

GSoC 2024 Apache OpenDAL OVFS Project Proposal

1 Project Abstract

Virtio is an open standard designed to enhance I/O performance between virtual machines (VMs) and host systems in virtualized environments. VirtioFS is an extension of the Virtio standard specifically crafted for file system sharing between VMs and the host. This is particularly beneficial in scenarios where seamless access to shared files and data between VMs and the host is essential. VirtioFS has been widely adopted in virtualization technologies such as QEMU and Kata Container.

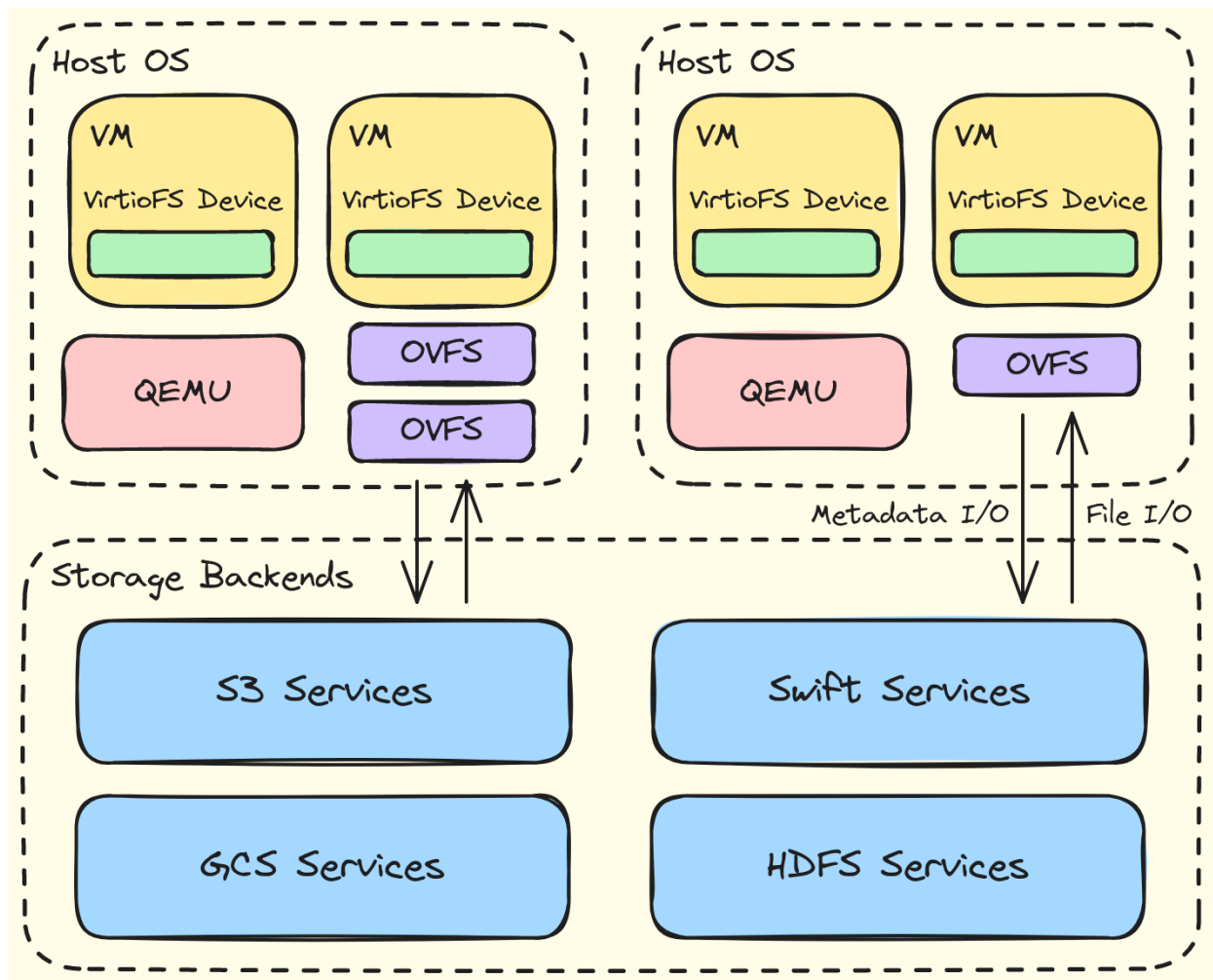
Apache OpenDAL is a data access layer that allows users to easily and efficiently retrieve data from various storage services in a unified manner. In this project, our goal is to reference virtiofsd (a standard vhost-user backend, a pure Rust implementation of VirtioFS based on the local file system) and implement VirtioFS based on OpenDAL.

Through this project, VMs can access numerous data services through the file system interface with the assistance of the OpenDAL service daemon deployed on the host, all without their awareness. It ensures the efficiency of file system reading and writing in VMs through VirtioFS support. This storage-system-as-a-service approach conceals the details of the distributed storage system from VMs. This ensures the security of storage services, as VMs do not need to be aware of the information, configuration and permission credentials of the accessed storage service. Additionally, it enables the utilization of a new backend storage system without reconfiguring all VMs.

2 Project Detailed Description

This chapter serves as an introduction to the overall structure of the project, outlining the design ideas and principles of critical components. It covers the OVFS architecture, interaction principles, design philosophy, file system operations based on various storage backend, cache pool design, configuration support, potential usage scenarios of OVFS, and the expected file system interface support.

2.1 The Architecture of OVFS



The picture above is the OVFS architecture diagram. OVFS is a file system implementation based on the VirtioFS protocol and OpenDAL. It serves as a bridge for semantic access to file system interfaces between VMs and external storage

systems. Leveraging the multiple service access capabilities and unified abstraction provided by OpenDAL, OVFS can conveniently mount shared directories in VMs on various existing distributed storage services.

The complete OVFS architecture consists of three crucial components:

- 1) VMs FUSE client that supports the VirtioFS protocol and implements the VirtioFS Virtio device specification. An appropriately configured Linux 5.4 or later can be used for OVFS. The VirtioFS protocol is built on FUSE and utilizes the VirtioFS Virtio device to transmit FUSE messages. In contrast to traditional FUSE, where the file system daemon runs in the guest user space, the VirtioFS protocol supports forwarding file system requests from the guest to the host, enabling related processes on the host to function as the guest's local file system.
- 2) A hypervisor that implements the VirtioFS Virtio device specification, such as QEMU. The hypervisor needs to adhere to the VirtioFS Virtio device specification, supporting devices used during the operation of VMs, managing the file system operations of the VMs, and delegating these operations to a specific vhost-user device backend implementation.
- 3) A vhost-user backend implementation, namely OVFS. This is a crucial aspect that requires particular attention in this project. This backend is a file system daemon running on the host side, responsible for handling all file system operations from VMs to access the shared directory. `virtiofsd` offers a practical example of a vhost-user backend implementation, based on pure Rust, forwarding VMs' file system requests to the local file system on the host side.

2.2 How OVFS Interacts with VMs and Hypervisor

The Virtio specification defines device emulation and communication between VMs and the hypervisor. Among these, the virtio queue is a core component of the communication mechanism in the Virtio specification and a key mechanism for achieving efficient communication between VMs and the hypervisor. The virtio

queue is essentially a shared memory area called vring between VMs and the hypervisor, through which the guest sends and receives data to the host.

Simultaneously, the Virtio specification provides various forms of Virtio device models and data interaction support. The vhost-user backend implemented by OVFS achieves information transmission through the vhost-user protocol. The vhost-user protocol enables the sharing of virtio queues through communication over Unix domain sockets. Interaction with VMs and the hypervisor is accomplished by listening on the corresponding sockets provided by the hypervisor.

In terms of specific implementation, the vm-memory crate, virtio-queue crate and vhost-user-backend crate play crucial roles in managing the interaction between OVFS, VMs, and the hypervisor.

The vm-memory crate provides encapsulation of VMs memory and achieves decoupling of memory usage. Through the vm-memory crate, OVFS can access relevant memory without knowing the implementation details of the VMs memory. Two formats of virtio queues are defined in the Virtio specification: split virtio queue and packed virtio queue. The virtio-queue crate provides support for the split virtio queue. Through the DescriptorChain package provided by the virtio-queue crate, OVFS can parse the corresponding virtio queue structure from the original vring data. The vhost-user-backend crate provides a way to start and stop the file system demon, as well as encapsulation of vring access. OVFS implements the vhost-user backend service based on the framework provided by the vhost-user-backend crate and implements the event loop for the file system process to handle requests through this crate.

2.3 OVFS Design Philosophy

In this section, we will present the design philosophy of the OVFS project. The concepts introduced here will permeate throughout the entire design and implementation of OVFS, fully manifesting in other sections of the proposal.

Stateless Services

The mission of OVFS is to provide efficient and flexible data access methods for VMs using Virtio and VirtioFS technologies. Through a stateless services design, OVFS can easily facilitate large-scale deployment, expansion, restarts, and error recovery in a cluster environment running multiple VMs. This seamless integration into existing distributed cluster environments means that users do not need to perceive or maintain additional stateful services because of OVFS.

To achieve stateless services, OVFS refrains from persisting any metadata information. Instead, it maintains and synchronizes all state information of the OVFS file system during operation through the backend storage system. There are two implications here: OVFS doesn't need to retain additional operational status during runtime; it doesn't require the maintenance of additional file system metadata when retrieving data from the backend storage system. Consequently, OVFS doesn't necessitate exclusive access to the storage system. It permits any other application to read and write data to the storage system when it serves as the storage backend for OVFS. Furthermore, OVFS ensures that the usage semantics of data in the storage system remain unchanged. All data in the storage system is visible and interpretable to other external applications.

Under this design, OVFS alleviates concerns regarding synchronization overhead and potential consistency issues stemming from data alterations in the storage system due to external operations, thereby reducing the threshold and risks associated with OVFS usage.

Storage System As A Service

We aspire for OVFS to serve as a fundamental storage layer within a VM cluster. With OVFS's assistance, VMs can flexibly and conveniently execute data read and write operations through existing distributed storage system clusters. OVFS enables the creation of different mount points for various storage systems. This service design model is extremely scalable and supports VMs to access multiple

existing storage systems. By accessing different mount points, VMs can seamlessly access various storage services.

This design pattern allows users to customize the data access pipeline of VMs in distributed clusters according to their needs and standardizes the data reading, writing, and synchronization processes of VMs. In case of a network or internal error in a mounted storage system, it will not disrupt the normal operation of other storage systems under different mount points.

User-Friendly Interface

OVFS must offer users a user-friendly operating interface. This entails ensuring that OVFS is easy to configure, intuitive, and controllable in terms of behavior. OVFS accomplishes this through the following aspects:

- 1) It's essential to offer configurations for different storage systems that align with OpenDAL. For users familiar with OpenDAL, there's no additional learning curve.
- 2) OVFS is deployed using a formatted configuration file format. The operation and maintenance of OVFS only require a TOML file with clear content.
- 3) Offer clear documentation, including usage and deployment instructions, along with relevant scenario descriptions.

2.4 File System Operations Based On Various Storage Backend

OVFS implements a file system model based on OpenDAL. A file system model that provides POSIX semantics should include access to file data and metadata, maintenance of directory trees (hierarchical relationships between files), and additional POSIX interfaces.

Lazy Metadata Fetch In OVFS

OpenDAL natively supports various storage systems, including object storage, file storage, key-value storage, and more. However, not all storage systems directly offer an abstraction of file systems. Take AWS S3 as an example, which provides object storage services. It abstracts the concepts of buckets and objects, enabling users to create multiple buckets and multiple objects within each bucket. Representing this classic two-level relationship in object storage directly within the nested structure of a file system directory tree poses a challenge.

To enable OVFS to support various storage systems as file data storage backends, OVFS will offer different assumptions for constructing directory tree semantics for different types of storage systems to achieve file system semantics. This design approach allows OVFS to lazily obtain metadata information without the need to store and maintain additional metadata. Additional metadata not only leads to synchronization and consistency issues that are challenging to handle but also complicated OVFS's implementation of stateless services. Stateful services are difficult to maintain and expand, and they are not suitable for potential virtualization scenarios of OVFS.

File System Operations Based On Object Storage Backend

The working principle of OVFS based on the object storage backend is to implement the file and directory in the file system through a single bucket in the object storage. When using OVFS, users need to specify the bucket used to store data. A comprehensive directory tree architecture is realized by treating the object name in the bucket as a full path in the file system and treating the slash character "/" as a directory separator. File system operations in the VMs can interact with the object storage system through similar escape operations to achieve file system interface based data reading and writing.

The following table lists the mapping of some file system operations based on the object storage system. Here we refer to the mapping method of FUSE operations to object storage system operations in Google Cloud Storage FUSE.

File System Operations	Object Storage Backend Operations
-------------------------------	--

read all entries under the directory with the full path "/xxx/yyy"	Objects.get("/xxx/yyy/") Objects.list("/xxx/yyy/")
create a directory with the full path "/xxx/yyy"	Objects.get("/xxx/yyy") Objects.get("/xxx/yyy/") Objects.insert("/xxx/yyy/")
remove a directory with the full path "/xxx/yyy" and delete all entries under the directory	Objects.get("/xxx/yyy/") Objects.list("/xxx/yyy/") Objects.delete("/xxx/yyy/local.txt") Objects.delete("/xxx/yyy/")
create a file named "zzz" in a directory with the full path "/xxx/yyy"	Objects.get("/xxx/yyy/") Objects.get("/xxx/yyy/zzz") Objects.get("/xxx/yyy/zzz/") Objects.insert("/xxx/yyy/zzz")
remove a file named "zzz" in a directory with the full path "/xxx/yyy"	Objects.get("/xxx/yyy/") Objects.get("/xxx/yyy/zzz") Objects.delete("/xxx/yyy/zzz")

File System Operations Based On File Storage Backend

Unlike distributed object storage systems, distributed file systems already offer operational support for file system semantics. Therefore, OVFS based on a distributed file system doesn't require additional processing of file system requests and can achieve file system semantics simply by forwarding requests.

Limitations Under OVFS Metadata Management

While OVFS strives to implement a unified file system access interface for various storage system backends, users still need to be aware of its limitations and potential differences. OVFS supports a range of file system interfaces, but this doesn't imply POSIX standard compliance. OVFS cannot support some file system calls specified in the POSIX standard.

2.5 Multi Granular Object Size Cache Pool

In order to improve data read and write performance and avoid the significant overhead caused by repeated transmission of hot data between the storage system and the host, OVFS needs to build a data cache in the memory on the host side.

Cache Pool Based On Multi Linked List

OVFS will create a memory pool to cache file data during the file system read and write process. This huge memory pool is divided into object sizes of different granularities (such as 4 kb, 16 kb, 64 kb, etc.) to adapt to different sizes of data file data blocks.

Unused cache blocks of the same size in the memory pool are organized through a linked list. When a cache block needs to be allocated, the unused cache block can be obtained directly from the head of the linked list. When a cache block that is no longer used needs to be recycled, the cache block is added to the tail of the linked list. By using linked lists, not only can the algorithmic complexity of allocation and recycling be $O(1)$, but furthermore, lock-free concurrency can be achieved by using CAS operations.

Write Back Strategy

OVFS manages the data reading and writing process through the write back strategy. Specifically, when writing data, the data is first written to the cache, and the dirty data will be gradually synchronized to the backend storage system in an asynchronous manner. When reading the file data, the data will be requested from the backend storage system after a cache miss or expiration, and the new data will be updated to the cache, and its expiration time will be set.

OVFS will update the dirty data in the cache to the storage system in these cases:

- 1) When VMs called `fysnc`, `fdatasync`, or used related flags during data writing.

2) The cache pool is full, and dirty data needs to be written to make space in the cache. This is also known as cache eviction, and the eviction order can be maintained using the LRU algorithm.

3) Cleaned by threads that regularly clean dirty data or expired data.

DAX Window Support (Experimental)

The VirtioFS protocol extends the DAX window experimental features based on the FUSE protocol. This feature allows memory mapping of file contents to be supported in virtualization scenarios. The mapping is set up by issuing a FUSE request to OVFS, which then communicates with QEMU to establish the VMs memory map. VMs can delete mapping in a similar manner. The size of the DAX window can be configured based on available VM address space and memory mapping requirements.

By using the mmap and memfd mechanisms, OVFS can use the data in the cache to create an anonymous memory mapping area and share this memory mapping with VMs to implement the DAX Window. The best performance is achieved when the file contents are fully mapped, eliminating the need for file I/O communication with OVFS. It is possible to use a small DAX window, but this incurs more memory map setup/removal overhead.

2.6 Flexible Configuration Support

Running QEMU With OVFS

As described in the architecture, deploying OVFS involves three parts: a guest kernel with VirtioFS support, QEMU with VirtioFS support, and the VirtioFS daemon (OVFS). Here is an example of running QEMU with OVFS:

```
host# ovfs --config-file=./config.toml
```

```
host# qemu-system \
```

```
-blockdev file,node-name=hdd,filename=<image file> \  
-device virtio-blk,drive=hdd \  
-chardev socket,id=char0,path=/tmp/vfsd.sock \  
-device vhost-user-fs-pci,queue-size=1024,chardev=char0,tag=<fs tag> \  
-object memory-backend-memfd,id=mem,size=4G,share=on \  
-numa node,memdev=mem \  
-accel kvm -m 4G
```

```
guest# mount -t virtiofs <fs tag> <mount point>
```

The configurations above will generate two devices for the VMs in QEMU. The block device named hdd serves as the backend for the virtio-blk device within the VMs. It functions to store the VMs' disk image files and acts as the primary device within the VMs. Another character device named char0 is implemented as the backend for the vhost-user-fs-pci device using the VirtioFS protocol in the VMs. This character device is of socket type and is connected to the file system daemon in OVFS using the socket path to forward file system messages and requests to OVFS.

It is worth noting that the configuration method largely refers to the configuration in virtiofsd. For the purpose and planning of the project, we ignored many VMs configurations related file system access permissions or boundary handling methods.

Enable Different Distributed Storage Systems

In order for OVFS to utilize the extensive service support provided by OpenDAL, the corresponding service configuration file needs to be provided when running OVFS. The parameters in the configuration file are used to support access to the storage system, including data root address and permission authentication. Below is an example of a configuration file, using a toml format similar to oli (a command line tool based on OpenDAL):

```
[socket_settings]  
socket_path = "/tmp/vfsd.sock"
```

```
[cache_settings]  
enabled_cache = true  
enabled_cache_write_back = false  
enabled_cache_expiration = true  
cache_expiration_time = "60s"  
  
[storage_settings]  
type = "s3"  
bucket = "<bucket>"  
region = "us-east-1"  
endpoint = "https://s3.amazonaws.com"  
access_key_id = "<access_key_id>"  
secret_access_key = "<secret_access_key>"
```

OVFS can achieve hot reloading by monitoring changes in the configuration file. This approach allows OVFS to avoid restarting the entire service when modifying certain storage system access configurations and mounting conditions.

2.7 Potential Usage Scenarios

In this section, we list some potential OVFS usage scenarios and application areas through the detailed description of the OVFS project in the proposal. It's worth mentioning that as the project progresses, more application scenarios and areas of advantage may expand, leading to a deeper understanding of the positioning of the OVFS project.

- 1) Unified data management basic software within distributed clusters.
- 2) The OVFS project could prove highly beneficial for large-scale data analysis applications and machine learning training projects. It offers a means for applications within VM clusters to read and write data, models, checkpoints, and logs through common file system interfaces across various distributed storage systems.

2.8 Expected File System Interface Support

Finally, the table below lists the expected file system interface support to be provided by OVFS, along with the corresponding types of distributed storage systems used by OpenDAL.

Command	Object Storage	File Storage	Key-Value Storage
stat	Support	Support	Not Support
touch/rm	Support	Support	Not Support
mkdir/rmdir	Support	Support	Not Support
read/write	Support	Support	Not Support
truncate	Support	Support	Not Support
readdir	Support	Support	Not Support
rename	Support	Support	Not Support
flush/fsync	Support	Support	Not Support
ln	Not Support	Not Support	Not Support
getxattr/setxattr	Not Support	Not Support	Not Support
chmod/chown	Not Support	Not Support	Not Support
access	Not Support	Not Support	Not Support

It is worth mentioning that although storage systems are simply divided into three types here, the file system interface support provided by OVFS needs to be discussed in detail based on the specific service type. In order to avoid introducing too much complexity here, we will use S3 and local file systems as references for object storage systems and file system storage systems respectively, and aim to implement the planning in the above table based on these two typical storage systems.

Since the data volume of an individual file may be substantial, contradicting the design of key-value storage, we do not intend to include support for key-value Storage in this project. In addition, the complex permission system control of Linux is not within the scope of this project. Users can restrict file system access behavior based on the configuration of storage system access permissions in the OVFS configuration file.

3 Deliverables

This chapter describes the items that the OVFS project needs to deliver during the implementation cycle of GSoC 2024.

1) A code repository that implements the functions described in the project details. The services implemented by OVFS in the code repository need to meet the following requirements: (1) VirtioFS implementation, well integrated with VMs and QEMU, able to correctly handle VMs read and write requests to the file system. (2) Supports the use of distributed object storage systems and distributed file systems as storage backends, and provides complete and correct support for at least one specific storage service type for each storage system type. S3 can be used as the target for object storage systems, and local file systems can be used as the target for file systems. (3) Supports related configurations of various storage systems. Users can configure storage system access and use according to actual needs. When an error occurs, users can use the configuration file to restart services.

2) Form an OVFS related test suite. Testing about the project should consist of two parts: (1) Unit testing in code components. Unit testing is the guarantee that the code and related functions are implemented correctly. This test implementation accompanies the entire code implementation process. (2) CI testing based on github actions. The OpenDAL project integrates a large number of CI tests to ensure the correct behavior of OpenDAL under various storage backends. OVFS needs to use good CI testing to check potential errors during code submission.

3) A performance test report of OVFS. The report needs to perform basic metadata operations, data reading and writing performance tests on the VMs mounted with

OVFS, and summarize the performance of OVFS through the test results. Reports can be based on file system performance testing tools such as fio, sysbench and mdtest, and compared with virtiofsd when necessary.

4) Documentation on the introduction and use of OVFS, and promote the inclusion of OVFS documentation into the official OpenDAL documentation when the GSoC project is completed.

4 Reference

- <https://opendal.apache.org/>
- <https://github.com/apache/opendal>
- <https://virtio-fs.gitlab.io/>
- <https://gitlab.com/virtio-fs/virtiofsd>
- <https://crates.io/crates/vm-memory>
- <https://crates.io/crates/virtio-queue>
- <https://crates.io/crates/vhost-user-backend>
- <https://www.qemu.org/>
- <https://katacontainers.io/>
- <https://github.com/hpc/ior>
- <https://github.com/axboe/fio>
- <https://github.com/akopytov/sysbench>
- <https://cloud.google.com/storage/docs/gcs-fuse>
- <https://groups.oasis-open.org/communities/tc-community-home2?CommunityKey=b3f5efa5-0e12-4320-873b-018dc7d3f25c>