Chen and Grigore Rosu. Java-MOP: A Monitoring Oriented Programming Environment for Java. In *TACAS*, volume 3440 of *LNCSS*, pages 546–550, 2005.
- 45 Feng Chen and Grigore Rosu. Mop: An Efficient and Generic Runtime Verification Framework. In *OOPSLA*, pages 569–588, 2007.
- 46 Feng Chen and Grigore Rosu. Parametric Trace Slicing and Monitoring. In *TACAS*, volume 5505 of *LNCSS*, pages 246–261, 2009.
- 47 Maria Christakis, Alkis Gotovos, and Konstantinos Sagonas. Systematic Testing for Detecting Concurrency Errors in Erlang Programs. In *ICST*, pages 154–163. IEEE Computer Society, 2013.
- 48 Christian Colombo and Yliès Falcone. Organising LTL Monitors over Distributed Systems with a Global Clock. *FMSD*, 49:109–158, 2016.
- 49 Christian Colombo, Adrian Francalanza, and Rudolph Gatt. Elarva: A Monitoring Tool for Erlang. In *RV*, volume 7186 of *LNCSS*, pages 370–374, 2011.
- 50 Christian Colombo, Adrian Francalanza, Ruth Mizzi, and Gordon J. Pace. polyLarva: Runtime Verification with Configurable Resource-Aware Monitoring Boundaries. In *SEFM*, volume 7504 of *LNCSS*, pages 218–232, 2012.
- 51 Christian Colombo and Gordon J. Pace. *Runtime Verification - A Hands-On Approach in Java*. Springer, 2022.
- 52 Christian Colombo, Gordon J. Pace, and Gerardo Schneider. LARVA — Safer Monitoring of Real-Time Java Programs (Tool Paper). In *SEFM*, pages 33–37, 2009.
- 53 Markus Dahm. Byte Code Engineering with the BCEL API. Technical report, Java Informationstage 99, 2001.
- 54 Jeffrey Dean and Sanjay Ghemawat. MapReduce: Simplified Data Processing on Large Clusters. *Commun. ACM*, 51:107–113, 2008.
- 55 Mathieu Desnoyers and Michel Dagenais. The LTTng Tracer: A Low Impact Performance and Behavior Monitor for GNU/Linux. Technical report, École Polytechnique de Montréal, 2006.
- 56 Jean Dollimore, Tim Kindberg, and George Coulouris. *Distributed Systems: Concepts and Design*. Addison-Wesley, 2005.
- 57 Eclipse/IBM. OpenJ9, 2021. URL: <https://www.eclipse.org/openj9>.
- 58 Antoine El-Hokayem and Yliès Falcone. Monitoring Decentralized Specifications. In *ISSTA*, pages 125–135, 2017.
- 59 Antoine El-Hokayem and Yliès Falcone. On the Monitoring of Decentralized Specifications: Semantics, Properties, Analysis, and Simulation. *ACM Trans. Softw. Eng. Methodol.*, 29:1:1–1:57, 2020.
- 60 Úlfar Erlingsson. *The Inlined Reference Monitor Approach to Security Policy Enforcement*. PhD thesis, Cornell University, US, 2004.
- 61 Úlfar Erlingsson and Fred B. Schneider. SASI Enforcement of Security Policies: A Retrospective. In *NSPW*, pages 87–95, 1999.
- 62 Yliès Falcone, Klaus Havelund, and Giles Reger. A Tutorial on Runtime Verification. In *Engineering Dependable Software Systems*, volume 34 of *NATO Science for Peace and Security Series, D: Information and Communication Security*, pages 141–175. IOS Press, 2013.
- 63 Yliès Falcone, Srdan Krstic, Giles Reger, and Dmitriy Traytel. A Taxonomy for Classifying Runtime Verification Tools. *STTT*, 23:255–284, 2021.
- 64 Yliès Falcone, Hosein Nazarpour, Saddek Bensalem, and Marius Bozga. Monitoring Distributed Component-Based Systems. In *FACS*, volume 13077 of *LNCSS*, pages 153–173, 2021.
- 65 Yliès Falcone, Hosein Nazarpour, Mohamad Jaber, Marius Bozga, and Saddek Bensalem. Tracing Distributed Component-Based Systems, a Brief Overview. In *RV*, volume 11237 of *LNCSS*, pages 417–425, 2018.
- 66 Apache Software Foundation. JMeter, 2020. URL: <https://jmeter.apache.org>.
- 67 Pierre Fraigniaud, Sergio Rajsbaum, and Corentin Travers. On the Number of Opinions Needed for Fault-Tolerant Run-Time Monitoring in Distributed Systems. In *RV*, volume 8734 of *LNCSS*, pages 92–107, 2014.

- 68 Adrian Francalanza. A Theory of Monitors. *Inf. Comput.*, 281:104704, 2021.
- 69 Adrian Francalanza, Luca Aceto, Antonis Achilleos, Duncan Paul Attard, Ian Cassar, Dario Della Monica, and Anna Ingólfssdóttir. A Foundation for Runtime Monitoring. In *RV*, volume 10548 of *LNCS*, pages 8–29, 2017.
- 70 Adrian Francalanza, Jorge A. Pérez, and César Sánchez. Runtime Verification for Decentralised and Distributed Systems. In *Lectures on RV*, volume 10457 of *LNCS*, pages 176–210. Springer, 2018.
- 71 Adrian Francalanza and Aldrin Seychell. Synthesising Correct Concurrent Runtime Monitors. *FMSD*, 46:226–261, 2015.
- 72 Sukumar Ghosh. *Distributed Systems: An Algorithmic Approach*. CRC, 2014.
- 73 Patrice Godefroid. Model Checking for Programming Languages using Verisort. In *POPL*, pages 174–186. ACM Press, 1997.
- 74 Susanne Graf, Doron A. Peled, and Sophie Quinton. Monitoring Distributed Systems Using Knowledge. In *FORTE*, volume 6722 of *LNCS*, pages 183–197, 2011.
- 75 Susan L. Graham, Peter B. Kessler, and Marshall K. McKusick. gprof: A Call Graph Execution Profiler. In *SIGPLAN Symposium on Compiler Construction*, pages 120–126. ACM, 1982.
- 76 Jim Gray. *The Benchmark Handbook for Database and Transaction Processing Systems*. Morgan Kaufmann, 1993.
- 77 Radu Grigore, Dino Distefano, Rasmus Lerchedahl Petersen, and Nikos Tzevelekos. Runtime Verification Based on Register Automata. In *TACAS*, volume 7795 of *LNCS*, pages 260–276, 2013.
- 78 Duncan A. Grove and Paul D. Coddington. Analytical Models of Probability Distributions for MPI Point-to-Point Communication Times on Distributed Memory Parallel Computers. In *ICA3PP*, volume 3719 of *LNCS*, pages 406–415, 2005.
- 79 Eric A. Hall. *Internet Core Protocols: The Definitive Guide*. O’Reilly Media, 2000.
- 80 Klaus Havelund, Giles Reger, Daniel Thoma, and Eugen Zalinescu. Monitoring Events that Carry Data. In *Lectures on Runtime Verification*, volume 10457 of *LNCS*, pages 61–102. Springer, 2018.
- 81 Carl Hewitt, Peter Boehler Bishop, and Richard Steiger. A Universal Modular ACTOR Formalism for Artificial Intelligence. In *IJCAI*, pages 235–245, 1973.
- 82 Yongqiang Huang and Hector Garcia-Molina. Exactly-Once Semantics in a Replicated Messaging System. In *ICDE*, pages 3–12. IEEE Computer Society, 2001.
- 83 Shams Mahmood Imam and Vivek Sarkar. Savina - An Actor Benchmark Suite: Enabling Empirical Evaluation of Actor Libraries. In *AGERE!@SPLASH*, pages 67–80, 2014.
- 84 Justin Iurman, Frank Brockners, and Benoit Donnet. Towards Cross-Layer Telemetry. In *ANRW*, pages 15–21. ACM, 2021.
- 85 Richard Jones, Antony Hosking, and Eliot Moss. *The Garbage Collection Handbook: The Art of Automatic Memory Management*. CRC, 2020.
- 86 Nicolai M. Josuttis. *SOA in Practice: The Art of Distributed System Design: Theory in Practice*. O’Reilly Media, 2007.
- 87 Saša Jurić. *Elixir in Action*. Manning, 2019.
- 88 Bill Kayser. What is the expected distribution of website response times?, 2017. URL: <https://blog.newrelic.com/engineering/expected-distributions-website-response-times>.
- 89 Gregor Kiczales, Erik Hilsdale, Jim Hugunin, Mik Kersten, Jeffrey Palm, and William G. Griswold. An Overview of AspectJ. In *ECOOP*, volume 2072 of *LNCS*, pages 327–353, 2001.
- 90 Moonzoo Kim, Mahesh Viswanathan, Sampath Kannan, Insup Lee, and Oleg Sokolsky. JavaMaC: A Run-Time Assurance Approach for Java Programs. *FMSD*, 24:129–155, 2004.
- 91 Hermann Kopetz. *Real-Time Systems: Design Principles for Distributed Embedded Applications (Real-Time Systems Series)*. Springer, 2011.
- 92 Ajay D. Kshemkalyani and Mukesh Singhal. *Distributed Computing: Principles, Algorithms, and Systems*. Cambridge University Press, 2011.
- 93 Roland Kuhn, Brian Hanafée, and Jamie Allen. *Reactive Design Patterns*. Manning, 2016.

- 94 Leslie Lamport. Time, Clocks, and the Ordering of Events in a Distributed System. *Commun. ACM*, 21(7):558–565, 1978.
- 95 Leslie Lamport, Robert E. Shostak, and Marshall C. Pease. The Byzantine Generals Problem. *ACM Trans. Program. Lang. Syst.*, 4:382–401, 1982.
- 96 Julien Lange and Nobuko Yoshida. Verifying Asynchronous Interactions via Communicating Session Automata. In *CAV*, volume 11561 of *LNCS*, pages 97–117, 2019.
- 97 Paul Lavery and Takuo Watanabe. An Actor-Based Runtime Monitoring System for Web and Desktop Applications. In *SNPD*, pages 385–390. IEEE Computer Society, 2017.
- 98 Philipp Lengauer, Verena Bitto, Hanspeter Mössenböck, and Markus Weninger. A Comprehensive Java Benchmark Study on Memory and Garbage Collection Behavior of DaCapo, DaCapo Scala, and SPECjvm2008. In *ICPE*, pages 3–14, 2017.
- 99 Bryon C. Lewis and Albert E. Crews. The Evolution of Benchmarking as a Computer Performance Evaluation Technique. *MIS Q.*, 9:7–16, 1985.
- 100 Jay Ligatti, Lujo Bauer, and David Walker. Edit Automata: Enforcement Mechanisms for Run-Time Security Policies. *Int. J. Inf. Sec.*, 4:2–16, 2005.
- 101 Zhen Liu, Nicolas Niclausse, and César Jalpa-Villanueva. Traffic Model and Performance Evaluation of Web Servers. *Perform. Evaluation*, 46:77–100, 2001.
- 102 Qingzhou Luo and Grigore Rosu. EnforceMOP: A Runtime Property Enforcement System for Multithreaded Programs. In *ISSTA*, pages 156–166, 2013.
- 103 Deep Medhi and Karthik Ramasamy. Chapter 3 - routing protocols: Framework and principles. In *Network Routing (Second Edition)*, The Morgan Kaufmann Series in Networking, pages 64–113. Morgan Kaufmann, 2018.
- 104 Silvana M. Melo, Jeffrey C. Carver, Paulo S. L. Souza, and Simone R. S. Souza. Empirical Research on Concurrent Software Testing: A Systematic Mapping Study. *Inf. Softw. Technol.*, 105:226–251, 2019.
- 105 Patrick O’Neil Meredith, Dongyun Jin, Dennis Griffith, Feng Chen, and Grigore Rosu. An Overview of the MOP Runtime Verification Framework. *STTT*, 14:249–289, 2012.
- 106 Patrick O’Neil Meredith and Grigore Rosu. Efficient Parametric Runtime Verification with Deterministic String Rewriting. In *ASE*, pages 70–80, 2013.
- 107 Microsoft. MSDN, 2021. URL: <https://msdn.microsoft.com>.
- 108 Ian Molyneaux. *The Art of Application Performance Testing 2e*. O’Reilly Media, 2014.
- 109 Menna Mostafa and Borzoo Bonakdarpour. Decentralized Runtime Verification of LTL Specifications in Distributed Systems. In *IPDPS*, pages 494–503, 2015.
- 110 Nicholas Nethercote and Julian Seward. Valgrind: A Framework for Heavyweight Dynamic Binary Instrumentation. In *PLDI*, pages 89–100. ACM, 2007.
- 111 Rumyana Neykova. *Multiparty Session Types for Dynamic Verification of Distributed Systems*. PhD thesis, Imperial College London, UK, 2017.
- 112 Rumyana Neykova and Nobuko Yoshida. Let it Recover: Multiparty Protocol-Induced Recovery. In *CC*, pages 98–108, 2017.
- 113 Rumyana Neykova and Nobuko Yoshida. Multiparty Session Actors. *LMCS*, 13, 2017.
- 114 Nicolas Niclausse. Tsung, 2017. URL: <http://tsung.erlang-projects.org>.
- 115 Scott Oaks. *Java Performance: In-Depth Advice for Tuning and Programming Java 8, 11, and Beyond*. CRC, 2020.
- 116 Martin Odersky, Lex Spoon, Bill Venners, and Frank Sommers. *Programming in Scala*. Artima Inc., 2021.
- 117 Kevin Quick. Thespian, 2020. URL: <https://thespianpy.com/doc>.
- 118 Giles Reger, Helena Cuenca Cruz, and David E. Rydeheard. MarQ: Monitoring at Runtime with QEA. In *TACAS*, volume 9035 of *LNCS*, pages 596–610, 2015.
- 119 Sartaj Sahni and George L. Vairaktarakis. The Master-Slave Paradigm in Parallel Computer and Industrial Settings. *J. Glob. Optim.*, 9:357–377, 1996.

- 120 Raja R. Sambasivan, Ilari Shafer, Jonathan Mace, Benjamin H. Sigelman, Rodrigo Fonseca, and Gregory R. Ganger. Principled Workflow-Centric Tracing of Distributed Systems. In *SoCC*, pages 401–414. ACM, 2016.
- 121 Torben Scheffel and Malte Schmitz. Three-Valued Asynchronous Distributed Runtime Verification. In *MEMOCODE*, pages 52–61, 2014.
- 122 Fred B. Schneider. Enforceable Security Policies. *ACM Trans. Inf. Syst. Secur.*, 3:30–50, 2000.
- 123 Joshua Schneider, David A. Basin, Frederik Brix, Srđan Krstić, and Dmitriy Traytel. Scalable Online First-Order Monitoring. *Int. J. Softw. Tools Technol. Transf.*, 23:185–208, 2021.
- 124 Koushik Sen, Grigore Rosu, and Gul Agha. Runtime Safety Analysis of Multithreaded Programs. In *ESEC / SIGSOFT FSE*, pages 337–346, 2003.
- 125 Koushik Sen, Grigore Rosu, and Gul Agha. Online Efficient Predictive Safety Analysis of Multithreaded Programs. *Int. J. Softw. Tools Technol. Transf.*, 8:248–260, 2006.
- 126 Koushik Sen, Abhay Vardhan, Gul Agha, and Grigore Rosu. Efficient Decentralized Monitoring of Safety in Distributed Systems. In *ICSE*, pages 418–427, 2004.
- 127 Eric Stenman. *The Erlang Runtime System*. 2023.
- 128 Sasu Tarkoma. *Overlay Networks: Toward Information Networking*. Auerbach, 2010.
- 129 The Pony Team. Ponylang, 2021. URL: <https://tutorial.ponylang.io>.
- 130 Ulf T. Wiger, Gösta Ask, and Kent Boortz. World-Class Product Certification using Erlang. *ACM SIGPLAN Notices*, 37(12):25–34, 2002.
- 131 Jiali Yao, Zhigeng Pan, and Hongxin Zhang. A Distributed Render Farm System for Animation Production. In *ICEC*, volume 5709 of *LNCS*, pages 264–269, 2009.
- 132 Teng Zhang, Greg Eakman, Insup Lee, and Oleg Sokolsky. Overhead-Aware Deployment of Runtime Monitors. In *RV*, volume 11757 of *LNCS*, pages 375–381, 2019.