



An Interference-Free Programming Model for Network Objects

Bertrand Meyer, Mischael Schill, Christopher M. Poskitt

► To cite this version:

Bertrand Meyer, Mischael Schill, Christopher M. Poskitt. An Interference-Free Programming Model for Network Objects. 18th International Conference on Coordination Languages and Models (COORDINATION), Jun 2016, Heraklion, Greece. pp.227-244, 10.1007/978-3-319-39519-7_14. hal-01631712

HAL Id: hal-01631712

<https://inria.hal.science/hal-01631712>

Submitted on 9 Nov 2017

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License

An Interference-Free Programming Model for Network Objects

Mischael Schill¹, Christopher M. Poskitt², and Bertrand Meyer^{3,4,5}

¹ Department of Computer Science, ETH Zürich, Switzerland

² Singapore University of Technology and Design, Singapore

³ Politecnico di Milano, Italy

⁴ Innopolis University, Russia

⁵ Université Paul Sabatier, France

Abstract. Network objects are a simple and natural abstraction for distributed object-oriented programming. Languages that support network objects, however, often leave synchronization to the user, along with its associated pitfalls, such as data races and the possibility of failure. In this paper, we present D-SCOOP, a distributed programming model that allows for interference-free and transaction-like reasoning on (potentially multiple) network objects, with synchronization handled automatically, and network failures managed by a compensation mechanism. We achieve this by leveraging the runtime semantics of a multi-threaded object-oriented concurrency model, directly generalizing it with a message-based protocol for efficiently coordinating remote objects. We present our pathway to fusing these contrasting but complementary ideas, and evaluate the performance overhead of the automatic synchronization in D-SCOOP, finding that it comes close to—or outperforms—explicit locking-based synchronization in Java RMI.

1 Introduction

Inter-device communication is becoming ubiquitous, and the number of connected devices is growing everyday. With this ubiquity comes an increasing demand for programmers to be able to write reliable distributed software, yet this is no simple task. Challenging errors such as data races and deadlocks can arise from subtle mistakes in synchronization code; and the failure of individual devices can block whole systems in the absence of appropriate recovery protocols.

Various language abstractions have been proposed to make it easier to write distributed programs. One such abstraction, natural for the object-oriented paradigm, is that of *network objects* [2]: objects whose methods can be invoked over a network. By handling communication in method calls, network objects allow for local and remote objects to be treated uniformly, without regard to where they are physically located. In principle an elegant generalization; in practice, languages supporting them are often lightweight on synchronization, leaving the user to manage it explicitly, and potentially exposing them to the aforementioned errors.

Many of these pitfalls of synchronization are not unique to distribution: they occur in multi-threaded concurrent programming too. Several languages and libraries attempt to make it easier and safer to write concurrent programs, providing their users with

high-level abstractions as diverse as transactional memory [23], block-dispatching [10], actors [1], and active objects [14]. Given the many shared synchronization challenges, a number of these abstractions have been successfully applied across novel distributed programming approaches, exemplified by languages such as Creol [12], JCoBox [21], and AmbientTalk [6].

A family of concurrency abstractions that (until the present paper) had not been generalized to distributed programming were those provided by SCOOP [25], despite their potential to naturally complement the network objects abstraction and to address some of its shortcomings. SCOOP is an object-oriented concurrency model that provides data-race freedom by construction, and strong guarantees about the order in which requests are executed by concurrently running processes. The synchronization provided by its runtime automatically excludes interfering calls, making it possible to reason independently about different blocks of code over multiple concurrent objects, almost as if each block is “sequential”. The ethos of the SCOOP approach—stick to the mental models programmers already know well (in this case sequential programming)—is aligned with that of the network objects abstraction, and challenged us to explore how they could complement the strengths of each other.

Our Contributions. The main outcome of this paper is D-SCOOP, a distributed programming model resulting from the fusion of the network objects abstraction with the runtime of the SCOOP concurrency model. The strong reasoning guarantees of the latter are directly generalized to provide interference-free and transaction-like reasoning on (potentially multiple) network objects, without the programmer having to worry about how to achieve it. The basis of this fusion is a message-based protocol for coordinating remote objects, which includes an efficient and novel two-phase locking algorithm for establishing the SCOOP order guarantees without prolonged periods of blocking. Furthermore, we adapt from transactional memory the recovery technique of compensations, in order for D-SCOOP to be able to restore consistency when clients fail mid-computation. This paper presents our pathway to fusing these independent, but complementary ideas. We furthermore evaluate a prototype implementation of D-SCOOP to investigate the performance overhead of its automatic synchronization mechanisms, finding that they come close to—and in some circumstances outperform—explicit locking-based synchronization in the Java RMI realization of network objects.

For the distributed programming community, this paper presents a programming model with interference-free and transaction-like reasoning for distributed objects, and a runtime that effectively handles the synchronization. For the SCOOP community, it presents a generalization of the classical SCOOP concurrency model to distribution in a way that maintains the guarantees of the core abstractions. For language designers, it presents a simple yet effective distributed programming abstraction (and descriptions of how we realized it) that could be transferred to other object-oriented languages.

Plan of the Paper. After introducing the necessary technical background of network objects and SCOOP (Section 2), we show how they fuse together in D-SCOOP, our distributed programming model (Section 3). We go into more depth on how objects are controlled to avoid interference (Section 4) and how compensation helps in managing failure (Section 5). Our prototype is then evaluated against Java RMI (Section 6), before we review some related work (Section 7) and conclude (Section 8).

2 Background: Network Objects and SCOOP

Our work combines networks objects—a distributed programming abstraction—with SCOOP, a concurrency model that handles synchronization in its runtime and provides strong reasoning guarantees. We present the necessary technical background of these concepts in the context of a running example.

Network Objects. A *network object* is an object whose methods can be invoked over a network. The abstraction is a simple but natural generalization of standard objects to distributed contexts: the programmer interacts with their interfaces in the same way as before, and without regard to where the object is physically located. Communication is handled in the method calls, and is typically synchronous to mimic regular method calls. Network objects first appeared in Modula-3 [2], and have since strongly influenced Java’s Remote Method Invocation (RMI) API as well as the Common Object Request Broker Architecture (CORBA) standard.

While implementations of network objects vary, the abstraction is typically lightweight on synchronization, leaving this difficulty to the user, to the point that multiple clients can concurrently execute the same method (introducing the possibility of data races). Simple mechanics such as `synchronized` in Java are not always sufficient to ensure atomicity. Consider for example the simple bank account transfer method in Listing 1, which allows some client to transfer an amount (`am`) of money from a source (`s`) account to a target (`t`) account. If the system is single-threaded and the accounts are local, then the method is correct. If the accounts can be accessed concurrently, then locks or other measures are required to ensure the atomicity of transfer. If however the accounts are remote and can be accessed concurrently as network objects, then we must adapt again.

```
transfer (s, t: ACCOUNT;
        am: NATURAL)
do
  if s.balance >= am then
    s.set_balance (s.balance - am)
    t.set_balance (t.balance + am)
  else -- Notify user
  end
end

transfer (s, t: separate ACCOUNT;
        am: NATURAL)
do
  if s.balance >= am then
    s.set_balance (s.balance - am)
    t.set_balance (t.balance + am)
  else -- Notify user
  end
end
```

Listing 1: Bank account transfer methods: sequential (left) and in SCOOP (right)

One solution is to use locks and expose them as network objects, but this poses risk, e.g. if a client loses its connection before having a chance to release its locks. Another solution is to hide the synchronization within additional methods in the account class, but this is still challenging to implement without introducing concurrency errors such as races or deadlocks. Either way, the simplicity of the network object abstraction suffers with the complexity of synchronizing correctly; hence our aim to elegantly integrate it with a concurrency model that can manage such complexity in its runtime.

SCOOP. SCOOP [25] is a concurrent object-oriented programming model that aims to preserve the well-understood modes of reasoning enjoyed by sequential programs, such as pre- and postcondition reasoning over blocks of code. Programmers are provided with simple abstractions for expressing concurrency, with the runtime itself responsible for correctly handling synchronization. We describe SCOOP in the context of its principal implementation for Eiffel [8], but remark that the ideas generalize to other object-oriented languages (e.g. Java [24]).

In SCOOP, every object is associated with a *process* (which we call its *handler*), a concurrent thread of execution with the exclusive right to call methods on the objects it handles. In this context, object references may point to objects with the same handler (*non-separate* objects) or to objects with distinct handlers (*separate* objects). Method calls on non-separate objects are executed immediately by the shared process. To make a call on a separate object, however, a *request* must be sent to the handler of that object to process it: if the method is a *command* (i.e. it does not return a result) then it is executed asynchronously, leading to concurrency; if it is a *query* (i.e. a result is returned and must be waited for) then it is executed synchronously. Note that processes cannot synchronize via shared memory: only by exchanging requests.

The possibility for objects to have different handlers is captured in the type system by the keyword **separate**. To request method calls on objects of **separate** type, programmers simply make the calls within *separate blocks*: these are the bodies of any methods that have separate objects as formal parameters. SCOOP provides guarantees about the order in which calls in these blocks are executed, so as to help programmers avoid concurrency errors. In particular, method calls on separate objects will be logged as requests by their handlers in the order that they are given in the program text; furthermore, there will be no intervening requests logged from other handlers. These guarantees exclude data races by construction, and allow programmers to apply sequential reasoning within separate blocks independently of the rest of the program.

Consider the concurrent version of transfer in Listing 1, in which bank account objects have concurrently running handlers. Suppose that a process calls the method `transfer (acc1, acc2, 100)` on **separate** accounts `acc1` and `acc2`. The body of the method contains two commands on these **separate** objects—thus, two asynchronously executed requests—that transfer the stated amount from the first account to the second. It also contains balance queries which are executed synchronously. The SCOOP guarantees ensure that while the process is inside the body of `transfer`, no other process can log intervening requests on `acc1` or `acc2`. As a result, it would not be possible for another process to observe the balances of the two accounts in an intermediate state, i.e. when the money has been withdrawn from the former but not credited to the latter. The body of `transfer` can thus be reasoned about sequentially and independently of the rest of the program. This additional control over the order in which requests are logged (i.e. that requests cannot be interrupted) is the key distinction SCOOP has over other message-passing-based models such as the actor model, or active objects.

SCOOP provides some more advanced concurrency mechanisms beyond the focus of this paper. Most notable are its generalization of method preconditions to support condition synchronization on **separate** objects, and its support for efficient data sharing between processes sharing memory via “passive” data objects that can be accessed

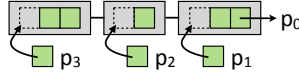


Figure 1: Three processes (p_1, p_2, p_3) logging requests on another (p_0)

directly (i.e. without the overhead of message-passing). We refer to [20] and [18] respectively for more detailed discussions of these concepts.

SCOOP Runtime. The concurrent programming abstractions presented rely on the existence of a runtime that can correctly and efficiently realize them. At the core of SCOOP’s runtime is a simple execution model for managing requests that are sent between processes. Each process is associated with a “queue of queues” [25], that is, a FIFO queue itself containing (possibly several) FIFO subqueues for storing incoming requests. Each of these subqueues represents a “private area” for some other process to log requests, in program text order, and without interference from other processes (since they have their own subqueues). Figure 1 visualizes three processes (p_1, p_2, p_3) simultaneously logging requests (green blocks) on another process (p_0). The process p_0 is handling the subqueues one-by-one in the order that they were created, and handles the requests within them in the order that they were logged there, hence ensuring the SCOOP reasoning guarantees.

Consider again the process that calls `transfer (acc1, acc2, 100)` on two separate accounts, `acc1` and `acc2`. Under the current runtime, the handlers of `acc1` and `acc2` both generate a private subqueue on which the calling process can log requests (i.e. the balance queries and `set_balance` commands) without interruption for the duration of the block. Should another process also need to log requests on an account, then a new private subqueue is generated for it and its requests can be logged without waiting.

We remark that earlier versions of the SCOOP runtime additionally provided timing guarantees by not allowing processes to enqueue requests concurrently [17]. A formal comparison with the current semantics is given in [5].

3 Overview of Distributed SCOOP

In this section we present D-SCOOP (for Distributed SCOOP), which combines network objects and the SCOOP synchronization semantics into a single, distributed programming model that maintains the simplicity of the original abstractions. We present an overview of its architecture and communication protocol, and explain how separate calls are generalized to potentially remote objects (Sections 4 and 5 describe in more detail how control of remote objects is achieved in D-SCOOP, and how the system compensates for unresponsive clients).

A prototype implementation of the D-SCOOP model is available online [7]. Our prototype builds upon the SCOOP support for Eiffel in EiffelStudio [8], which implements the model using threads and shared memory. D-SCOOP generalizes the implementation, allowing for multiple instances of potentially remote SCOOP programs to communicate, under-the-hood, by asynchronous message passing.

Architecture. In D-SCOOP, an instance of a SCOOP program is called a *node*. A node can open a connection to another node through a network socket, which is then shared

by all of its processes. A node can request the *index object* of another node, which is a user-defined object that typically provides the API of the node, or some form of registry. It is valid for a node to not supply an index object, typically if it is a client in a client-server style setup. To be able to accept incoming connections from other nodes, a node must start a *server* and provide its own index object (or a factory that generates them). Every node in a D-SCOOP network has a unique *identifier* (ID), which is independent of any other IDs such as IP addresses. Object references in D-SCOOP include this node ID, along with their object and process identifiers (as in classical SCOOP), with the latter important for determining the number of processes involved in a separate block.

The nodes in D-SCOOP networks communicate, via their connections, using an asynchronous message-passing scheme. Messages conform to a protocol and can be one of two types: a *request*⁶ or a *reply*. Requests are sent from a *client* node to a *supplier*, defining work for the supplier to do. Replies are sent back from the supplier to the client indicating the outcome.

Within nodes, we rely on existing mechanisms of SCOOP for garbage collecting local objects and processes. D-SCOOP however must also account for objects used by multiple nodes. To achieve this, we use a distributed garbage collection algorithm similar to that of Birrell et al. [3].

Requests and Replies. Messages in the D-SCOOP communication protocol have *subjects* which convey their intended semantics. Messages that are requests can have one of many different subjects which we outline in the following. Replies however only indicate success (`OK`) or failure (`FAIL`), sometimes with additional arguments, such as the result of a query call.

The simplest request subjects are `HELLO`, `PING` and `INDEX`, which respectively initialize a connection between nodes, test whether an existing one is still alive, and request the index object of the supplier node (which typically provides an API of methods for retrieving more objects).

A number of requests are required to realize a separate block involving remote objects. A `PRELOCK` request announces that a process in a client node wishes to log calls on one or more processes in a supplier node. When a supplier is ready, the client can issue a `LOCK` request to announce it is now entering the separate block. Following this, it can issue requests corresponding to asynchronous method calls (`CALL`), synchronous calls (`SCALL`), and queries (`QCALL`). To announce leaving the separate block, the client sends an `UNLOCK` request. (We describe in more detail how these requests establish control in Section 4.)

Requests with the subjects `SHARE` and `RELEASE` are respectively used for obtaining and revoking permission for given object references to be shared with third party nodes. They are used by D-SCOOP for garbage collecting.

Finally, `AWAIT` and `READY` requests are used to implement condition synchronization on remote objects. In short: if the condition does not hold, the client process issues an `AWAIT` request before going to sleep. This instructs the supplier to wake it up with a `READY` request once the state of the remote objects changes, so that the condition can be checked again.

⁶ Note that these are distinct from the requests used for inter-process communication in SCOOP.

Message Handling. Incoming messages are handled by the request handlers of D-SCOOP nodes in multiple stages, depending on their subjects. If an incoming message has the subject `HELLO`, `PING`, `SHARE`, or `RELEASE`, then it is handled directly. If a message is a reply, then it is relayed to the appropriate process within the node. Messages addressed to other nodes are relayed.

For messages concerning separate blocks and condition synchronization, a more careful treatment is required. In D-SCOOP, every node has a special designated *proxy process* for handling incoming lock and call requests. Associated with these proxy processes are *proxy objects*, which are surrogates (or placeholders) for actual remote objects, holding references to them. This additional layer is used to catch special contexts in which calls are treated differently. For lack of space we do not go into detail, but mention two of the most important: callbacks (see [20]), and a SCOOP extension for passive data objects (see [18]).

To minimize the overhead of proxy processes and objects, they are created only when needed and removed when they are not. For example, if not existing already, receiving a `LOCK` request with some given object identifiers will trigger the creation of a proxy process on that node and proxies for those objects. And when no longer in use by local processes, they can be collected by the local SCOOP garbage collector.

Remote Calls in Separate Blocks. The communication protocol presented is ultimately the glue that allows for network objects to be used within the SCOOP framework. Our aim was to make the fusion of these concepts as seamless as possible: programmers should not need to be aware of the communication protocol for network objects, and the core abstractions of SCOOP should not need to be fundamentally reinvented to accommodate the extension.

In D-SCOOP we were able to maintain the original abstractions provided by separate blocks, while also providing a natural generalization to support objects residing on other nodes. When a process needs to make a call on a **separate** object, there are now three possible cases to distinguish. If the target object shares the same process (and thus, obviously, the same node), the call is executed immediately—as in SCOOP. If the target object has a distinct process but on the same node, the process logs a request in a private subqueue for the caller (see Section 2)—as in SCOOP. If the target object has a distinct process on a remote node, however, the D-SCOOP communication protocol comes into play, and a `CALL` message is sent to the remote node.

4 Controlling Remote Objects

We have presented an overview of the D-SCOOP architecture, its messaging protocol, and its generalization of **separate** blocks to support calls on remote objects. In this section, we describe how control of remote objects and thus distributed separate blocks are achieved.

In D-SCOOP, separate blocks are handled in three phases: (i) the *prelock phase*, for ensuring a correct ordering; (ii) the *issuing phase*, for enqueueing calls; and (iii) the *execution phase*, for executing calls. The issuing phase happens strictly after the prelock phase. While the execution phase cannot start before the issuing phase, the two can otherwise overlap due to asynchronicity.

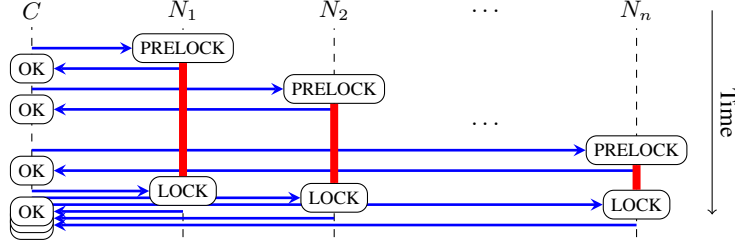


Figure 2: Prelock phase: a process on node C is entering a separate block involving **separate** objects on remote nodes $N_1, \dots, N_n$

Prelock Phase. In standard SCOOP, if a process enters a separate block, the processes handling the **separate** objects generate private subqueues for logging calls (see Section 2 and Figure 1). In D-SCOOP however, if a process enters a separate block involving **separate** objects on remote nodes, messages must be sent to trigger the generation of subqueues in a way that preserves the usual reasoning guarantees. We refer to this messaging phase as the *prelock phase*.

A client node seeking to enter a separate block involving remote objects must first announce its intention by sending `PRELOCK` requests to the nodes they reside on. This is done in a fixed order (a global order based on node IDs) to avoid deadlocks, and one-at-a-time; an `OK` reply must be received before the next `PRELOCK` is sent. Once the last such request is successful, the client node announces that it is entering the separate block and will start issuing calls. This announcement is made via `LOCK` requests, which can be sent asynchronously in any order. By replying with `OK`, the supplier nodes are acknowledging that the involved processes have created private subqueues and are ready to enqueue calls from the client. Figure 2 exemplifies this phase for a client node C that wishes to enter a separate block involving remote objects on supplier nodes $N_1, \dots, N_n$. Here, an arrow denotes the transmission of a message, with its subject given at the end (additional parameters are not visualized).

When multiple nodes are entering prelock phases involving common supplier nodes, blocking must occur in order to maintain the separate block order guarantees. In particular, if a `PRELOCK` message is sent but the supplier is already involved in the prelock phase of a competing node, then the system blocks on that message. Instead of blocking for the whole of the competing node’s separate block, D-SCOOP permits a more fine-grained and efficient solution. In particular, it only blocks until the competing node leaves its prelock phase and starts issuing calls. That is to say, D-SCOOP only blocks while “setting up” the subqueues in a correct order; competing issuing phases can otherwise safely run concurrently.

Issuing and Execution Phases. The prelock phase ends and the issuing phase begins when the final `LOCK` request is successful. At this point, the processes handling all the involved remote objects are ready to enqueue calls. In most circumstances, commands on remote objects are requested via asynchronous `CALL` messages, and queries are requested via synchronous `QCALL` messages. The supplier nodes enqueue commands and immediately reply with an `OK`. When a query is received however, the supplier node en-

```

withdraw (s: separate ACCOUNT; am: NATURAL)
do
  if s.balance >= am then
    s.set_balance (s.balance - am)
  else -- Notify user
  end
end
end

```

Listing 2: Bank account withdrawal method in D-SCOOP

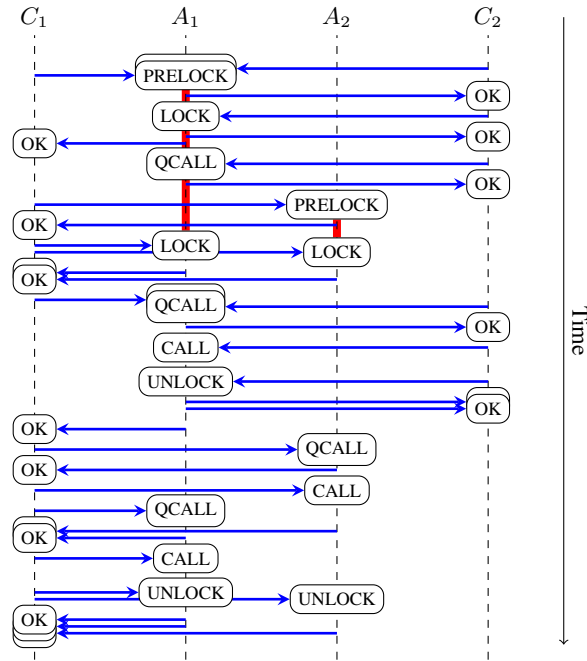


Figure 3: All three phases: a process on C_1 calls transfer on A_1 and A_2 ; a process on C_2 concurrently calls withdraw on A_1

queues it, but only replies once it has been executed (passing the result in an additional parameter of the **OK** message).

The execution phase begins with the execution of the first logged call. If all the calls are asynchronous, it can take place strictly after the issuing phase. The issuing phase ends on sending the **UNLOCK** message; the execution phase ends on processing it.

Example Communication. We return to our running bank account example, which we extend with a simple method `withdraw` (Listing 2) for withdrawing a given amount from a given account that we assume to be remote. The method first synchronously queries the remote object to check that the balance is sufficient, before asynchronously decreasing the balance.

Suppose we have a running D-SCOOP system with two bank accounts on different nodes (A_1, A_2). Suppose now that a client node (C_1) is trying to transfer an amount

from A_1 to A_2 , while another client node (C_2) is trying to withdraw an amount from A_1 . Recall that the bodies of both methods are separate blocks (involving, respectively, **separate** accounts on A_1 , A_2 and A_1). Figure 3 visualizes the messages exchanged in one possible behavior.

Observe that both clients initially send a `PRELOCK` request to A_1 . The request from C_2 is received first and is therefore answered immediately; meanwhile, C_1 blocks. Since C_2 only seeks control over a process on A_1 , it proceeds to send a `LOCK` request, thus completing its prelock phase and generating its private subqueue on A_1 . This allows C_1 to unblock and its first `PRELOCK` request finally succeeds.

Since the prelock phase of one client can take place in parallel to the issuing and execution phases of another, C_2 already starts issuing calls before C_1 concludes its prelock phase. In particular, it requests the balance query (via `QCALL`) which is executed synchronously (and the balance amount returned). Following this, C_1 requests a `PRELOCK` on A_2 (which is uncontended), before completing its prelock phase by sending `LOCK` requests to A_1 and A_2 .

At this point, both C_1 and C_2 issue balance queries (`QCALL`)—the former is evaluating its conditional guard, and the latter is evaluating the expression in the input of `s.set_balance (s.balance - am)`. Since C_2 completed its prelock first, its private subqueue on A_1 is ahead of the subqueue for C_1 , and so its call is executed first. Following this, C_2 requests an asynchronous command (`CALL`) to update the balance, and then exits its separate block via an `UNLOCK` request. Once acknowledged, C_2 knows that the whole transaction (balance and then `set_balance`) was successful, and its effects become visible to other clients. Once the `OK` corresponding to its earlier `QCALL` arrives, C_1 can resume issuing the remaining calls in its separate block before exiting via `UNLOCK` requests to A_1 and A_2 .

Note that the reasoning guarantees of the separate blocks have been maintained. The calls are executed in program text order and without intervening calls from other nodes: within a separate block, multiple balance calls in sequence thus always return the same result. The combination of the prelock phase and the underlying queue of queues semantics prevents the possibility of interleavings that break this.

5 Compensating for Failure

Our presentation of D-SCOOP has thus far focused on the challenge and intricacies of combining the network objects abstraction with a concurrency model and runtime. In this section, we turn our attention to a topic that cannot be ignored in the setting of distributed computing: coping with failure.

While failure can often be managed simply—a fixed timeout is used, for example, to manage it in prelock phases—failure in the middle of a separate block, when only some of the side-effecting commands have been issued, needs a more elaborate solution. We introduce *compensation*, D-SCOOP’s mechanism for reacting to such failure, and demonstrate its use on our running example.

Compensation. In D-SCOOP, upon failure of a supplier, the client is informed using exceptions, and can react to it appropriately in a **rescue**-clause. However, the suppliers in separate blocks are in general oblivious to the status of the client. Our solution is

(a) Client-defined compensation	(b) Supplier-defined compensation
<pre> ... t.compensate (agent t.set_balance (t.balance)) t.set_balance (t.balance + am) ... </pre>	<pre> set_balance (nb: NATURAL) do compensate (agent set_balance (balance)) balance := nb end </pre>

Listing 3: A `set_balance` method together with possible compensation

to introduce *compensation*, a supplier-side mechanism for reacting to client nodes that become unresponsive or disconnect prematurely. The technique registers user-provided closures on suppliers that, before releasing objects controlled by disconnected clients, are executed to restore consistency.

The basic technique is adapted from well-established usage in transactions, in particular, for recovering from long-running transactions or transactions with side effects. It fits naturally with the D-SCOOP model, given that separate blocks are transaction-like in the sense that other clients cannot observe the `separate` objects in intermediate states. One can think of a `LOCK` and `UNLOCK` pair as being the beginning and end of a transaction; after `UNLOCK` is acknowledged, all changes become visible.

The scope of compensation is the issuing phase, and encompasses all executed calls on processes that have been acquired during the prelock phase (and only those processes). In the case of nested separate blocks, the outer block has to take into account that the effects of the inner block are already visible if an `UNLOCK` was issued. This is different to most definitions of nested transactions, in which the inner transaction always finishes together with the outer transaction.

Defining Compensation. Compensation closures are provided by the user as the input of special methods for registering compensation. (We remark that closures are given with the Eiffel keyword `agent`, and can refer to existing methods.) It is possible to define them in the client or the supplier. A client-defined compensation closure is registered before the call to the method to be compensated (and is ignored by the supplier if no request follows). A supplier-defined compensation closure is provided within the called method. The latter comes with the advantage that compensation is defined together with the method, but the former allows for more flexibility: different compensations can be defined depending on where the call is made, which is particularly useful for methods that do not always need compensation.

Consider the simple method `set_balance` for bank account objects (Listing 3) which sets the balance of an account to some provided input. The listing also includes examples of how to make it compensable. On the left is a snippet of the body of `transfer`, now annotated with client-defined compensation before the call. On the right is supplier-defined compensation, provided at the beginning of the method body. In both cases, the `balance` argument to the closure (`agent`) is evaluated to the original balance, so it will restore the old balance if called.

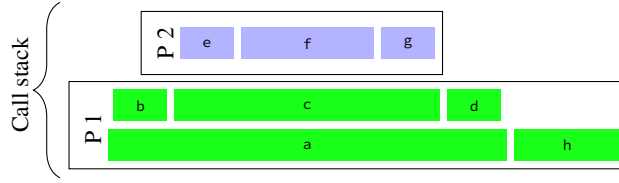


Figure 4: Example call stack

Implementing Compensation. Upon receiving a `LOCK` request, a supplier node stores the IDs of the newly requested processes in a stack. This stack is mainly used to identify which processes need to be released upon `UNLOCK`. Each of the process entries also contains a reference to a set of compensation closures, extracted from the program text. These closures are accompanied by relative timestamps, so that within all the sets for this client each number is unique and a later registration has a strictly higher number than an earlier one. Whenever a process is unlocked normally (i.e. not due to premature disconnection) the respective set is cleared. However, if a client node disconnects prematurely, all sets associated with the client are merged and then ordered by the timestamps. The execution of the compensation closures is done in reverse order.

Figure 4 shows the call stack caused by a remote client calling the method `a` and then `h`. The targets of `a`, `b`, `c`, `d` and `h` are owned by process P_1 , while the targets of the calls `e`, `f`, and `g` are owned by process P_2 . During the execution of `c`, P_1 acquires control over P_2 to execute. After `a` is finished, the client sends another request to execute `h` before releasing P_1 .

We now take a look at three failure scenarios, all of them due to a premature disconnect by the client. If the client disconnects before `a` is executed, nothing happens. The client’s control over P_1 is simply lifted. The second case is more complex: if the client disconnects while `a` is executing, the calls `a`, `b`, $\dots$ `g` are all executed as requested. Since P_1 is issuing the `UNLOCK` request to P_2 before finishing itself, the changes done by `e`, `f`, `g` are visible. The disconnect then causes the compensation closures of `d`, `c`, `b`, `a` to be executed before control over P_1 is released. Consequently, the compensation of `c` has to deal with the fact that the changes due to `e`, `f`, `g` are already visible.

If the client issued the call to `h` but got lost before sending the `UNLOCK` request, the situation is similar, with the one difference being that the compensation of `h` is executed before the others.

6 Evaluation

We evaluated D-SCOOP against Java RMI to gauge its performance against a well-established and widely used approach based on network objects. We sought to collect evidence towards answering two questions. First, is there a performance overhead associated with the automatic synchronization in D-SCOOP, and does it become incommensurate with the effort to manually write synchronization code? Second, do the language abstractions of D-SCOOP facilitate simpler code?

Example Selection. D-SCOOP and Java RMI have many differences: not only in the model, but also in terms of the underlying programming languages (Eiffel and Java)

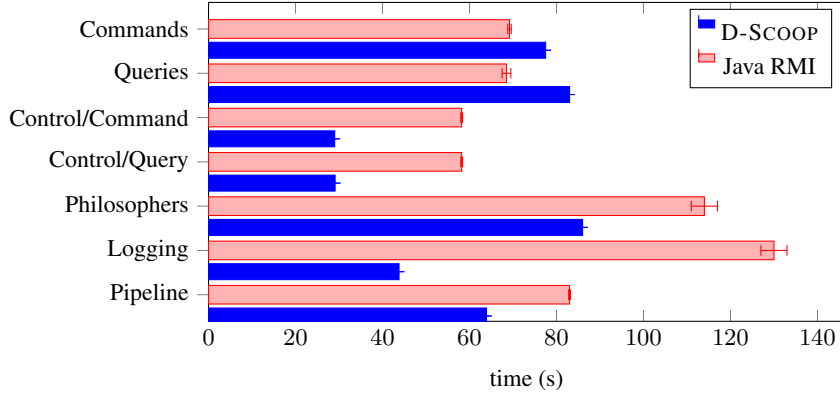


Figure 5: Benchmark results: each run involved several thousand iterations (see [7])

which have many points of variation regarding performance and compilers. In this context, we devised a set of four microbenchmarks isolated to comparing the performance of calls: (i) *command call*, in which a single client sends a series of command calls to the supplier; (ii) *query call*, analogous, but with query calls; (iii) *locking and command call*, in which a few clients compete to control a supplier object and send a single command call; and (iv) *locking and query call*, analogous, but with a single query call.

In addition to microbenchmarks, we also evaluated D-SCOOP against Java RMI on three larger examples. First, *dining philosophers*, a classical example where multiple objects (forks) are repeatedly controlled. For this benchmark, all philosophers and forks reside on different nodes, and we assume that eating, using the fork, and thinking take no time. Second, a more practical example: a *log server*, in which various events are logged. Here, there are multiple log servers for redundancy, meaning that copies of logs can still be retrieved if one fails. To ensure a consistent ordering across servers, a client must control all of them before adding the entries. In our benchmark, three clients repeatedly generate a simple log message, gain control across the servers, and then place it. Third, a *pipeline* representing distributed services. Each stage waits until the previous stages are ready before retrieving data and processing it. Each stage provides one operation of the well known formula $\sqrt{a^2 + b^2}$. We measured the time the final stage needed for a specific number of calculations.

For Java RMI, explicit locking was used to establish a comparable flexibility in the clients. Furthermore, the Java code explicitly orders the locks so as to avoid deadlocks. The source code of the examples and of D-SCOOP itself can be found on our supplementary material webpage [7].

Performance. Overall, we found that despite the potential overhead of automatic synchronization, D-SCOOP’s performance is competitive with—and can be superior to—explicit locking-based synchronization in Java RMI. The results of the performance evaluation are listed in Figure 5 and are the averages of 30 runs; we used two off-the-shelf laptops connected by an ethernet cable. The microbenchmarks show that the performance of both D-SCOOP and Java RMI is similar when just issuing commands or queries. D-SCOOP commands are a bit quicker than D-SCOOP queries due to them be-

Table 1: Code complexity						
	Classes		Features		Instructions	
	RMI	D-SCOOP	RMI	D-SCOOP	RMI	D-SCOOP
Microbenchmarks	3	2	8	6	19	13
Dining philosophers	3	2	6	3	18	10
Logging	6	3	16	9	23	10
Pipelines	2	1	10	16	62	42

ing asynchronous, whereas in RMI both are synchronous. When it comes to the control microbenchmarks, the built-in synchronization in D-SCOOP allows for a more significant improvement in speed, both for synchronous and asynchronous calls. However, the synchronization overhead prevents the asynchronous advantage of Control/Command translating into faster performance than Control/Query.

For both the dining philosophers and the logging example, the fact that the prelock phase can be done in parallel with the issuing and execution phase of another client proves to be a significant advantage in comparison to RMI. In addition, the logging example shows the advantage of asynchronous calls in D-SCOOP. The underlying semantics make it possible to ensure control over multiple nodes and have multiple clients issuing asynchronous calls at the same time. The pipeline example has less congestion around the protected objects; here, the advantage of D-SCOOP lies solely in slightly fewer messages sent due to more powerful synchronization mechanics.

Simplicity. Our second question asked whether the language abstractions also yield simpler code. For our seven benchmarks, we recorded: (i) the number of classes involved, excluding primitive types, classes, and strings, and ignoring the RMI remote interface; (ii) the number of features (i.e. attributes and methods), ignoring the Java “getters” in RMI since they just return an otherwise counted attribute; and (iii) the number of written instructions, excluding boilerplate code. This ensures that the differences are only due to synchronization. Table 1 lists the results.

As can be seen, the solutions in D-SCOOP are much more compact across the three measurements. In the case of advanced techniques such as condition synchronization—an in-depth discussion is omitted for brevity—the complexity of RMI increases further still. Note that not included in the RMI examples are compensation and the automatic releasing of locks, since they are difficult to achieve in that framework. Also, although the usage of a lock or semaphore is counted as a class, its features are not counted in the feature column since they are already provided by the library. We remark that these numbers only indicate that D-SCOOP programs are more compact than their RMI counterparts. What we leave to future work is a study of users themselves to determine whether the D-SCOOP abstractions are easier to read and program with, regardless of their compactness. (An existing SCOOP study is encouraging [19].)

7 Related Work

There is a wide selection of work addressing concurrency and distribution in the object-oriented paradigm. Here, we highlight some work that is closest to our own.

The active object [14] design pattern (which inherits from the actor model [1]), like SCOOP, decouples method calls from method executions. Such objects are associated with their own processes, which can send messages to each other asynchronously, introducing concurrency. Despite the similarity to SCOOP, active objects lack the guarantee of interference-freedom when multiple objects are involved. Furthermore, non-active objects have to be protected manually, and there is no built-in support for condition synchronization (although it is possible to use the observer pattern to actively notify waiting processes). SCOOP can be seen as an advanced form of active objects: objects are by default active, but multiple objects can share the same process. In addition, the SCOOP synchronization mechanisms ensure the absence of intervening calls and also protect non-active objects [18]. Condition synchronization is simple (via method pre-conditions) and does not require signaling.

There have been some successful attempts to generalize ideas from active objects and the actor model to distributed programming frameworks, with some prominent examples including Creol [12], AmbientTalk [6], and JCoBox [21]. The latter partitions the object space into “coboxes”, each with a common thread of control to improve safety; an approach similar to the processes of SCOOP and D-SCOOP. Caromel et al. [4] consider a way of unifying threads and objects to support simpler reasoning about distributed computing, and provide a formal calculus. An important distinction of D-SCOOP in comparison to other frameworks is the impossibility of interrupting requests sent to multiple (potentially distributed) objects controlled by different threads, giving the model its transaction-like semantics.

Network objects [2] share some similarities with active objects, although calls to them are traditionally synchronous to mimic standard method calls, and calls to local network objects are usually handled by the calling process. Creol exemplifies different synchronization approaches possible with active objects, and their natural extension to network objects. Some languages, such as E [16], avoid blocking entirely to ensure deadlock-freedom. This, in our view, can lead to complex behavior that is difficult to understand from the point of view of classical sequential programming. By making synchronization simpler to use, D-SCOOP potentially reduces (but does not eliminate) the risk of deadlocks.

For dealing with failures, the programming language Argus [15] supports “atomic objects” that can be used in a transaction. In contrast, our compensation approach is not limited to pure data-objects.

8 Conclusion

This paper made a case for combining network objects with synchronization models. We presented D-SCOOP, a distributed programming model obtained by combining the network objects abstraction with the runtime semantics of the object-oriented concurrency model SCOOP. We presented an efficient two-phase locking protocol that generalized the strong reasoning guarantees of SCOOP to network objects, allowing for interference-free and transaction-like reasoning on (potentially multiple) remotely located objects, without the programmer having to explicitly manage their synchronization. Furthermore, we proposed a compensation mechanism by which D-SCOOP pro-

grams can recover from failure. The evaluation of our prototype implementation [7] suggested that D-SCOOP remains competitive against—and can outperform—explicit locking-based synchronization in Java RMI, a well-established realization of network objects, with the automatic synchronization mechanisms also allowing for more compact code.

In future work, we plan to improve the efficiency of D-SCOOP with respect to intra-object parallelism [11, 13]. We will investigate concepts such as slicing [22], and the possible integration of software transactional memory [9]. We will also investigate whether performance can be improved, by (safely) relaxing the requirement that one node communicates with another via a single connection. Finally, we want to formalize the D-SCOOP semantics using [5] to test extensions, and provide a formal proof that the protocol and algorithms correctly generalize the SCOOP guarantees.

Acknowledgements. We thank Sebastian Nanz for his invaluable support throughout this project. We also thank Carlo A. Furia and the anonymous referees for their helpful comments and criticisms. The underlying research was partially funded by ERC Grant CME #291389.

References

1. Agha, G.: ACTORS: A Model of Concurrent Computation in Distributed Systems. MIT Press (1986)
2. Birrell, A., Nelson, G., Owicki, S.S., Wobber, E.: Network objects. In: Proc. SOSP 1993. pp. 217–230. ACM (1993)
3. Birrell, A., et al.: Distributed garbage collection for network objects. Tech. rep., Systems Research Center (1993)
4. Caromel, D., Henrio, L., Serpette, B.P.: Asynchronous sequential processes. *Information and Computation* 207(4), 459–495 (2009)
5. Corrodi, C., Heußner, A., Poskitt, C.M.: A graph-based semantics workbench for concurrent asynchronous programs. In: Proc. FASE 2016. LNCS, vol. 9633, pp. 31–48. Springer (2016)
6. Dedecker, J., Cutsem, T.V., Mostinckx, S., D’Hondt, T., Meuter, W.D.: Ambient-oriented programming in AmbientTalk. In: Proc. ECOOP 2006. LNCS, vol. 4067, pp. 230–254. Springer (2006)
7. Distributed SCOOP website. <http://cme.ethz.ch/scoop/dscoop/>
8. Eiffel Documentation: Concurrent Eiffel with SCOOP. <https://www.eiffel.org/doc/solutions/Concurrent%20programming%20with%20SCOOP>, acc.: Apr. 2016
9. Eugster, P., Vaucouleur, S.: Composing atomic features. *Science of Computer Programming* 63(2), 130–146 (2006)
10. Grand Central Dispatch (GCD) Reference. https://developer.apple.com/library/mac/documentation/Performance/Reference/GCD_libdispatch_Ref/index.html, acc.: Apr. 2016
11. Henrio, L., Huet, F., István, Z.: Multi-threaded active objects. In: Proc. COORDINATION 2013. LNCS, vol. 7890, pp. 90–104. Springer (2013)
12. Johnsen, E.B., Owe, O., Yu, I.C.: Creol: A type-safe object-oriented model for distributed concurrent systems. *Theoretical Computer Science* 365(1-2), 23–66 (2006)
13. Johnsen, E.B., Blanchette, J.C., Kyas, M., Owe, O.: Intra-object versus inter-object: Concurrency and reasoning in Creol. In: Proc. TTSS 2008. ENTCS, vol. 243, pp. 89–103 (2009)

14. Lavender, R.G., Schmidt, D.C.: Active object: An object behavioral pattern for concurrent programming. In: Vlissides, J.M., Coplien, J.O., Kerth, N.L. (eds.) *Pattern Languages of Program Design 2*, pp. 483–499. Addison-Wesley (1996)
15. Liskov, B.: Distributed programming in Argus. *Communications of the ACM (CACM)* 31(3), 300–312 (1988)
16. Miller, M.S., Tribble, E.D., Shapiro, J., Laboratories, H.P.: Concurrency among strangers: Programming in E as plan coordination. In: *Proc. TGC 2005. LNCS*, vol. 3705, pp. 195–229. Springer (2005)
17. Morandi, B., Schill, M., Nanz, S., Meyer, B.: Prototyping a concurrency model. In: *Proc. ACSD 2013*. pp. 170–179. IEEE (2013)
18. Morandi, B., Nanz, S., Meyer, B.: Safe and efficient data sharing for message-passing concurrency. In: *Proc. COORDINATION 2014. LNCS*, vol. 8459, pp. 99–114. Springer (2014)
19. Nanz, S., Torshizi, F., Pedroni, M., Meyer, B.: Design of an empirical study for comparing the usability of concurrent programming languages. In: *Proc. ESEM 2011*. pp. 325–334. IEEE Computer Society (2011)
20. Nienaltowski, P.: Practical framework for contract-based concurrent object-oriented programming. Doctoral dissertation, ETH Zürich (2007)
21. Schäfer, J., Poetzsch-Heffter, A.: JCoBox: Generalizing active objects to concurrent components. In: *Proc. ECOOP 2010. LNCS*, vol. 6183, pp. 275–299. Springer (2010)
22. Schill, M., Nanz, S., Meyer, B.: Handling parallelism in a concurrency model. In: *Proc. MUSEPAT 2013. LNCS*, vol. 8063, pp. 37–48. Springer (2013)
23. Shavit, N., Touitou, D.: Software transactional memory. *Distributed Computing* 10(2), 99–116 (1997)
24. Torshizi, F.A., Ostroff, J.S., Paige, R.F., Chechik, M.: The SCOOP concurrency model in Java-like languages. In: *Proc. CPA 2009. Concurrent Systems Engineering Series*, vol. 67, pp. 7–27. IOS Press (2009)
25. West, S., Nanz, S., Meyer, B.: Efficient and reasonable object-oriented concurrency. In: *Proc. ESEC/FSE 2015*. pp. 734–744. ACM (2015)