

Malacology: A Programmable Storage System

Michael A. Sevilla[†], Noah Watkins[†], Ivo Jimenez,
Peter Alvaro, Shel Finkelstein, Jeff LeFevre, Carlos Maltzahn

University of California, Santa Cruz

{msevilla, jayhawk, ivo}@soe.ucsc.edu, {palvaro, shel}@ucsc.edu, {jlefevre, carlosm}@soe.ucsc.edu

Abstract

Storage systems need to support high-performance for special-purpose data processing applications that run on an evolving storage device technology landscape. This puts tremendous pressure on storage systems to support rapid change both in terms of their interfaces and their performance. But adapting storage systems can be difficult because unprincipled changes might jeopardize years of code-hardening and performance optimization efforts that were necessary for users to entrust their data to the storage system. We introduce the programmable storage approach, which exposes internal services and abstractions of the storage stack as building blocks for higher-level services. We also build a prototype to explore how existing abstractions of common storage system services can be leveraged to adapt to the needs of new data processing systems and the increasing variety of storage devices. We illustrate the advantages and challenges of this approach by composing existing internal abstractions into two new higher-level services: a file system metadata load balancer and a high-performance distributed shared-log. The evaluation demonstrates that our services inherit desirable qualities of the back-end storage system, including the ability to balance load, efficiently propagate service metadata, recover from failure, and navigate trade-offs between latency and throughput using leases.

CCS Concepts • **Information systems** → **Distributed storage**; • **Software and its engineering** → **File systems management**; *Software functional properties*

Keywords Distributed Storage, Programmability, Ceph

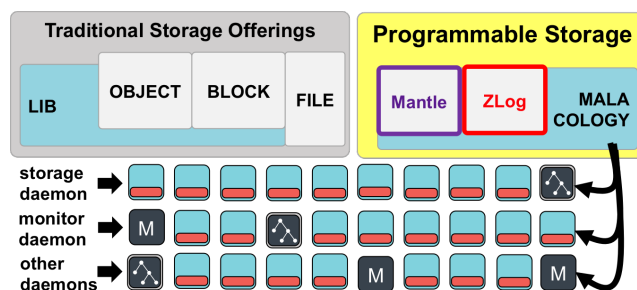


Figure 1: Scalable storage systems have storage daemons which store data, monitor daemons (M) that maintain cluster state, and service-specific daemons (e.g., file system metadata servers). Malacology enables the programmability of internal abstractions (bold arrows) to re-use and compose existing subsystems. With Malacology, we built new higher-level services, ZLog and Mantle, that sit alongside traditional user-facing APIs (file, block, object).

1. Introduction

A storage system implements abstractions designed to persistently store data and must exhibit a high level of correctness to prevent data loss. Storage systems have evolved around storage devices that often were orders of magnitude slower than CPU and memory, and therefore could dominate overall performance if not used carefully. Over the last few decades members of the storage systems community have developed clever strategies to meet correctness requirements while somewhat hiding the latency of traditional storage media [12]. To avoid lock-in by a particular vendor, users of storage systems have preferred systems with highly standardized APIs and lowest common denominator abstract data types such as blocks of bytes and byte stream files [4].

A number of recent developments have disrupted traditional storage systems. First, the falling prices of flash storage and the availability of new types of non-volatile memory that are orders of magnitude faster than traditional spinning media are moving overall performance bottlenecks away from storage devices to CPUs and networking, and pressure storage systems to shorten their code paths and incorporate new optimizations [21, 22]. Second, emerging “big data” applications demand interface evolution to support flexible consistency as well as flexible structured data representa-

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

EuroSys '17, April 23–26, 2017, Belgrade, Serbia

© 2017 ACM. ISBN 978-1-4503-4938-3/17/04...\$15.00

DOI: <http://dx.doi.org/10.1145/3064176.3064208>

[†] These authors contributed equally to this work.

tions. [3]. Finally, production-quality scalable storage systems available as open source software have established and are continuing to establish new, *de-facto* API standards at a faster pace than traditional standards bodies [31, 40].

The evolutionary pressure placed on storage systems by these trends raises the question of whether there are principles that storage systems designers can follow to evolve storage systems efficiently, without jeopardizing years of code-hardening and performance optimization efforts. In this paper we investigate an approach that focuses on identifying and exposing existing storage system resources, services, and abstractions that in a generalized form can be used to *program* new services. This ‘dirty-slate’ approach of factoring out useful code lets programmers re-use subsystems of the back-end storage system, thus inheriting their optimizations, established correctness, robustness, and efficiency. ‘Clean-slate’ approaches could be implemented faster but they do so at the expense of “throwing away” proven code.

Contribution 1: We define a programmable storage system to be a storage system that facilitates the re-use and extension of existing storage abstractions provided by the underlying software stack, to enable the creation of new services via composition. A programmable storage system can be realized by exposing existing functionality (such as file system and cluster metadata services and synchronization and monitoring capabilities) as interfaces that can be “glued together” in a variety of ways using a high-level language. Programmable storage differs from *active storage* [35]—the injection and execution of code within a storage system or storage device—in that the former is applicable to any component of the storage system, while the latter focuses at the data access level. Given this contrast, we can say that active storage is an example of how one internal component (the storage layer) is exposed in a programmable storage system.

To illustrate the benefits and challenges of this approach we have designed and evaluated Malacology, a programmable storage system that facilitates the construction of new services by re-purposing existing subsystem abstractions of the storage stack. We build Malacology in Ceph, a popular open source software storage stack. We choose Ceph to demonstrate the concept of programmable storage because it offers a broad spectrum of existing services, including distributed locking and caching services provided by file system metadata servers, durability and object interfaces provided by the back-end object store, and propagation of consistent cluster state provided by the monitoring service (see Figure 1). Malacology is expressive enough to provide the functionality necessary for implementing new services.

Malacology includes a set of interfaces that can be used as building blocks for constructing novel storage abstractions, including:

1. An interface for managing strongly-consistent time-varying **service metadata**.

2. An interface for installing and evolving domain-specific, cluster-wide **data I/O** functionality.
3. An interface for managing access to **shared resources** using a variety of optimization strategies.
4. An interface for **load balancing** resources across the cluster.
5. An interface for **durability** that persists policies using the underlying storage stack’s object store.

Contribution 2: We implement two distributed services using Malacology to demonstrate the feasibility of the programmable storage approach:

1. A high-performance distributed shared log service called ZLog, that is an implementation of CORFU [6]
2. An implementation of Mantle, the programmable load balancing service [37]

The remainder of this paper is structured as follows. First, we describe and motivate the need for programmable storage by describing current practices in the open source software community. Next we describe Malacology by presenting the subsystems within the underlying storage system that we re-purpose, and briefly describe how those system are used within Malacology (Section 4). Then we describe the services that we have constructed in the Malacology framework (Section 5), and evaluate our ideas within our prototype implementation (Section 6). We conclude by discussing future and related work.

2. Application-Specific Storage Stacks

Building storage stacks from the ground up for a specific purpose results in the best performance. For example, GFS [18] and HDFS [38] were designed specifically to serve MapReduce and Hadoop jobs, and use techniques like exposing data locality and relaxing POSIX constraints to achieve application-specific I/O optimizations. Another example is Boxwood [32], which experimented with B-trees and chunk stores as storage abstractions to simplify application building. Alternatively, general-purpose storage stacks are built with the flexibility to serve many applications by providing standardized interfaces and tunable parameters. Unfortunately, managing competing forces in these systems is difficult and users want more control from the general-purpose storage stacks without going as far as building their storage system from the ground up.

To demonstrate a recent trend towards more application-specific storage systems we examine the state of programmability in Ceph. Something of a storage Swiss army knife, Ceph simultaneously supports file, block, and object interfaces on a single cluster [1]. Ceph’s Reliable Autonomous Distributed Object Storage (RADOS) system is a cluster of object storage daemons that provide Ceph with data durability and integrity using replication, erasure-coding, and

scrubbing [50]. Ceph already provides some degree of programmability; the object storage daemons support domain-specific code that can manipulate objects on the server that has the data local. These “interfaces” are implemented by composing existing low-level storage abstractions that execute atomically. They are written in C++ and are statically loaded into the system.

The Ceph community provides empirical evidence that developers are already beginning to embrace programmable storage. Figure 2 shows a dramatic growth in the production use of domain-specific interfaces in the Ceph community since 2010. In that figure, classes are functional groupings of methods on storage objects (e.g. remotely computing and caching the checksum of an object extent). What is most remarkable is that this trend contradicts the notion that API changes are a burden for users. Rather it appears that gaps in existing interfaces are being addressed through ad-hoc approaches to programmability. In fact, Table 1 categorizes existing interfaces and we clearly see a trend towards reusable services.

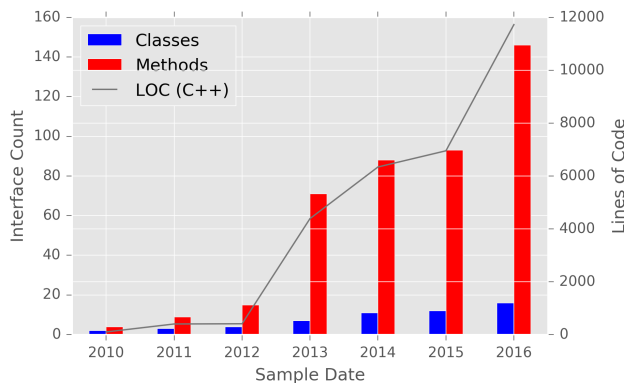


Figure 2: [source] Since 2010, the growth in the number of co-designed object storage interfaces in Ceph has been accelerating. This plot is the number of object classes (a group of interfaces), and the total number of methods (the actual API end-points).

Category	Example	#
Logging	Geographically distribute replicas	11
Metadata Management	Snapshots in the block device OR Scan extents for file system repair	74
Locking	Grants clients exclusive access	6
Other	Garbage collection, reference counting	4

Table 1: A variety of object storage classes exist to expose interfaces to applications. # is the number of methods that implement these categories.

The takeaway from Figure 2 is that programmers are already trying to use programmability because their needs, whether they be related to performance, availability, consistency, convenience, etc., are not satisfied by the existing default set of interfaces. The popularity of the custom object interface facility of Ceph could be due to a number of reasons, such as the default algorithms/tunables of the storage system being insufficient for the application’s performance goals,

programmers wanting to exploit application-specific semantics, and/or programmers knowing how to manage resources to improve performance. A solution based on application-specific object interfaces is a way to work around the traditionally rigid storage APIs because custom object interfaces give programmers the ability to tell the storage system about their application: if the application is CPU or I/O bound, if it has locality, if its size has the potential to overload a single node, etc. Programmers often know what the problem is and how to solve it, but until the ability to modify object interfaces, they had no way to express to the storage system how to handle their data.

Our approach is to expose more of the commonly used, code-hardened subsystems of the underlying storage system as interfaces. The intent is that these interfaces, which can be as simple as a redirection to the persistent data store or as complicated as a strongly consistent directory service, should be used and re-used in many contexts to implement a wide range of services. By making programmability a ‘feature’, rather than a ‘hack’ or ‘workaround’, we help standardize a development process that now is largely ad-hoc.

3. Challenges

Implementing the infrastructure for programmability into existing services and abstractions of distributed storage systems is challenging, even if one assumes that the source code of the storage system and the necessary expertise for understanding it is available. Some challenges include:

- Storage systems are generally required to be highly available so that any complete restarts of the storage system to reprogram them is usually unacceptable.
- Policies and optimizations are usually hard-wired into the services and one has to be careful when factoring them to avoid introducing additional bugs. These policies and optimizations are usually cross-cutting solutions to concerns or trade-offs that cannot be fully explored at the time the code is written (as they relate to workload or hardware). Given these policies and optimizations, decomposition of otherwise orthogonal internal abstractions can be difficult or dangerous.
- Mechanisms that are often only exercised according to hard-wired policies and not in their full generality have hidden bugs that are revealed as soon as those mechanisms are governed by different policies. In our experience introducing programmability into a storage system proved to be a great debugging tool.
- Programmability, especially in live systems, implies changes that need to be carefully managed by the system itself, including versioning and propagation of those changes without affecting correctness.

To address these challenges we present Malacology, our prototype programmable storage system. It uses the programmable storage design approach to evolve storage systems efficiently and without jeopardizing years of code-hardening and performance optimization efforts. Although Malacology uses the internal abstractions of the underlying storage system, including its subsystems, components, and implementations, we emphasize that our system still addresses the general challenges outlined above.

The main challenge of designing a programmable storage system is choosing the right internal abstractions and picking the correct layers for exposing them. A programmable storage system is not defined by what abstractions are exposed, rather a programmable storage system adheres to the design approach of exposing interfaces so administrators can have better control of the storage system. The interfaces presented in this paper are abstractions that we found useful for building our prototype services ZLog and Mantle, yet they may not provide the best trade-offs for all higher-level services. For example, if consensus is correctly exposed one could implement high-level features like versioning, serialization, or various flavors of strongly consistent data management on top; but perhaps a low-level consensus interface is suited well for a particular set of applications. These questions are not answered in this paper and instead we focus on showing the feasibility of building such a system, given advances in the quality and robustness of today’s storage stacks.

The Malacology prototype we present has been implemented on Ceph. While there are other systems on top of which Malacology could be implemented (see Table 2), we choose Ceph because it is a production quality system and because it is open source. The large developer community ensures that code is robust and the visibility of the code lets us expose any interface we want. In the next section we describe the Ceph components that we expose as Malacology interfaces.

4. Malacology: A Programmable Storage System

The guiding principle is to re-use existing services and compose them so that these services can be *programmed*. We accomplish programmability of a service by exporting bindings for an interpreted programming language so that programming can occur without having to restart the storage system (see also below, Section 4.4). Table 2 shows the internal services from Ceph that we expose in the Malacology prototype via Lua [26] bindings and Figure 3 compares what was already present in Ceph (gray boxes) to the Malacology interfaces we added. Section 5 will describe the higher-level services we built with these interfaces.

Lua is a portable embedded scripting language and we choose it as the interpreted language for Malacology because it offers superior performance and productivity trade-offs, including a JIT-based implementation that is well known for

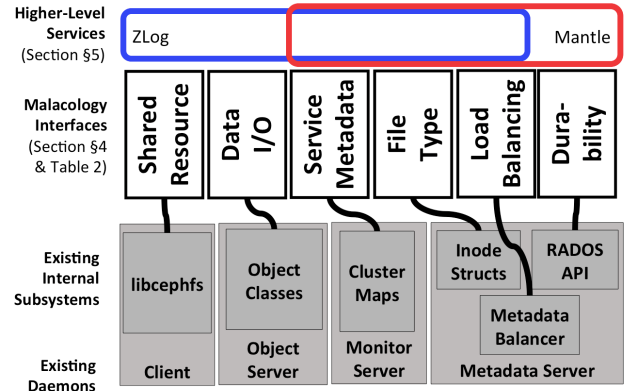


Figure 3: Malacology is implemented on the daemons and clients that run in a Ceph cluster. Interfaces expose internal subsystems and are used as building blocks for higher-level services.

near native performance. Additionally, Lua has been used extensively in game engines, and systems research [45], including storage systems where it has been effectively used both on [17, 20, 48] and off [37] the performance critical path. Finally, the flexibility of the runtime allows execution sandboxing in order to address security and performance concerns. We will now discuss the common subsystems used to manage storage systems and how Malacology makes them programmable.

4.1 Service Metadata Interface

Keeping track of state in a distributed system is an essential part of any successful service and a necessary component in order to diagnose and detect failures, when they occur. This is further complicated by variable propagation delays and heterogeneous hardware in dynamic environments. Service metadata is information about the daemons in the system and includes membership details, hardware layout (e.g., racks, power supplies, etc.), data layout, and daemon state and configuration. It differs from traditional file system metadata which is information about files. For the rest of the paper when we use the phrase “metadata server” or “metadata service”, we are referring to the daemon(s) that manages file system metadata (not service metadata).

Existing Ceph Implementation: a consistent view of cluster state among server daemons and clients is critical to provide strong consistency guarantees to clients. Ceph maintains cluster state information in per-subsystem data structures called “maps” that record membership and status information. A Paxos [30] monitoring service is responsible for integrating state changes into cluster maps, responding to requests from out-of-date clients and synchronizing members of the cluster whenever there is a change in a map so that they all observe the same system state. As a fundamental building block of many system designs, consensus abstractions such as Paxos are a common technique for maintaining consistent data versions, and are a useful system to expose.

Interface	Section	Example in Production Systems	Example in Ceph	Provided Functionality
Service Metadata	§4.1	Zookeeper/Chubby coordination [13, 25]	cluster state management [29]	consensus/consistency
Data I/O	§4.2	Swift in situ storage/compute [34]	object interface classes [46]	transaction/atomicity
Shared Resource	§4.3.1	MPI collective I/O, burst	POSIX metadata protocols	serialization/batching
File Type	§4.3.2	MPI supports architecture-specific code [42]	file striping strategy	data/metadata access
Load Balancing	§4.3.3	VMWare’s VM migration [16, 23]	migrate POSIX metadata [49]	migration/sampling
Durability	§4.4	S3/Swift interfaces (RESTful API)	object store library [50]	persistence/safety

Table 2: Common internal abstractions. The same “internal abstractions” are common across large-scale systems because they provide functionality that solves general distributed systems problems. Here we list examples of what these internal abstractions are used for in “production systems” and in Ceph. Malacology provides these internal abstractions as interfaces (Section 4) that higher level services (Section 5) can use.

The default behavior of the monitor can be seen as a Paxos-based notification system, similar to the one introduced in [13], allowing clients to identify when new values (termed epochs in Ceph) are associated to given maps. Since Ceph does not expose this service directly, as part of our Malacology implementation, we expose a key-value service designed for managing service metadata that is built on top of the consensus engine. Since the monitor is intended to be out of the high-performance I/O path, a general guideline is to make use of this functionality infrequently and to assign small values to maps.

Malacology: we expose a strongly-consistent view of time-varying service metadata as an interface rather than a hidden internal component. This is shown in Figure 4, where object interfaces and load balancer policies use the Service Metadata interface. Malacology provides a generic API for adding arbitrary values to existing subsystem cluster maps. As a consequence of this, applications can define simple but useful service-specific logic to the strongly-consistent interface, such as authorization control (just specific clients can write new values) or triggering actions based on specific values (e.g. sanitize values). The higher-level services we implement in Section 5 make use of this functionality to register, version and propagate dynamic code (Lua scripts) for new object interfaces defined in storage daemons (Section 4.2) and policies in the load balancer (Section 4.3). Using this service guarantees that interface definitions are not only made durable, but are transparently and consistently propagated throughout the cluster so that clients are properly synchronized with the latest interfaces.

Impact: provides core functionality because it lets daemons come to a consensus on system critical state. Bugs in the internal subsystems or omitting this from services that need this type of consistency affects correctness.

4.2 Data I/O Interface

Briefly described in Section 2, Ceph supports application-specific object interfaces [50]. The ability to offload computation can reduce data movement, and transactional interfaces can significantly simplify construction of complex storage interfaces that require uncoordinated parallel access.

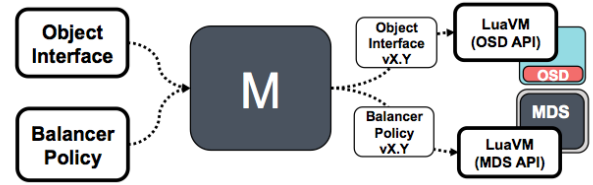


Figure 4: Malacology allows users to dynamically define object/file system metadata interfaces by composing the object storage daemon (OSD) and metadata server (MDS) subsystems with an embedded Lua VM. It uses the Service Metadata interface in the monitor (M) to propagate interfaces/versions across the cluster

Existing Ceph Implementation: an object interface is a plugin structured similarly to that of an RPC in which a developer creates a named function within the cluster that clients may invoke. In the case of Ceph each function is implicitly run within the context of an object specified when the function is called. Developers of object interfaces express behavior by creating a composition of native interfaces or other custom object interfaces, and handle serialization of function input and output. A wide range of native interfaces are available to developers such as reading and writing to a byte stream, controlling object snapshots and clones, and accessing a sorted key-value database. These native interfaces may be transactionally composed along with application specific logic to create semantically rich interfaces. An example would be an interface that atomically updates a matrix stored in the bytestream and an index of the matrix stored in the key-value database.

Malacology: the implementation of Ceph’s object abstraction, although powerful, does not readily support programmability. Supporting only C/C++ for object interface developers, Ceph requires distribution of compiled binaries for the correct architecture, adding a large barrier of entry for developers and system administrators. Second, having no way to dynamically unload modules, any changes require a full restart of a storage daemon which may have serious performance impacts due to loss of cached data. And finally, the security limitations of the framework limit the use of object interfaces to all but those with administrative level access and deep technical expertise.

To address these concerns, Malacology takes advantage of Lua extensions contributed by the Ceph community. This allows new object interfaces to be dynamically loaded into the system and modified at runtime, resulting in a object storage API with economy of expression, which at the same time provides the full set of features of the original object interface implementation. New object interfaces that are expressed in thousands of lines of code can be implemented in approximately an order of magnitude less code [17]. While the use of Lua does not prevent deployment of malicious code, certain types of coding mistakes can be handled gracefully, and access policies are used to limit access to trusted users [26].

Impact: helps applications optimize performance by pushing behavior to lower parts of the storage stack, thereby minimizing hops and distributing computation.

4.3 Distributed Metadata Interfaces

File systems provide clients with the familiar POSIX file abstraction. While this guarantees strong consistency it comes at the cost of scalability, increased complexity, and lower performance. In general, distributed file systems protect resources by providing hierarchical indexing and distributed locking services.

4.3.1 Shared Resource Interface

File system metadata servers manage client sessions, allowing clients to obtain locks (e.g. file byte ranges), and capabilities (e.g. to cache file data). Clients and metadata servers use a cooperative protocol in which clients voluntarily release resources back to the file system metadata service in order to implement sharing policies.

Existing Ceph Implementation: the locking service implements a capability-based system that expresses what data and file system metadata clients are allowed to access as well as what state they may cache and modify locally. While designed for the file abstraction, indexing, locking, and caching are all common services that are useful to a broad spectrum of applications. Distributed applications that share centralized resources (e.g. a database or directory) face similar challenges which are often solved using application-specific sharding.

Malacology: while the current policy for sharing access and voluntarily releasing resources is largely best-effort, Malacology supports generalized policies between metadata servers and clients that can be used to implement fairness or priority.

Impact: provides core functionality to protect and provide exclusive access for any shared resource. May hurt performance if the resource in question does not require strong consistency.

4.3.2 File Type Interface

Applications that manage large amounts of file system metadata (e.g. users or database snapshots) often require a naming service. The metadata service exposes a POSIX file system hierarchy where files and directories are represented as inode data structures.

Existing Ceph Implementation: CephFS is the POSIX compliant file system that uses Ceph. Inodes are quite large (1KB for an inode, 400 bytes for a directory entry, and 700 bytes for a directory) and contain CephFS-specific policies like how to stripe data across RADOS.

Malacology: allows new inode types to be defined such that applications can create domain-specific interfaces to inodes that may modify locking and capability policies. We will show how this is used in Section 5.2.1 when we discuss a distributed shared-log built on Malacology.

Impact: this interface is both a feature and a performance optimization. It is a feature because it allows developers to add support for different storage types, such as how to read new file formats or what consistency semantics to use for a specific subtree in the hierarchical namespace. It is also a performance optimization because future programmers can add optimizations for processing specific types of files into the inode itself.

4.3.3 Load Balancing Interface

Many large scale storage systems separate file system metadata and data I/O so that the corresponding services can scale independently. Metadata requests transfer small amounts of data and they happen relatively frequently so many systems employ separate file system metadata clusters.

Existing Ceph Implementation: addresses the challenge of balancing file system metadata load with a separate metadata cluster. This cluster uses load balancing policies to migrate directory inodes around the cluster to alleviate overloaded servers [49]. The policies use metrics based on system state (e.g. CPU and memory utilization) and statistics collected by the cluster (e.g. the popularity of an inode). Ceph uses dynamic subtree partitioning to move variable sized namespace subtrees. These units can be shipped anywhere (i.e., to any metadata server of any capacity) at any time for any reason. The original balancer was designed with hard-coded policies and tunables.

Malacology: the existing load balancing mechanisms are exposed through an API and programmers can customize the behavior through a domain specific language. These mechanisms include the ability to migrate, partition, and measure load. Using the Service Metadata and Durability interfaces, this Load Balancing interface can safely version bal-

ancer policies, save balancer policies in the back-end object store and centralize warnings/errors. When combined with the File Type interface, the Load Balancing interface can express policies for handling a variety of multi-tenant workloads.

Impact: helps applications optimize performance by allowing them to specify how to partition, replicate, and distribute metadata in response to overloaded servers.

4.4 Durability Interface

Object stores protect data using techniques like erasure coding, replication, and data scrubbing. For scalability, many of these features are implemented using a peer-to-peer protocol that allows object storage daemons to operate autonomously without a centralized coordinator.

Existing Ceph Implementation: provides storage by striping and replicating data across RADOS [50], the reliable distributed object store. RADOS protects data using common techniques such as erasure coding, replication, and scrubbing. For example, when the number of placement groups change, the object storage daemons re-balance and re-shard data in the background in a process called placement group splitting. During placement group splitting, object storage daemons communicate directly with each other to converge on a new data layout. In order to reduce load on the monitoring service, the object storage daemons use a gossip protocol to efficiently propagate changes to cluster maps throughout the system, and autonomously initiate recovery mechanisms when failures are discovered.

Malacology: metadata service policies and object storage interfaces are stored durably within RADOS and are managed by storing references in the object server daemon maps. Since the cluster already propagates a consistent view of these data structures, we use this service to automatically install interfaces in object storage daemons, and install policies within the metadata server daemons such that clients and daemons are synchronized on correct implementations without restarting.

Impact: this is a feature because it adds data safety and persistence to system metadata; while nice to have it does not necessarily effect correctness.

5. Services Built on Malacology

In this section we describe two services built on top of Malacology. The first is Mantle, a framework for dynamically specifying file system metadata load balancing policies. The second system, ZLog, is a high-performance distributed shared-log. In addition to these services, we will demonstrate how we combine ZLog and Mantle to implement service-aware metadata load balancing policies.

5.1 Mantle: Programmable Load Balancer

Mantle [37] is a programmable load balancer that separates the metadata balancing policies from their mechanisms. Administrators inject code to change how the metadata cluster distributes metadata. Our previous work showed how to use Mantle to implement a single node metadata service, a distributed metadata service with hashing, and a distributed metadata service with dynamic subtree partitioning.

The original implementation was “hard-coded” into Ceph and lacked robustness (no versioning, durability, or policy distribution). Re-implemented using Malacology, Mantle now enjoys (1) the versioning provided by Ceph’s monitor daemons and (2) the durability and distribution provided by Ceph’s reliable object store. Re-using the internal abstractions with Malacology resulted in a 2× reduction in source code compared to the original implementation.

5.1.1 Versioning Balancer Policies

Ensuring that the version of the current load balancer is consistent across the physical servers in the metadata cluster was not addressed in the original implementation. The user had to set the version on each individual server and it was trivial to make the versions inconsistent. Maintaining consistent versions is important for cooperative balancing policies, where local decisions are made assuming properties about other instances in the cluster.

With Malacology, Mantle stores the version of the current load balancer in the Service Metadata interface. The version of the load balancer corresponds to an object name in the balancing policy. Using the Service Metadata interface means Mantle inherits the consistency of Ceph’s internal monitor daemons. The user changes the version of the load balancer using a new CLI command.

5.1.2 Making Balancer Policies Durable

The load balancer version described above corresponds to the name of an object in RADOS that holds the actual Lua balancing code. When metadata server nodes start balancing load, they first check the latest version from the metadata server map and compare it to the balancer they have loaded. If the version has changed, they dereference the pointer to the balancer version by reading the corresponding object in RADOS. This is in contrast to the original Mantle implementation which stored load balancer code on the local file system – a technique which is unreliable and may result in silent corruption.

The balancer pulls the Lua code from RADOS synchronously; asynchronous reads are not possible because of the architecture of the metadata server. The synchronous behavior is not the default behavior for RADOS operations, so we achieve this with a timeout: if the asynchronous read does not come back within half the balancing tick interval the operation is canceled and a Connection Timeout error is

returned. By default, the balancing tick interval is 10 seconds, so Mantle will use a 5 second timeout.

This design allows Mantle to immediately return an error if anything RADOS-related goes wrong. We use this implementation because we do not want to do a blocking object storage daemon read from inside the global metadata server lock. Doing so would bring down the metadata server cluster if any of the object storage daemons are not responsive.

Storing the balancers in RADOS is simplified by the use of an interpreted language for writing balancer code. If we used a language that needs to be compiled, like the C++ object classes in the object storage daemon, we would need to ensure binary compatibility, which is complicated by different operating systems, distributions, and compilers.

5.1.3 Logging, Debugging, and Warnings

In the original implementation, Mantle would log all errors, warnings, and debug messages to a log stored locally on each metadata server. To get the simplest status messages or to debug problems, the user would have to log into each metadata server individually, look at the logs, and reason about causality and ordering.

With Malacology, Mantle re-uses the centralized logging features of the monitoring service. Important errors, warnings, and info messages are collected by the monitoring subsystem and appear in the monitor cluster log so instead of users going to each node, they can watch messages appear at the monitor daemon. Messages are logged sparingly, so as not to overload the monitor with frivolous debugging but important events, like balancer version changes or failed subsystems, show up in the centralized log.

5.2 ZLog: A Fast Distributed Shared Log

The second service implemented on Malacology is ZLog, a high-performance distributed shared-log that is based on the CORFU protocol [6]. The shared-log is a powerful abstraction used to construct distributed systems, such as metadata management [7] and elastic database systems [8–10]. However, existing implementations that rely on consensus algorithms such as Paxos funnel I/O through a single point introducing a bottleneck that restricts throughput. In contrast, the CORFU protocol is able to achieve high throughput using a network counter called a *sequencer*, that decouples log position assignment from log I/O.

While a full description of the CORFU system is beyond the scope of this paper, we briefly describe the custom storage device interface, sequencer service, and recovery protocol, and how these services are instantiated in the Malacology framework.

5.2.1 Sequencer

High-performance in CORFU is achieved using a sequencer service that assigns log positions to clients by reading from a volatile, in-memory counter which can run at a very high throughput and at low latency. Since the sequencer is cen-

tralized, ensuring serialization in the common case is trivial. The primary challenge in CORFU is handling the failure of the sequencer in a way that preserves correctness. Failure of the sequencer service in CORFU is handled by a recovery algorithm that recomputes the new sequencer state using a CORFU-specific custom storage interface to discover the tail of the log, while simultaneously invalidating stale client requests using an epoch-based protocol.

Sequencer interface. The sequencer resource supports the ability to *read()* the current tail value and get the *next()* position in the log which also atomically increments the tail position. We implement the sequencer service in Malacology using the File Type interface. This approach has the added benefit of allowing the metadata service to handle naming, by representing each sequencer instance in the standard POSIX hierarchical namespace. The primary challenge in mapping the sequencer resource to the metadata service is handling serialization correctly to maintain the global ordering provided by the CORFU protocol.

Initially we sought to directly model the sequencer service in Ceph as a non-exclusive, non-cacheable resource, forcing clients to perform a round-trip access to the resource at the authoritative metadata server for the sequencer inode. Interestingly, we found that the capability system in Ceph reduces metadata service load by allowing clients that open a shared file to temporarily obtain an exclusive cached copy of the resource, resulting in a round-robin, best-effort batching behavior. When a single client is accessing the sequencer resource it is able to increment the sequencer locally. Any competing client cannot query the sequencer until the metadata service has granted it access.

While unexpected, this discovery allowed us to explore an implementation strategy that we had not previously considered. In particular, for bursty clients, and clients that can tolerate increased latency, this mode of operation may allow a system to achieve much higher throughput than a system with a centralized sequencer service. We utilize the programmability of the metadata service to define a new policy for handling capabilities that controls the amount of time that clients are able to cache the sequencer resource. This allows an administrator or application to control the trade-off between latency and throughput beyond the standard best-effort policy that is present in Ceph by default.

In Section 6 we quantify the trade-offs of throughput and latency for an approach based on a round-robin batching mode, and compare this mode to one in which the metadata server mediates access to the sequencer state when it is being shared among multiple clients. Quantifying these trade-offs should provide administrators with guidelines for setting the tunables for different “caching” modes of the sequencer.

Balancing policies. As opposed to the batching mode for controlling access to the sequencer resource, more predictable latency can be achieved by treating the sequencer inode as a shared non-cacheable resource, forcing clients

to make a round-trip to the metadata service. However, the shared nature of the metadata service may prevent the sequencer from achieving maximum throughput. To address this issue we use the Load Balancing interface to construct a service-specific load balancing policy. As opposed to a balancing policy that strives for uniform load distribution, a ZLog-specific policy may utilize knowledge of inode types to migrate the sequencer service to provisioned hardware during periods of contention or high demand.

5.2.2 Storage Interface

The storage interface is a critical component in the CORFU protocol. Clients independently map log positions that they have obtained from the sequencer service (described in detail in the next section) onto storage devices, while storage devices provide an intelligent *write-once*, random *read* interface for accessing log entries. The key to correctness in CORFU lies with the enforcement of up-to-date epoch tags on client requests; requests tagged with out-of-date epoch values are rejected, and clients are expected to request a new tail from the sequencer after refreshing state from an auxiliary service. This mechanism forms the basis for sequencer recovery.

In order to repopulate the sequencer state (*i.e.* the cached, current tail of the log) during recovery of a sequencer, the maximum position in the log must be obtained. To do this, the storage interface exposes an additional *seal* method that atomically installs a new epoch value and returns the maximum log position that has been written.

Since the sequencer service does not resume until the recovery process has completed, there cannot be a race with clients appending to the log, and the immutability of the log allow reads to never block during a sequencer failure. Recovery of the sequencer process itself may be handled in many ways, such as leader election using an auxiliary service like Paxos. In our implementation, the recovery is the same as (and is inherited from) the CephFS metadata service. Handling the failure of a client that holds the sequencer state is similar, although a timeout is used to determine when a client should be considered unavailable.

6. Evaluation

Our evaluation demonstrates the feasibility of building new service abstractions atop programmable storage, focusing on the performance of the internal abstractions exposed by Malacology and used to construct the Mantle and ZLog services. We also discuss latent capabilities we discovered in this process that let us navigate different trade-offs within the services themselves. First, we benchmark scenarios with high sequencer contention by examining the interfaces used to map ZLog onto Malacology; specifically, we describe the sequencer implementation and the propagation of object and data interfaces interfaces. Next, we benchmark scenarios in

which the storage system manages multiple logs by using Mantle to balance sequencers across a cluster.

Since this work focuses on the programmability of Malacology, the goal of this section is to show that the components and subsystems that support the Malacology interfaces provide reasonable relative performance, as well as to give examples of the flexibility that Malacology provides to programmers. This section uses a principled approach for evaluating tunables of the interfaces and the trade-offs we discuss should be acknowledged when building higher-level services.

6.1 Mapping ZLog onto Malacology

We evaluate Malacology by exploring one possible mapping of the ZLog implementation of CORFU onto Ceph in which we re-use (1) the metadata service to manage naming and synchronization of the sequencer resource by treating the resource as an inode, and (2) the monitoring sub-subsystem to distribute and install application-specific I/O interfaces required of the CORFU protocol. In Section 6.2 we then demonstrate how the re-use of the inode abstraction for implementing the sequencer resource enables load balancing policies to migrate the sequencer resource in heavy-load situation.

6.1.1 Sequencer Implementation

We evaluate the feasibility of using the metadata service to implement a sequencer resource that is responsible for maintaining a total ordering of the log. Clients contact the sequencer to obtain the tail of the log and then independently initiate I/O, thus we measure both the throughput and latency of obtaining new tail positions which bounds client append performance.

The sequencer is implemented using the File Type interface so that the sequencer state (a 64-bit integer) is embedded in the inode of a file. A total ordering of the log is imposed by the re-use of the capability service that can be used to grant exclusive access of inode state to clients. The metadata service is responsible for maintaining exclusivity and granting access. Figure 5 (a) shows the behavior of the system in which a best-effort policy is used. The two colors represent points in time that the clients were able to access the resource. The best-effort policy shows a high degree of interleaving between clients but the system spends a large portion of time re-distributing the capability, reducing overall throughput.

In order to control the performance of the system we implement a policy that (1) restricts the length of time that a client may maintain exclusive access and (2) limits the number of log positions that a client may generate without yielding to other clients waiting for access. The behavior of these two modes is illustrated in Figures 5 (b) and (c), respectively.

Figure 6 demonstrates a configurable trade-off between throughput and latency. In the experiment two clients are run

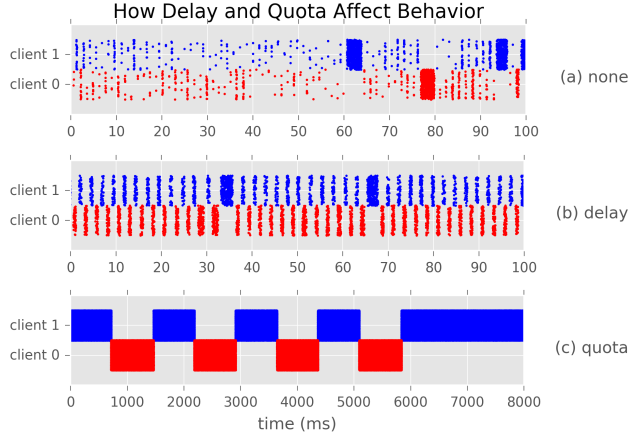


Figure 5: [source] Each dot is an individual request, spread randomly along the y axis. The default behavior is unpredictable, “delay” lets clients hold the lease longer, and “quota” gives clients the lease for a number of operations.

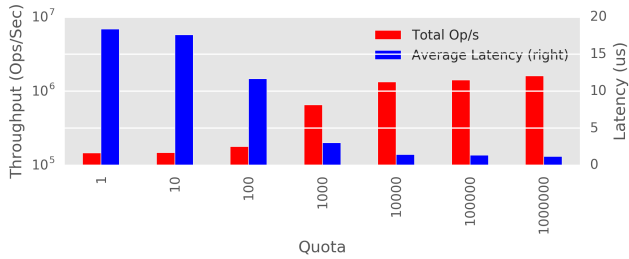


Figure 6: [source] Sequencer throughput by re-using various services. The highest performance is achieved using a single client with exclusive, cacheable privilege. Round-robin sharing of the sequencer resource is affected by the amount of time the resource is held, with best-effort performing the worst.

each with a fixed 0.25 second maximum reservation on the capability, and we vary the size of the log position quota running each configuration for two minutes. The total operations per second is the combined throughput of the two clients, and the average latency is the number of microseconds required to obtain a new log position. With a small quota more time is spent exchanging exclusive access, while a large quota reservation allows clients to experience a much lower latency because they experience isolated access for a longer period of time.

To get a better picture of latency, Figure 7 shows the CDF of latency for each client in all experiment configurations. At the 99th percentile clients accessed the sequencer in less than a millisecond. The CDF is cropped at the 99.999th percentile due to large outliers that we believe occur in instances in which the metadata server is performing I/O while it is in the process of re-distributing the capability to another client.

Malacology exposes the internal capability management service and allows users to navigate latency and throughput trade-offs. Other approaches to designing the sequencer service also exist, such as using a centralized service in which each access is a network round-trip. In contrast to the mech-

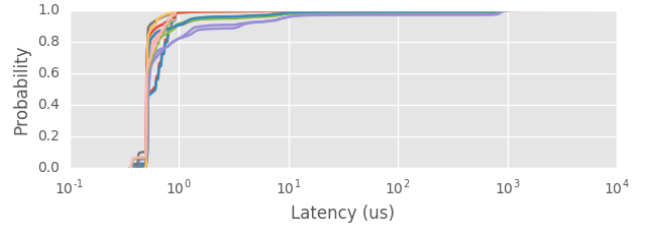


Figure 7: [source] The latency distribution shows that if the clients hold their capabilities longer, performance decreases. Changing the delay can help the sequencer and clients strike a balance between throughput and latency.

anism we explored which is appropriate for clients with bursty workloads, it may be easier to provide predictable performance using a centralized service, and we will be exploring in the future how this can be achieved using the capability system.

6.1.2 Interface Propagation

Domain-specific data interfaces (Section 2) allow co-design between applications and the storage system. Malacology supports custom object interfaces in RADOS that require interface implementations to be installed on the storage devices in the system, supporting the evolution of interfaces through automatic system-wide versioning and installation through the service metadata interface (Section 4.1). We evaluate the performance of installing a new interface version in the cluster, which is an important metric for applications that frequently evolve interfaces.

We demonstrate the feasibility of utilizing the Ceph monitoring sub-system by evaluating the performance of installing and distributing interface updates. Figure 8 shows the CDF of the latency of interface updates. The interfaces are Lua scripts embedded in the cluster map and distributed using a peer-to-peer gossip protocol. The latency is defined as the elapsed time following the Paxos proposal for an interface update until each object storage daemon makes the update live (the cost of the Paxos proposal is configurable and is discussed below). The latency measurements were taken on the nodes running object server daemons, and thus exclude the client round-trip cost. In each of the experiments 1000 interface updates were observed.

Figure 8 shows the lower bound cost for updates in a large cluster. In the experiment labeled “120 OSD (RAM)” a cluster of 120 object storage daemons (OSDs) using an in-memory data store were deployed, showing a latency of less than 54 ms with a probability of 90% and a worst case latency of 194 ms. These costs demonstrate the penalty of distributing the interface in a large cluster. In practice the costs include, in addition to cluster-wide propagation of interface updates, the network round-trip to the interface management service, the Paxos commit protocol itself, and other factors such as system load. By default Paxos proposals occur periodically with a 1 second interval in order to accumulate

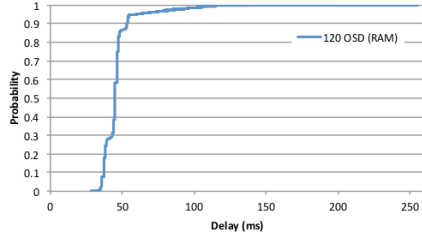


Figure 8: [source] Cluster-wide interface update latency, excluding the Paxos proposal cost for committing the Service Metadata interface.

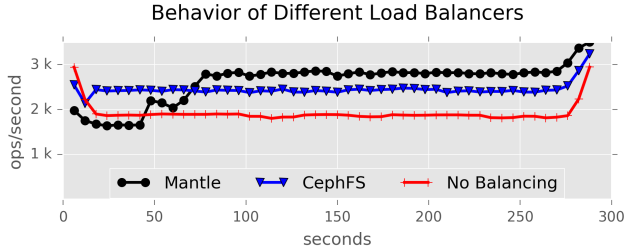


Figure 9: [source] CephFS/Mantle load balancing have better throughput than co-locating all sequencers on the same server. Sections 6.2.1 and 6.2.2 quantify this improvement; Section 6.2.3 examines the migration at 0-60 seconds.

updates. In a minimum, realistic quorum of 3 monitors using hard drive-based storage, we were able to decrease this interval to an average of 222 ms.

6.2 Load Balancing ZLog Sequencers with Mantle

In practice, a storage system implementing CORFU will support a multiplicity of independent totally-ordered logs for each application. For this scenario co-locating sequencers on the same physical node is not ideal but building a load balancer that can migrate the shared resource (e.g., the resource that mediates access to the tail of the log) is a time-consuming, non-trivial task. It requires building subsystems for migrating resources, monitoring the workloads, collecting metrics that describe the utilization on the physical nodes, partitioning resources, maintaining cache coherence, and managing multiple sequencers. The following experiments demonstrate the feasibility of using the mechanisms of the Malacology Load Balancing interface to inherit these features and to alleviate load from overloaded servers.

The experiments are run on a cluster with 10 nodes to store objects, one node to monitor the cluster, and 3 nodes that can accommodate sequencers. Instead of measuring contention at the clients like Section 6.1.1, these experiments measure contention at the sequencers by forcing clients to make round-trips for every request. We implement this using the Shared Resource interface that forces round-trips. Because the sequencer’s only function is to hand out positions for the tail of the log, the workload is read-heavy.

First, we show how the ZLog service can orchestrate multiple sequencers using the Malacology Load Balancing interface. Figure 9 shows the throughput over time of different

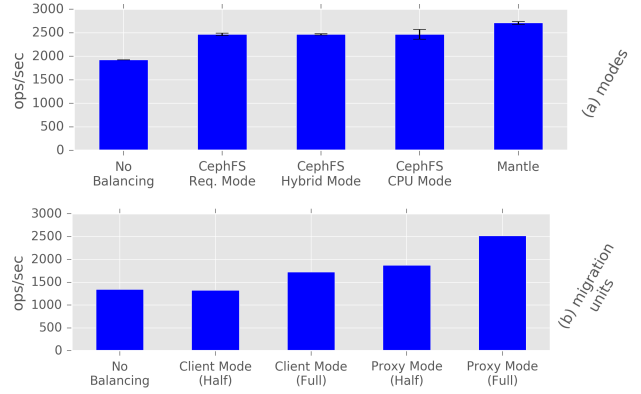


Figure 10: [source, source] In (a) all CephFS balancing modes have the same performance; Mantle uses a balancer designed for sequencers. In (b) the best combination of mode and migration units can have up to a 2× improvement.

load balancers as they migrate 3 sequencers (with 4 clients) around the cluster; “No Balancing” keeps all sequencers on one server, “CephFS” migrates sequencers using the hard-coded CephFS load balancers, and “Mantle” uses a custom load balancer we wrote specifically for sequencers. The increased throughput for the CephFS and Mantle curves between 0 and 60 seconds are a result of migrating the sequencer(s) off overloaded servers.

In addition to showing that migrating sequencers improves performance, Figure 9 also demonstrates features that we will explore in the rest of this section. Sections 6.2.1 and 6.2.2 quantify the differences in performance when the cluster stabilizes at time 100 seconds and Section 6.2.3 examines the slope and start time of the re-balancing phase between 0 and 60 seconds by comparing the aggressiveness of the balancers.

6.2.1 Feature: Balancing Modes

Next, we quantify the performance benefits shown in Figure 9. To understand why load balancers perform differently we need to explain the different balancing modes that the load balancer service uses and how they stress the subsystems that receive and forward client requests in different ways. In Figure 9, the CephFS curve shows the performance of the balancing mode that CephFS falls into *most of the time*. CephFS currently has 3 modes for balancing load: CPU mode, workload mode, and hybrid mode. All three have the same structure for making migration decisions but vary based on the metric used to calculate load. For this sequencer workload the 3 different modes all have the same performance, shown in Figure 10 (a), because the load balancer falls into the same mode a majority of the time. The high variation in performance for the CephFS CPU Mode bar reflects the uncertainty of using something as dynamic and unpredictable as CPU utilization to make migration decisions. In addition to the suboptimal performance and unpredictability, another problem is that all the CephFS bal-

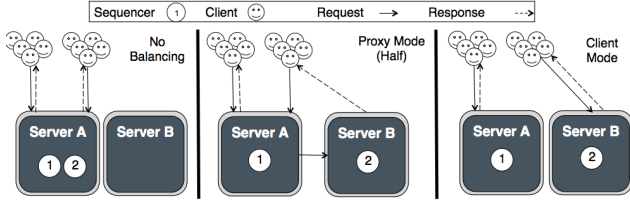


Figure 11: In client mode clients sending requests to the server that houses their sequencer. In proxy mode clients continue sending their requests to the first server.

ancers behave the same. This prevents administrators from properly exploring the balancing state space.

Mantle gives the administrator more control over balancing policies; for the Mantle bar in Figure 10 (a) we use the Load Balancing interface to program logic for balancing read-heavy workloads, resulting in better throughput and stability. When we did this we also identified two balancing modes relevant for making migration decisions for sequencers.

Using Mantle, the administrator can put the load balancer into “proxy mode” or “client mode”. In proxy mode one server receives all requests and farms off the requests to slave servers; the slave servers do the actual tail finding operation. In client mode, clients interact directly with the server that has their sequencer. These modes are illustrated in Figure 11. “No Balancing” is when all sequencers are co-located on one physical server – performance for that mode is shown by the “No Balancing” curve in Figure 9. In “Proxy Mode”, clients continue sending requests to server A even though some of the sequencers have been migrated to another server. Server A redirects client requests for sequencer 2 to server B. “Proxy Mode (Half)” is shown in Figure 9; in this scenario, half of the sequencers have migrated off the first server. Alternatively, “Proxy Mode (Full)”, which is not pictured, is when all the sequencers migrate off the first server. “Client Mode”, shown on the far right of Figure 11, shows how clients for sequencer 2 contact server B without a redirect from server A.

Figure 12 shows the throughput over time of the two different modes for an environment with only 2 sequencers (again 4 clients each) and 2 servers. The curves for both sequencers in Figure 12(a) start at less than 1000 ops/second and at time 60 seconds Mantle migrates Sequencer 1 to the slave server. Performance of Sequencer 2 decreases because it stayed on the proxy which now processes requests for Sequencer 2, and forwards requests for Sequencer 1. The performance of Sequencer 1 improves dramatically because distributing the sequencers in this way separates (1) the handling of the client requests and (2) finding the tail of the log and responding to clients. Doing both steps is too heavy weight for one server and sequencers on slave nodes can go faster if work is split up; this phenomenon is not uncommon and has been observed in chain replication [44].

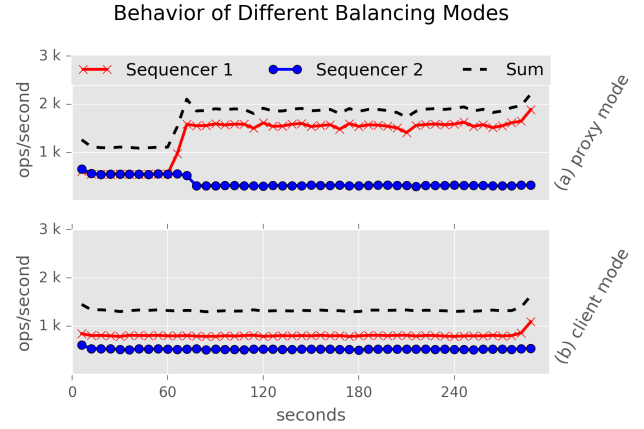


Figure 12: [source] The performance of proxy mode achieves the highest throughput but at the cost of lower throughput for one of the sequencers. Client mode is more fair but results in lower cluster throughput.

Cluster throughput improves at the cost of decreased throughput for Sequencer 2. Figure 12(b) is set to sequencer mode manually (no balancing phase) and shows that the cluster throughput is worse than the cluster throughput of proxy mode. That graph also shows that Sequencer 2 has less throughput than Sequencer 1. In this case, the scatter-gather process used for cache coherence in the metadata protocols causes strain on the server housing Sequencer 2 resulting in this uneven performance.

6.2.2 Feature: Migration Units

Another factor that affects performance in this environment is how much load is on each server; these experiments quantify that effect by programming the Load Balancing interface to control the amount of load to migrate. We call this metric a “migration unit”. Expressing this heuristic is not easily achievable with outward facing tunable parameters (i.e. system knobs) but with Mantle’s programmable interface, we can force the load balancer to change its migration units. To force the balancer into the Proxy Mode (Half) scenario in Figure 11, which uses migration units equal to half the load on the current server, we can use: `targets[whoami+1] = mds[whoami]["load"]/2`.

This code snippet uses globally defined variables and tables from the Mantle API [37] to send half of the load on the current server (whoami) to the next ranked server (whoami + 1); the targets array is a globally defined table that the balancer uses to do the migrations. Alternatively, to migrate all load a time step, we can remove the division by 2.

Figure 10 (b) shows the performance of the modes using different migration units. Recall that this setup only has 2 sequencers and 2 servers, so performance may be different at scale. Even so, it is clear that client mode does not perform as well for read-heavy workloads. We even see a throughput improvement when migrating all load off the first server, leaving the first server to do administrative tasks (this is com-

mon in the metadata cluster because the first server does a lot of the cache coherence work) while the second server does all the processing. Proxy mode does the best in both cases and shows large performance gains when completely decoupling client request handling and operation processing in Proxy Mode (Full). The parameter that controls the migration units helps the administrator control the sequencer co-location or distribution across the cluster. This trade-off was explored extensively in the Mantle paper but the experiments we present here are indicative of an even richer set of states to explore.

6.2.3 Feature: Backoff

Tuning the aggressiveness of the load balancer decision making is also a trade-off that administrators can control and explore. The balancing phase from 0 to 60 seconds in Figure 9 shows different degrees of aggressiveness in making migration decisions; CephFS makes a decision 10 seconds into the run and throughput jumps to 2500 ops/second while Mantle takes more time to stabilize. This conservative behavior is controlled by programming the balancer to (1) use different conditions for when to migrate and (2) using a threshold for sustained overload.

We control the conditions for when to migrate using `when()`, a callback in the Mantle API. For the Mantle curve in Figure 9 we program `when()` to wait for load on the receiving server to fall below a threshold. This makes the balancer more conservative because it takes 60 seconds for cache coherence messages to settle. The Mantle curve in Figure 9 also takes longer to reach peak throughput because we want the policy to wait to see how migrations affect the system before proceeding; the balancer does a migration right before 50 seconds, realizes that there is a third underloaded server, and does another migration.

The other way to change aggressiveness of the decision making is to program into the balancer a threshold for sustained overload. This forces the balancer to wait a certain number of iterations after a migration before proceeding. In Mantle, the policy would use the `save state` function to do a countdown after a migration. Behavior graphs and performance numbers for this backoff feature is omitted for space considerations, but our experiments confirm that the more conservative the approach the less overall throughput.

Malacology pulls the load balancing service out of the storage system to balance sequencers across a cluster. This latent capability also gives future programmers the ability to explore the different load balancing trade-offs including: load balancing modes to control forwarding vs. client redirection, load migration units to control sequencer distribution vs. co-location, and backoffs to control conservative vs. aggressive decision making.

7. Future Work

Malacology is a first step towards showing how general-purpose storage systems can be adapted to target special-purpose applications. By encapsulating storage system functionality as reusable building blocks, we enable application developers to leverage storage capabilities based on interfaces that are proven and understandable. However, creation and composition of interfaces is complex; constructs must be combined safely in order to provide correctness, performance and security. We will study additional Malacology-based services in order to learn techniques that support safe composition.

Some higher-level services that we plan to build using the interfaces in Table 2 are: an elastic cloud database, a data processing engine, and a data layout manager. Approaches proposed so far use the Data I/O interface to push down predicates and computation, the File Type interface to maintain access paths and metadata efficiently, and the Durability interface to manage ingestion and movement. Using the programmable storage approach helps us build higher-level services that work well with the storage system not in spite of it.

Our experience with ZLog and Mantle demonstrates that the labor of wrapping existing services in reusable interfaces is justified by the power and flexibility that this encapsulation affords to programmers. In exchange for this flexibility, however, programmers may forfeit the *protection from change* afforded by narrow storage interfaces such as the POSIX API. To implement applications on programmable storage systems such as Malacology, programmers must find solutions by navigating a complex design space, simultaneously addressing functional correctness, performance and fault tolerance. Worse still, their solutions may be sensitive to changes in the underlying environment, such as hardware upgrades, software version changes and evolving workloads. For example, a major version change in Ceph required us to rewrite significant parts of ZLog to maintain acceptable performance. Each such evolution costs developer time and risks introducing bugs.

We are actively exploring the use of high-level *declarative* languages based on Datalog [2] to program data access and storage APIs. Using this approach, a systems programmer can specify the functional behavior in a relational (or algebraic) language, allowing an optimizer to search through the space of functionally equivalent physical implementations and select a good execution plan, re-optimizing when storage characteristics or statistics change. Much like query planning and optimization in database systems [24], this approach will separate the concerns of correctness and performance, protecting applications (which usually evolve slowly) against changes in more dynamic storage system.

8. Related Work

Programmability of operating systems and networking resources, including distributed storage systems is not new, but we are not aware of work that makes generalization of existing services into programmable resources a key principle in storage systems design.

Programmable storage systems can be viewed as an infrastructure for creating abstractions to better separate policies from mechanisms. This idea is not new. Software-defined networks (SDNs) create such an abstraction by separating the control plane from the data plane (see for example [27]). This notion of control/data separation was also applied in software-defined storage (SDS) [41, 43]. Similarly, IOStack [19] is providing policy-based provisioning and filtering in OpenStack Swift. According to a SNIA white paper [14], the primary goal of SDS is to control and facilitate flexible and dynamic provisioning of storage resources of different kinds, including flash memory and disk drives, to create a virtual mapping between common storage abstractions (e.g. files, objects, and blocks) and storage devices taking data service objectives in terms of protection, availability, performance, and security into account. A programmable storage system exposes internal abstractions so that end users (not necessarily operators) can create new services on top of the storage stack. Thus, our notion of programmable storage differs from “software-defined storage” (SDS) in terms of goals and scope, although definitions of SDS are still in flux.

Another view of programmable storage systems is one of tailoring systems resources to applications [5]. Related efforts include the Exokernel [15], SPIN [11] and Vino [36] projects; the latter two addressed the ability of injecting code into the kernel to specialize resource management. Another approach is to pass hints between the different layers of the I/O stack to bridge the semantic gap between applications and storage [5, 33, 39].

Malacology uses the same Active and Typed Storage module presented in DataMods [47]; Asynchronous Service and File Manifolds can be implemented with small changes to the Malacology framework, namely asynchronous object calls and Lua stubs in the inode, respectively.

9. Conclusion

Programmable storage is a viable method for eliminating duplication of complex error-prone software used as workarounds for storage system deficiencies. We propose that systems expose their services in a safe way allowing application developers to customize system behavior to meet their needs while not sacrificing correctness. To illustrate the benefits of this approach we presented Malacology², a programmable storage system that facilitates the construction

of new services by re-purposing existing subsystem abstractions of the storage stack.

Acknowledgments

We thank the EuroSys reviewers for their hard work, attentiveness, and genuinely helpful suggestions. We especially thank Mahesh Balakrishnan for shepherding the paper. This work was partially funded by the Center for Research in Open Source Software³, the DOE Award DE-SC0016074, and the NSF Award 1450488.

Note: this paper follows The Popper Convention⁴ [28]. All the experiments presented here are available on the repository associated to this article⁵. For every figure, a [source] link points to a Jupyter notebook that shows the analysis from where the graph was obtained; its parent folder contains all the associated artifacts.

References

- [1] Ceph Architecture. URL <http://docs.ceph.com/docs/master/architecture>.
- [2] P. Alvaro, N. Conway, J. M. Hellerstein, and W. R. Marczak. Consistency Analysis in Bloom: A CALM and Collected Approach. In *Proceedings 5th Biennial Conference on Innovative Data Systems Research*, CIDR '11, Asilomar, CA, January 2011.
- [3] Apache Parquet Contributors. Parquet Columnar Storage Format, <http://parquet.io>.
- [4] M. Armbrust, A. Fox, R. Griffith, A. D. Joseph, R. Katz, A. Konwinski, G. Lee, D. Patterson, A. Rabkin, I. Stoica, and M. Zaharia. A View of Cloud Computing. *Communications of the ACM*, vol. 53, 2010.
- [5] A. C. Arpaci-Dusseau and R. H. Arpaci-Dusseau. Information and Control in Gray-box systems. In *Proceedings of the 18th ACM Symposium on Operating Systems Principles*, SOSP '01, Banff, Alberta, Canada, 2001.
- [6] M. Balakrishnan, D. Malkhi, V. Prabhakaran, T. Wobber, M. Wei, and J. D. Davis. CORFU: A Shared Log Design for Flash Clusters. In *Proceedings of the 9th USENIX Conference on Networked Systems Design and Implementation*, NSDI '12, San Jose, CA, April 2012.
- [7] M. Balakrishnan, D. Malkhi, T. Wobber, M. Wu, V. Prabhakaran, M. Wei, J. D. Davis, S. Rao, T. Zou, and A. Zuck. Tango: Distributed Data Structures Over a Shared Log. In *Proceedings of the 24th ACM Symposium on Operating Systems Principles*, SOSP '13, Farmington, PA, November 2013.
- [8] P. A. Bernstein, C. W. Reid, and S. Das. Hyder – A Transactional Record Manager for Shared Flash. In *Proceedings 5th Biennial Conference on Innovative Data Systems Research*, CIDR '11, Asilomar, CA, January 2011.
- [9] P. A. Bernstein, C. W. Reid, M. Wu, and X. Yuan. Optimistic Concurrency Control by Melding Trees. In *Proceedings of*

³<http://cross.ucsc.edu>

⁴<http://falsifiable.us>

⁵<https://github.com/michaelsevilla/malacology-popper/tree/v2.1>

²<http://programmability.us>

the 37th International Conference on Very Large Data Bases, VLDB '11, August 2011.

- [10] P. A. Bernstein, S. Das, B. Ding, and M. Pilman. Optimizing Optimistic Concurrency Control for Tree-Structured, Log-Structured Databases. In *Proceedings of the ACM International Conference on Management of Data, SIGMOD '15*, Melbourne, Australia, May 2015.
- [11] B. N. Bershad, S. Savage, P. Pardyak, E. G. Sirer, M. E. Ficzynski, D. Becker, C. Chambers, and S. Eggers. Extensibility Safety and Performance in the SPIN Operating System. In *Proceedings of the 15th ACM Symposium on Operating Systems Principles, SOSP '95*, Copper Mountain, CO, December 1995.
- [12] E. Brewer, L. Ying, L. Greenfield, R. Cypher, and T. T'so. Disks for Data Centers. Technical Report, Google, 2016.
- [13] M. Burrows. The Chubby Lock Service for Loosely-Coupled Distributed Systems. In *Proceedings of the 7th Symposium on Operating Systems Design and Implementation, OSDI '06*, Seattle, WA, November 2006.
- [14] M. Carlson, A. Yoder, L. Schoeb, D. Deel, C. Pratt, C. Lionetti, and D. Voigt. Software Defined Storage. *SNIA Whitepaper*, January 2015.
- [15] D. R. Engler, M. F. Kaashoek, and J. J. O'Toole. Exokernel: An Operating System Architecture for Application-Level Resource Management. In *Proceedings of the 15th ACM Symposium on Operating Systems Principles, SOSP '95*, Copper Mountain, CO, December 1995.
- [16] Epping, Duncan and Denneman, Frank. VMware vSphere 5.1 Clustering Deepdive, accessed 03/21/2014, <http://www.vmware.com/product/drs>.
- [17] R. Geambasu, A. A. Levy, T. Kohno, A. Krishnamurthy, and H. M. Levy. Comet: An Active Distributed Key-Value Store. In *Proceedings of the 9th USENIX conference on Operating Systems Design and Implementation, OSDI '10*, Vancouver, Canada, October 2010.
- [18] S. Ghemawat, H. Gobioff, and S.-T. Leung. The Google File System. In *Proceedings of the 14th ACM Symposium on Operating Systems Principles, SOSP '03*, Bolton Landing, NY, October 2003. ACM.
- [19] R. Gracia-Tinedo et al. IOSTack: Software-Defined Object Storage. *IEEE Internet Computing*, 20(3):10–18, May-June 2016.
- [20] M. Grawinkel, T. Sub, G. Best, I. Popov, and A. Brinkmann. Towards Dynamic Scripted pNFS Layouts. In *Proceedings of the 2012 SC Companion: High Performance Computing, Networking Storage and Analysis, SCC '12*, Salt Lake City, UT, November 2012.
- [21] J. Gray. Tape is Dead, Disk is Tape, Flash is Disk, RAM Locality is King. *CIDR 2007 - Gong Show Presentation*, January 2007.
- [22] J. Gray and B. Fitzgerald. Flash Disk Opportunity for Server Applications. *Queue*, vol. 6, July 2008.
- [23] A. Gulati, G. Shanmuganathan, A. Holler, and I. Ahmad. Cloud-Scale Resource Management: Challenges and Techniques. In *Proceedings of the 3rd USENIX Conference on Hot Topics in Cloud Computing, HotCloud '11*, Portland, OR, June 2011.
- [24] J. M. Hellerstein and M. Stonebraker. Anatomy of a Database System. *Readings in Database Systems*, January 2005.
- [25] P. Hunt, M. Konar, F. P. Junqueira, and B. Reed. ZooKeeper: Wait-free Coordination for Internet-Scale Systems. In *Proceedings of the USENIX Annual Technical Conference, ATC '10*, Boston, MA, June 2010.
- [26] R. Ierusalimsky, L. H. De Figueiredo, and W. Celes Filho. Lua - An Extensible Extension Language. *Software Practical Experiences*, 26(6):635–652, 1996.
- [27] S. Jain, A. Kumar, S. Mandal, J. Ong, L. Poutievski, A. Singh, S. Venkata, J. Wanderer, J. Zhou, M. Zhu, J. Zolla, U. Hlzle, S. Stuart, and A. Vahdat. B4: Experience with a Globally-Deployed Software Defined WAN. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications, SIGCOMM '13*, Hong Kong, China, August 2013.
- [28] I. Jimenez, M. Sevilla, N. Watkins, C. Maltzahn, J. Lofstead, K. Mohror, R. Arpaci-Dusseau, and A. Arpaci-Dusseau. Popper: Making Reproducible Systems Performance Evaluation Practical, UCSC-SOE-16-10. Technical Report UCSC-SOE-16-10, UC Santa Cruz, May 2016.
- [29] L. Joao. Ceph's New Monitor Changes. URL <https://ceph.com/dev-notes/cephs-new-monitor-changes>.
- [30] L. Lamport. The Part-Time Parliament. *ACM Transactions on Computer Systems*, 16(2):133–169, May 1998.
- [31] Linux Foundation. Kinetic Open Storage Project, 2015. URL <https://www.openkinetic.org/>.
- [32] J. MacCormick, N. Murphy, M. Najork, andramohan A. Thekkath, and L. Zhou. Boxwood: Abstractions as the Foundation for Storage Infrastructure. In *Proceedings of the 6th USENIX Symposium on Operating Systems Design and Implementation, OSDI '04*, San Francisco, CA, December 2004.
- [33] M. Mesnier, F. Chen, and J. B. Akers. Differentiated Storage Services. In *Proceedings of the 23rd ACM Symposium on Operating Systems Principles, SOSP '11*, Cascais, Portugal, October 2011.
- [34] Rackspace. ZeroVM and OpenStack Swift. URL <http://www.zerovm.org/zerocloud.html>.
- [35] E. Riedel, G. A. Gibson, and C. Faloutsos. Active Storage For Large-Scale Data Mining and Multimedia. In *Proceedings of the 24th international Conference on Very Large Databases, VLDB '98*, New York, NY, July 1998.
- [36] M. I. Seltzer, Y. Endo, C. Small, and K. A. Smith. Dealing with Disaster: Surviving Misbehaved Kernel Extensions. In *Proceedings of the 2nd Symposium on Operating Systems Design and Implementation, OSDI '96*, Seattle, WA, October 1996.
- [37] M. A. Sevilla, N. Watkins, C. Maltzahn, I. Nassi, S. A. Brandt, S. A. Weil, G. Farnum, and S. Fineberg. Mantle: A Programmable Metadata Load Balancer for the Ceph File System. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis, SC '15*, November 2015.

- [38] K. V. Shvachko, H. Kuang, S. Radia, and bert Chansler. The Hadoop Distributed File System. In *Proceedings of the 26th Symposium on Mass Storage Systems and Technologies, MSST '10*, Incline Village, NV, May 2010.
- [39] M. Sivathanu, V. Prabhakaran, F. I. Popovici, T. E. Denehy, A. C. Arpaci-Dusseau, and R. H. Arpaci-Dussea. Semantically-Smart Disk Systems. In *Proceedings of the 2nd USENIX Conference on File and Storage Technologies, FAST '03*, San Francisco, CA, March 2003.
- [40] SNIA. Implementing Multiple Cloud Storage APIs, November 2014. URL <http://www.sniacloud.com/?p=88>.
- [41] I. Stefanovici, B. Schroeder, G. O'Shea, and E. Thereska. sRoute: Treating the Storage Stack Like a Network. In *Proceedings of the 15th USENIX Conference on File and Storage Technologies, FAST '16*, Santa Clara, CA, February 2016.
- [42] R. Thakur, W. Gropp, and E. Lusk. On Implementing MPI-IO Portably and with High Performance. In *Proceedings of the th Workshop on I/O in Parallel and Distributed Systems, IOPADS '99*, Atlanta, Georgia, May 1999.
- [43] E. Thereska, H. Ballani, G. O'Shea, T. Karagiannis, A. Rowstron, T. Talpey, R. Black, and T. Zhu. IOFlow: A Software-Defined Storage Architecture. In *Proceedings of the 24th ACM Symposium on Operating Systems Principles, SOSP '13*, Farmington, PA, November 2013.
- [44] R. van Renesse and F. B. Schneider. Chain Replication for Supporting High Throughput and Availability. In *Proceedings of the 6th Conference on Symposium on Operating Systems Design and Implementation, OSDI '04*, San Francisco, CA, December 2004.
- [45] L. Vieira Neto, R. Ierusalimsky, A. L. de Moura, and M. Balmer. Scriptable Operating Systems with Lua. In *Proceedings of the 10th ACM Symposium on Dynamic Languages, DLS '14*, New York, NY, 2014.
- [46] N. Watkins. Dynamic Object Interfaces with Lua. URL <http://ceph.com/rados/dynamic-object-interfaces-with-lua>.
- [47] N. Watkins, C. Maltzahn, S. Brandt, and A. Manzanares. DataMods: Programmable File System Services. In *Proceedings of the 6th Workshop on Parallel Data Storage, PDSW '12*, Salt Lake City, Utah, November 2012.
- [48] N. Watkins, C. Maltzahn, S. Brandt, I. Pye, and A. Manzanares. In-Vivo Storage System Development. In *EuroPar: Parallel Processing Workshops*, Aachen, Germany, August 2013.
- [49] S. A. Weil, K. T. Pollack, S. A. Brandt, and E. L. Miller. Dynamic Metadata Management for Petabyte-Scale File Systems. In *Proceedings of the 17th ACM/IEEE Conference on Supercomputing, SC '04*, Pittsburgh, PA, November 2004.
- [50] S. A. Weil, A. W. Leung, S. A. Brandt, and C. Maltzahn. RADOS: A Scalable, Reliable Storage Service for Petabyte-Scale Storage Clusters. In *Proceedings of the 2nd International Workshop on Petascale Data Storage, PDSW '07*, Reno, NV, November 2007.