

# Evolving the Cloud Block Store with Performance, Elasticity, Availability, and Hardware Offloading

ERCI XU, Alibaba Group, Hangzhou, China  
WEIDONG ZHANG, Alibaba Group, Beijing, China  
QIUPING WANG, Alibaba Group, Hangzhou, China  
XIAOLU ZHANG, Alibaba Group, Hangzhou, China  
YUESHENG GU, Alibaba Group, Hangzhou, China  
ZHENWEI LU, Alibaba Group, Hangzhou, China  
TAO OUYANG, Alibaba Group, Hangzhou, China  
GUANQUN DONG, Alibaba Group, Hangzhou, China  
WENWEN PENG, Alibaba Group, Hangzhou, China  
ZHE XU, Alibaba Group, Hangzhou, China  
SHUO ZHANG, Alibaba Group, Hangzhou, China  
DONG WU, Alibaba Group, Hangzhou, China  
YILEI PENG, Alibaba Group, Hangzhou, China  
TIANYUN WANG, Alibaba Group, Hangzhou, China  
HAORAN ZHANG, Alibaba Group, Hangzhou, China

---

Authors' Contact Information: Erci Xu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: jostep90@gmail.com; Weidong Zhang (Corresponding author), Alibaba Group, Beijing, Beijing, China; e-mail: iszhangwd@hotmail.com; Qiuping Wang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: wangqiuping.wqp@alibaba-inc.com; Xiaolu Zhang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: xiaoluzhang.zxl@alibaba-inc.com; Yuesheng Gu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: yuesheng.gu@alibaba-inc.com; Zhenwei Lu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: zhenwei.lu@alibaba-inc.com; Tao Ouyang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: tao.oyt@alibaba-inc.com; Guanqun Dong, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: guanqun.dgq@alibaba-inc.com; Wenwen Peng, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: wenwen.pww@alibaba-inc.com; Zhe Xu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: xiuda.xz@alibaba-inc.com; Shuo Zhang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: zs157140@alibaba-inc.com; Dong Wu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: haoyu.wd@alibaba-inc.com; Yilei Peng, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: yilei.pyl@alibaba-inc.com; Tianyun Wang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: tianyun.wty@alibaba-inc.com; Haoran Zhang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: haoming.zhr@alibaba-inc.com; Jiesheng Wang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: gavin.wjs@alibaba-inc.com; Wenyan Yan, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: daniel.ywy@alibaba-inc.com; Yuanyuan Dong, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: yuanyuan.dyy@alibaba-inc.com; Wenhui Yao, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: wenhui.yao@alibaba-inc.com; Zhongjie Wu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: alanwu.wzj@alibaba-inc.com; Lingjun Zhu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: lingjun.zlj@alibaba-inc.com; Chao Shi, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: chao.shi@alibaba-inc.com; Yinhu Wang, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: xuanyou.wyh@alibaba-inc.com; Rong Liu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: todd.liur@alibaba-inc.com; Junping Wu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: junping.wu@alibaba-inc.com; Jiaji Zhu, Alibaba Cloud, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: jiaji.zhu@alibaba-inc.com; Jiesheng Wu, Alibaba Group, Hangzhou, Zhejiang, China; e-mail: jiesheng.wu@alibaba-inc.com. 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 the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from [permissions@acm.org](mailto:permissions@acm.org). © 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM 1553-3077/2025/02-ART10  
<https://doi.org/10.1145/3705925>

JIASHENG WANG, Alibaba Group, Hangzhou, China  
WENYUAN YAN, Alibaba Group, Hangzhou, China  
YUANYUAN DONG, Alibaba Group, Hangzhou, China  
WENHUI YAO, Alibaba Group, Hangzhou, China  
ZHONGJIE WU, Alibaba Group, Hangzhou, China  
LINGJUN ZHU, Alibaba Group, Hangzhou, China  
CHAO SHI, Alibaba Group, Hangzhou, China  
YINHU WANG, Alibaba Group, Hangzhou, China  
RONG LIU, Alibaba Group, Hangzhou, China  
JUNPING WU, Alibaba Group, Hangzhou, China  
JIAJI ZHU, Alibaba Cloud, Alibaba Group, Hangzhou, China  
JIESHENG WU, Alibaba Group, Hangzhou, China

---

In this paper, we qualitatively and quantitatively discuss the design choices, production experience, and lessons in building the Elastic Block Storage (EBS) at ALIBABA CLOUD over the past decade. To cope with hardware advancement and users' demands, we shift our focus from design simplicity in EBS1 to high performance and space efficiency in EBS2, and finally reducing network traffic amplification in EBS3. In addition to the architectural evolutions, we also summarize development lessons and experiences as four topics, including: (i) achieving high elasticity in latency, throughput, IOPS, and capacity; (ii) improving availability by minimizing the blast radius of individual, regional, and global failure events; (iii) identifying the motivations and key tradeoffs in various hardware offloading solutions; and (iv) identifying the pros/cons of alternative solutions and explaining why seemingly promising ideas would not work in practice.

CCS Concepts: • **Computer systems organization** → **Cloud computing**; **Reliability**; • **Software and its engineering** → **Ultra-large-scale systems**;

Additional Key Words and Phrases: Block storage, elasticity, availability, hardware offloading, cloud storage

#### ACM Reference Format:

Erci Xu, Weidong Zhang, Qiuping Wang, Xiaolu Zhang, Yuesheng Gu, Zhenwei Lu, Tao Ouyang, Guanqun Dong, Wenwen Peng, Zhe Xu, Shuo Zhang, Dong Wu, Yilei Peng, Tianyun Wang, Haoran Zhang, Jiasheng Wang, Wenyuan Yan, Yuanyuan Dong, Wenhui Yao, Zhongjie Wu, Lingjun Zhu, Chao Shi, Yinhu Wang, Rong Liu, Junping Wu, Jiaji Zhu, and Jiesheng Wu. 2025. Evolving the Cloud Block Store with Performance, Elasticity, Availability, and Hardware Offloading. *ACM Trans. Storage* 21, 2, Article 10 (February 2025), 29 pages. <https://doi.org/10.1145/3705925>

---

## 1 Introduction

**Elastic Block Storage (EBS)** service is a cornerstone in today's cloud [16, 18, 19]. In EBS, the storage service is in the form of virtual block devices with high performance, availability, and elasticity. The most outstanding characteristic of EBS architecture is the compute-to-storage disaggregation where the virtual machines (compute end) and disks (storage end) are not physically co-located but interconnected via datacenter networks.

In this paper, we start by revisiting the evolutions behind the three generations of EBS at ALIBABA CLOUD [16]. EBS1 marks our initial step in adopting the compute-to-storage philosophy. In EBS1, there are two notable design choices: in-place update from **virtual disks (VDs)** to physical disks, and the exclusive management of virtual disks. First, EBS1 directly maps a VD inside the **virtual machine (VM)** as a series of 64 MiB Ext4 files in the backend storage server. Moreover, EBS1 employs a fleet of stateless BlockServers to manage VDs where each VD is exclusively handled by

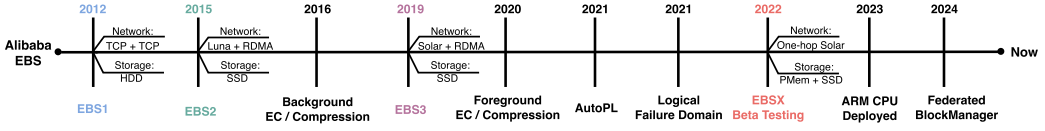


Fig. 1. Alibaba EBS timeline.

a BlockServer. While EBS1 had been successfully deployed on more than 300 HDD-backed clusters, its limitations also unfolded. The straightforward virtualization led to severe space amplification and performance bottlenecks.

We then developed EBS2 with two significant changes: the log-structured design, and VD segmentation. First, we employed the Pangu [33] distributed file system as our storage backend, and redesigned the BlockServers to convert VDs' all writes to sequential appends. By switching to a log-structured layout, EBS2 still used three-way replication for incoming writes but could transparently perform data compression and **erasure coding (EC)** in the background during **garbage collection (GC)**. Moreover, EBS2 split VDs into finer segments (32 GiB each), thus shifting the mapping between VDs and BlockServers from VD level to Segment level. With the above two changes, EBS2 was able to reduce the space efficiency from 3 (i.e., three-way replication) in EBS1 to 1.29 on average in the field. Moreover, supercharged with SSDs, an EBS2-backed VD can achieve up to 1 M IOPS and 4,000 MiB/s throughput with 100  $\mu$ s-level latency on average. Unfortunately, EBS2 also faced a significant challenge. That is, the traffic amplification factor increased to 4.69, namely 3 (foreground replication write) plus 1 (background GC read) and 0.69 (background EC/-compression write).

Hence, we built EBS3 to reduce traffic amplification using online (i.e., foreground) EC/compression via two techniques: **Fusion Write Engine (FWE)**, and FPGA-based hardware compression. FWE aggregates write requests from different segments (if necessary) to meet the size requirement of EC and compression. Moreover, EBS3 offloads the compute-intensive compression to a customized FPGA for acceleration. As a result, EBS3 can reduce the storage amplification factor from 1.29 to 0.77 (after compression) and the traffic amplification factor from 4.69 to 1.59 while still maintaining performance similar to EBS2. Since release, EBS3 has been deployed on more than 100 clusters, serving over 500K VDs.

Figure 1 outlines the chronological progression of Alibaba EBS since 2012. We highlight the time of major releases (i.e., EBS1 to EBS3), the integration of key techniques (e.g., Luna, our user-space TCP stack [44]) and the adoption of advanced hardware (e.g., Persistent Memory in EBSX). The evolution of EBS demonstrates a shift in focus from performance to space and traffic efficiency. Nevertheless, simply altering the high-level architecture is not enough. Next, we further discuss our lessons in building a high-performance and robust EBS from four perspectives, including elasticity, availability, hardware offloading, and alternative solutions.

One representative feature of a cloud block store is elasticity—providing VDs with varying performance and capacity levels. There are two key aspects: identifying boundaries and achieving fine-grained tuning. First, we discover that the average and tail latency are dominated by different factors—hardware overhead and software processing, respectively. Thus, we build corresponding solutions including EBSX, a one-hop architecture backed by persistent memory for minimizing average latency, and the use of dedicated threads for I/O to alleviate tail latency. Furthermore, we realize that throughput and IOPS are bounded by similar mechanisms. In the frontend Block-Client, we optimize the stack by moving the processing from kernel to user-space and then to hardware offloading in FPGA. In the backend BlockServer, we utilize high parallelism to achieve efficient throughput/IOPS control. As for space elasticity, EBS not only provides wide ranges of

storage space (i.e., from 1 GiB to 64 TiB) but also supports features including flexible resizing and fast cloning.

We then move on to discuss threats to the availability of EBS, especially under large-scale deployment. We begin by categorizing three levels of failure events, including individual, regional, and global, that can lead to one, several, or all VD in a cluster (temporarily) ceasing service. With VD segmentation and segment migration since EBS2, field diagnosis indicates that regional events become more frequent and an individual event can now easily cascade into regional or even global ones. Therefore, on the control plane, we developed a Federated BlockManager to organize the VDs into mini groups and use CentralManager for coordination. Additionally, on the data plane, we have built the logical failure domain to limit the destinations of segment migration.

In the third topic, we highlight the importance of offloading key control/data paths to hardware for acceleration. We use the offloading evolutions of both BlockClient and BlockServer as examples to discuss the tradeoffs between different options. Specifically, BlockClient started with FPGA offloading to accelerate storage/network virtualization. However, impacted by FPGA instability under large deployment (e.g., 22% of downtime caused by FPGA-related issues), our BlockClient dropped the FPGA-based approach and adopted the ASIC-based solution. Conversely, BlockServer, which also initially chose FPGA for speeding up EC/compression, opted for the many-core ARM CPU as the next step due to the flexibility and cost requirement.

We extensively evaluate the three generations of EBS designs by testing raw VD performance, running application-based benchmarking and comparing deployment statistics. The experiments verify our analysis of the pros/cons of each generation.

Finally, we lead a series of discussions on the lessons learned and potential future directions for EBS (Section 7). For example, we explain why seemingly promising ideas, such as extending EBS1 with segmentation but without a log-structured design, eventually failed, and discuss the possibilities of alternative solutions (e.g., building EBS with open-source software). We end this paper with a short discussion of related work and a conclusion.

## 2 Architecture Evolution: A Shift of Focus

### 2.1 EBS1: An Initial Foray

EBS1 marked our first step into offering an elastic block store based on a disaggregated architecture. By placing the compute and storage in different clusters and connecting them via the datacenter network, this design offers flexibility in deployment, scaling, and evolution. Such philosophy has been widely adopted by many vendors, such as Microsoft Azure [25] and Google datacenters [29].

**Compute end.** Figure 2 provides a high-level overview of EBS1. A compute cluster comprises multiple servers where each server runs one BlockClient and can host several VMs. In addition, a VM can mount one or more VDs. Users can access the VD as a normal block device and the host server forwards the I/O requests to the storage clusters via the BlockClient.

**Storage end.** EBS1 used a different fleet of dedicated servers for storage. First, we build the BlockManager (a set of three nodes backed by Paxos) and a group of BlockServers. The BlockManager maintains VDs' metadata, such as the capacity and snapshot versions. The BlockServers handle I/O requests of multiple VDs assigned by the BlockManager. Note that the BlockManager can reassign a VD to another BlockServer during failover. Second, we further introduce a data abstraction, called *chunk*. The **Logical Block Address (LBA)** space of a user's VD is divided into a series of 64 MiB chunks. Similarly, the ChunkManager (a set of three nodes backed by Paxos) manages the chunks' metadata. The ChunkServer stores a 64 MiB chunk as a 64 MiB Ext4 file (called DataFile) and performs in-place updates to the chunk for write requests. Each chunk is three-way replicated

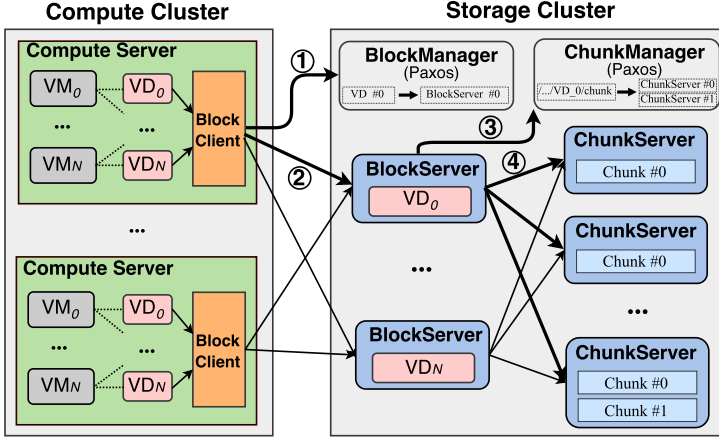


Fig. 2. The system architecture of EBS1 (Section 2.1). **VD**: Virtual Disk. **VM**: Virtual Machine. *BlockManagers and ChunkManagers all run three-instance Paxos groups. Each VM can host multiple VDs.*

on three different ChunkServers. For efficiency, we use thin provisioning—allocating space only when the user writes data to the VD.

**Network.** There are mainly two sets of network in EBS1. The frontend network connects the compute and storage clusters. The backend network inside the storage clusters connects the BlockServers to the ChunkServers. Both use Clos topology and rely on the 10 Gbps network with the kernel TCP/IP stack.

**Data-flow.** When a VM issues a new write request to its VD, BlockClient first contacts the BlockManager to locate the corresponding BlockServer for this request (① in Figure 2). Then, the BlockClient forwards the write request to the BlockServer (②). The BlockServer further asks the ChunkManager to determine the three ChunkServers (③) and persists the data accordingly (④). In practice, the BlockClient caches the VD-to-BlockServer mappings, and BlockServers cache chunk-to-ChunkServer mappings (i.e., skipping ① and ③).

**Limitations.** EBS1, released in 2012, has served over 1 million VDs and stored hundreds of PBs of data across hundreds of deployed clusters. Its straightforward and mature designs (e.g., in-place update and N-to-1 mapping between VDs and BlockServers)—while expediting development and deployment—limit the performance and efficiency. For example, to reduce the space overhead, we wish to use data compression and EC. However, compression non-deterministically alters the size of data which breaks the direct mapping of in-place updates. In addition, EC has a minimum size requirement (e.g., when the stripe unit is 4 KiB, EC(4,2) requires at least 16 KiB) and thus can result in significant write amplification, especially for small I/O requests. Another limitation is that, under the N-to-1 mapping, the performance of a VD is ultimately bounded by the performance of the corresponding BlockServer which can suffer from hotspot issues under burst workloads. In addition, with HDD and traditional kernel TCP/IP stack, we find it difficult to quantify and guarantee SLOs to users.

## 2.2 EBS2: Speedup with Space Efficiency

**Overview.** Figure 3 presents the high-level architecture of EBS2. The most significant change is that EBS2 no longer directly handles the data persistence or manages the consensus protocol. Instead, it builds on top of a distributed storage system—named Pangu [33]—which provides

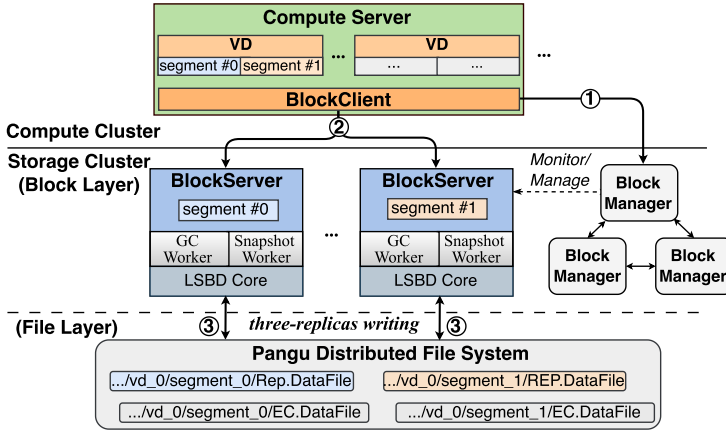


Fig. 3. The overview of EBS2. **LSBD**: Log-Structured Block Device. **REP.DataFile**: DataFile with three-way replication. **EC.DataFile**: DataFile with EC(8,3) encoding.

append-only file semantics and distributed lock services (based on a customized Raft protocol). The BlockServers employ a log-structured design [36] and translate the VD's write requests into Pangu append-only writes, thus enabling efficient data compression and EC during background garbage collection. We also split the VD address space into fixed-size segments, allowing one VD to be served by multiple BlockServers. Also, with segmentation, failover is no longer at the granularity of the whole VD but a segment—BlockManager migrates the impacted segment to another BlockServer. In addition, the BlockManager directly uses Pangu distributed lock service instead of Paxos for leader election.

As a result, we modified the I/O procedures as follows. After receiving a VD's request, BlockClient first retrieves the segment's address from BlockManager (① in Figure 3, which can be skipped by caching), and then forwards the I/O requests to the target BlockServer (②). BlockServer employs the **Log-Structured Block Device (LSBD)** Core to convert I/O requests into Pangu APIs and then calls an embedded Pangu client for persisting or fetching data (③). Note that, since EBS2, a BlockServer and a Pangu's ChunkServer, while co-located on the same physical server, are logically independent processes and rely on backend network for transferring data (i.e., not enforcing locality).

**Disk segmentation.** Figure 4 illustrates how EBS2 partitions the VD's LBA into several 128 GiB segment groups each of which further comprises four 32 GiB segments. BlockServers in EBS2 operate at the granularity of segments. Further, EBS2 organizes the segment group as a series of data sectors and allocates them to the segments in a round-robin fashion. Finally, EBS2 associates one segment with 8 - 16 DataFiles (512 MiB by default) to support concurrent writes. DataFile is essentially a Pangu file designed to persist a portion of a segment's data. These different levels of parallelism help EBS2 alleviate the hotspot accessing in VDs.

**Log-structured Block Device.** In EBS2, we developed an LSBD Core to support the append-only semantics of the underlying Pangu and thus split traffic into frontend (i.e., client I/Os) and backend (i.e., GC and compression). Figure 5 shows the frontend I/O flow including persisting users' data as a series of 4 KiB blocks and 64B metadata pairs (①), responding to users (②), recording updates in the transaction file (③), and updating the in-memory index map (④). The index map is essentially a **log-structured merge tree (LSM-tree)** to speed up the locating process by storing a mapping from VD's LBA to the corresponding DataFile ID, offset and length. The TxnFile accelerates the

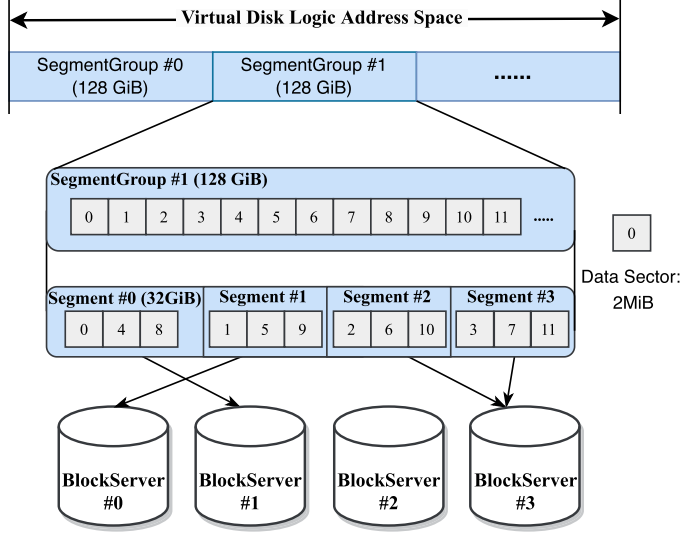
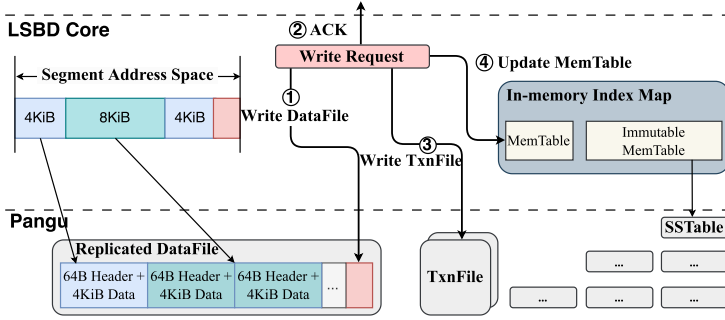


Fig. 4. The disk segmentation design of EBS2 (Section 2.2).

Fig. 5. The data organization and persistence format of LSBD. *TxnFile*: TransactionFile.

index map rebuild upon segment migrations. EBS can recover the in-memory index map by reading the latest I/Os' LBA-to-DataFile mappings from the TxnFile, without the need of tail scanning the data blocks in the DataFiles. Note that DataFile, TxnFile, and the SSTables in the LSM-tree are all Pangu files.

**GC with EC/Compression.** EBS2 runs GC at the granularity of *DataFile* (see Figure 6). When stale data within DataFiles reaches the threshold, EBS2 initiates performing GC by collecting valid data from the dirtiest DataFiles under the same segment and combining them as new DataFiles. EBS2 finishes GC by updating the TxnFile and the in-memory index map. During GC, we also convert the replicated “REP.DataFiles” to space-efficient “EC.DataFiles” with (8,3) EC and LZ4 compression.

Given that compression can alter the size of data, we structured the EC.DataFile into three main components: a DataFile Header, a series of CompressedBlocks, and an Offset Table. The Compressed DataFile Header includes a magic number (marking the start of the DataFile), version and checksum. Each CompressedBlock contains CompressionHeader (CmpHdr) and CompressionBody (CmpBdy). The CmpHdr records the timestamp, the compression algorithm (LZ4 by default), the size of CmpBdy, and the checksum. CmpBdy contains the compressed data and metadata (i.e., 4 KiB + 64B before compression). At the end of the Compressed DataFile,

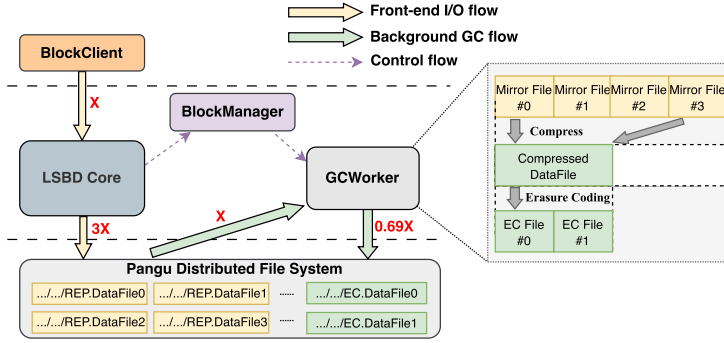


Fig. 6. The garbage collection in EBS2.

we enclose the mapping between the original LBA in VD and the location in the Compressed DataFile as OffsetTable. When reading the compressed data, EBS2 first locates data by querying OffsetTable, then reads and decompresses the data.

We leveraged the opportunity of GC to perform the transformation of erasure coding and compression. If needed, EBS2 can schedule special types of GC tasks that preferably select “Rep.DataFiles” (replicated, non-compressed data) over existing “EC.DataFiles” (erasure-coded, compressed data), in order to make up more storage space for incoming writes.

The garbage percentage thresholds used to trigger GC in production vary, depending on the cluster storage usage and the workload. We deployed a set of optimizations for improving GC efficiency (e.g., placement based on inferring the block invalidation time [37]). In production, for the most stressed clusters, the write amplification due to GC (i.e., the number of bytes written by BlockServer and GCWorker over the number of bytes written by BlockServer) is less than 1.5.

**BlockManager with higher availability.** The integration of Pangu enhances the availability of EBS2’s control plane. First, through the Pangu lock service, the BlockManager can continue serving clients even in the face of two out of three node failures. Second, EBS2 now stores the VDs’ metadata in a persistent and replicated key-value store as Pangu files, while EBS1 stores the VDs’ metadata in local disks where data loss could lead to an extended repair time.

**Network.** EBS2 uses a similar network setup as EBS1 except for two fundamental differences. First, for the frontend network, we replace the kernel TCP with our user-space TCP implementation (called Luna [44]) over a  $2 \times 25$  Gbps network. Luna achieves high performance (up to  $3.5\times$  throughput improvement and 53% latency reduction) by leveraging a run-to-completion thread model and a zero-copy memory model. Second, for the backend network, we use a  $2 \times 25$  Gbps RDMA network to meet the demanding SLAs [28]. Note that the two changes above only affect the data path. For the control path, we still use the kernel TCP (e.g., RPCs between BlockManagers and BlockServers).

**Snapshot.** The architectural changes in EBS2 also facilitate creating snapshots for VDs. With the out-of-place update, creating a snapshot in EBS2 no longer blocks foreground traffic. Instead, when receiving a snapshot request, BlockManager simply records a timestamp and asks the snapshot workers in BlockServers to upload the updates between the last snapshot timestamp and the latest one. As a result, generating a snapshot for 20GB of new data only takes around 30 seconds in EBS2, much less than the average 33 minutes in EBS1.

Figure 7 shows the process of taking snapshots and cloning disks in EBS2. First, SnapshotWorker enables the sharing of DataFile and IndexFile between the Snapshot Copy and virtual disk via *Hard Link*. Then, Snapshot Copy is divided into fixed-size data blocks before being uploaded offline to

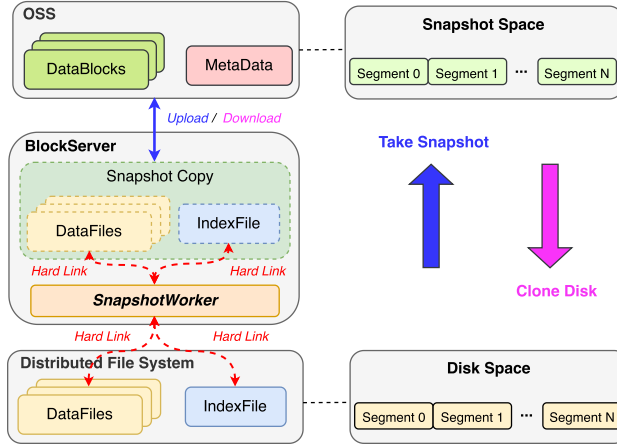


Fig. 7. The process of taking snapshots and cloning disks in EBS2.0.

the **object storage service (OSS)** for backup. For cloning disks, only one copy of the Snapshot data is downloaded and created as a Snapshot Copy within the same storage cluster. All virtual disks share the same Snapshot Copy via *Hard Link*. This not only reduces the traffic waste caused by downloading snapshot data when cloning multiple disks, but also supports the creation of numerous disks with one snapshot in a short time.

**Other background I/O.** Apart from GC, one important background task is data scrubbing, which performs periodical scanning to detect anomalies such as disk corruption and CPU silent data errors. To minimize the performance impacts, we have separated the scrubbing traffic from the GC tasks, capped the scrubbing traffic at 10 MiB/s (i.e., scanning all DataFiles every 15 days), and leveraged the heartbeat mechanism for monitoring the scrubbing progress.

**Deployment.** We released EBS2 in 2015 and subsequently scaled to more than 500 clusters for 2 million VDs. EBS2 could provide a virtual disk with an average write latency of  $100\ \mu\text{s}$  ( $12\times$  reduction than EBS1), a maximum IOPS of 1 M ( $50\times$  increase), and a maximum throughput of 4,000 MiB/s ( $13\times$  increase) for the guest OS. In the field, the compression ratio of the LZ4 algorithm is between 43.3% ~ 54.7%, and the average compression ratio is 50.1%. With Compress-EC in GC, EBS2 can reduce space usage from 3 to 0.69 replicas. Since un-GCed DataFiles are still stored with three-way replication, the average number of replicas in EBS2 is 1.29 in the field on average. For better management, we also reduced the cluster size from 700 servers in EBS1 to around 100 in EBS2.

**Limitations.** EBS2 successfully improves the space efficiency but incurred heavy traffic amplification. Compared with EBS1, EBS2 increases the overall traffic from 3 (i.e., from three-way replication) to 4.69 (i.e., 3 from the frontend plus 1.69 from the backend GC), yielding only 15.5% of the network bandwidth for serving the VDs' requests. To alleviate this issue, one promising solution is to adopt online EC/Compression, which means to directly store users' data in erasure-coded and compressed format.

The challenges are twofold. First, erasure coding requires the raw data blocks to be at least 16 KiB to achieve high compression and encoding efficiency. However, in the field, nearly 70% of write requests are smaller than 16 KiB. Moreover, EBS aims to deliver  $100\ \mu\text{s}$  write latency, and accumulating enough data in such a short interval can be difficult. For example, in order to perform compression and erasure coding for all user writes (i.e., accumulating 16 KiB data blocks within

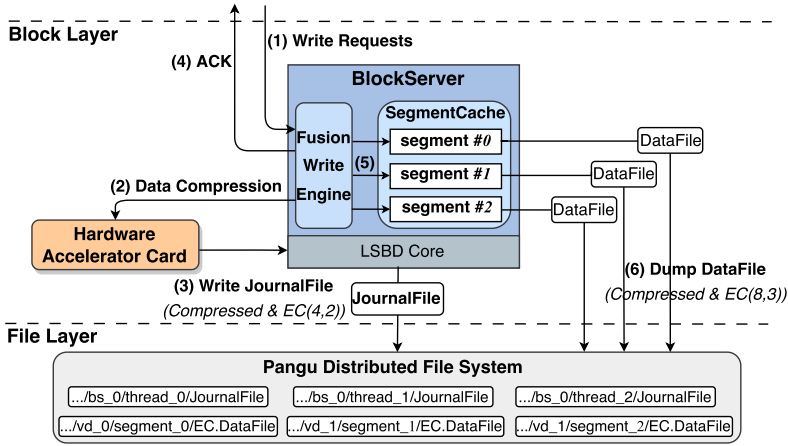


Fig. 8. The architecture and I/O flow of EBS3.

each 100  $\mu$ s interval), a segment needs to have a write throughput over 160 MiB/s—surpassing 90% of segments in production. Simply padding zeroes can result in an even higher traffic/space amplification. Second, even with the latency-optimized LZ4 algorithm, compressing a 16 KiB-sized data block still requires 25  $\mu$ s for CPUs, and such overhead escalates significantly for larger ones, rendering an unacceptable performance penalty for our service.

### 2.3 EBS3: Foreground EC/Compression

**Overview.** EBS3 achieves online EC/Compression by utilizing a **Fusion Write Engine (FWE)** to merge small writes and adopt an FPGA to offload the compression computations. Specifically, EBS3 first leverages FWE to accumulate writes from different segments of different VDs (i.e., step ① in Figure 8). FWE then combines these incoming writes as DataBlocks and sends them to the FPGA-based accelerator for data compression (②). EBS3 then calls Pangu to persist the compressed DataBlocks as JournalFiles with EC(4,2), namely ③. After persisting JournalFiles, EBS3 sends acks to the VD indicating the I/O completion (step ④). Then EBS3 copies the uncompressed data and preserves them within the BlockServer’s memory by segment (i.e., SegmentCaches, step ⑤). When the data in a SegmentCache reaches the threshold (512 KiB by default), EBS3 compresses the data via the host CPU, appends them to the DataFile with EC(8,3), and updates the TxnFile and in-memory index map (step ⑥).

For read requests, EBS3 first queries the SegmentCaches in the BlockServers as they have the latest data. If not found, EBS3 would further read the in-memory index map (i.e., the LSM-tree). Note that JournalFiles are EC-ed with compressed data from various segments of different VDs. Directly fetching data from JournalFiles can result in severe read amplification (i.e., decompressing with heavy scanning). Therefore, JournalFile is write-only during runtime and only readable during failover to recover yet-to-be-dumped data upon crash.

**Fusion Write Engine.** Usually, when receiving a batch of small write requests, FWE waits until the total amount of data reaches a threshold (16 KiB by default) and then forms a DataBlock before sending it to the FPGA for compression. We set the waiting timeout as 8  $\mu$ s (i.e., the interval between NIC pollings). Moreover, we discover that insisting on merging smaller writes (e.g., 4 KiB) can result in higher 99th tail latency (i.e., 220%). Therefore, for clusters that have infrequent small writes (e.g., certain clusters, on average, with only 3.72% of the workload are 4 KiB writes), we do not aggregate 4 KiB writes but directly append them with the traditional three-way replication.

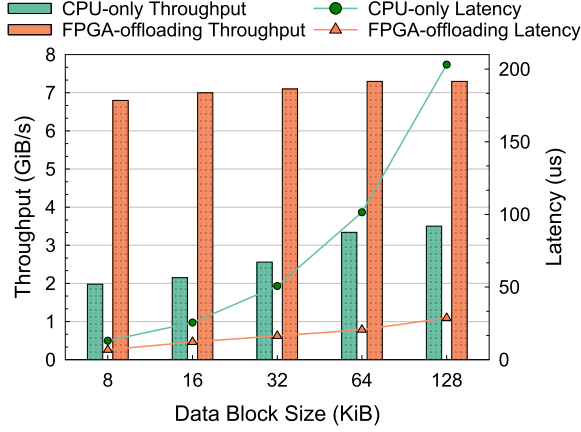


Fig. 9. The compression performance comparison of FPGA-offloading and CPU-only with 8 cores compression based on Silesia Compression Corpus.

**FPGA-based compression offloading.** We employ a customized FPGA to accelerate (de)compression, which includes a submission queue to buffer newly-formed DataBlocks and a completion queue to poll the results of hardware compression. We implement a scheduler inside the FPGA to split the DataBlocks as fixed-sized (e.g., 4 KiB) slices and employ multiple execution units to perform (de)compression tasks on these slices in parallel. To ensure data integrity, we place an end-to-end CRC check within the driver. After FPGA returns compressed data, EBS3 would immediately decompress the data and verify data integrity via CRC checking. Note this is a necessary overhead as, during failover, JournalFiles are the only data source.

Figure 9 shows the latency and maximum throughput of FPGA-offloading and CPU-only compression across different data block sizes, based on Silesia Compression Corpus [26]. The latency distribution of FPGA-offloading ranges from 29~65  $\mu$ s. Notably, when the data block size is 16 KiB, the latency of FPGA-offloading reduces by 78% compared to CPU-only. Further, FPGA-offloading achieves a maximum throughput of 7.3 GiB/s, whereas CPU-only compression is only 3.5 GiB/s. As data block size increases, the FPGA-offloading solution leads to larger performance gains. As for CPU savings, FPGA-offloading can save 8 CPU cores.

**Network.** EBS3 adopts higher linkspeed (i.e., 2×100 Gbps) network for both frontend and backend. In addition, we further developed Solar [34], a UDP-based transmission protocol. By leveraging the hardware offloading on our **Data Processing Units (DPUs)**, Solar can pack each storage data block as a network packet, thereby achieving CPU/PCIe bypassing, easy receive-side buffer management, and fast multi-path recovery.

**Deployment.** EBS3 has been deployed in over 100 storage clusters, serving more than 500K virtual disks since released in 2019. EBS3 offers comparable performance to EBS2. The incorporation of foreground EC/Compression in EBS3 enables all data to be stored immediately with high storage efficiency except in a few corner cases. As a result, the space efficiency (i.e., replica per data) in the field further drops from 1.29 in EBS2 to 0.77 in EBS3. In addition, the FPGA-based compression offloading can achieve 7.3 GiB/s throughput per card and the overhead ranges from 29~65  $\mu$ s. The overall traffic amplification drops from 4.69 in EBS2 to 1.59 in EBS3 (based on field statistics and numbers may slightly vary due to compression ratio differences across workloads).

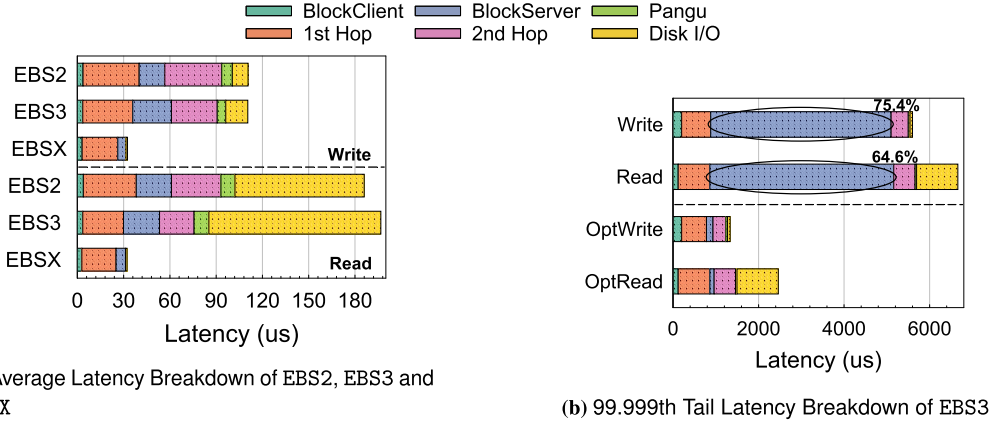


Fig. 10. 8 KiB-Sized Avg. and Tail Latency Breakdown of EBS. **1st hop**: network latency from compute to storage end. **2nd hop**: network latency from BlockServer to Pangu.

### 3 Elasticity: A Tale of Four Metrics

The capabilities (e.g., capacity and throughput) of a common block device (e.g., HDD and SSD) are usually bounded by the physical properties, such as encapsulation or interface. Backed by the cloud, EBS can provide VDs with much higher flexibility. In this section, we will share our experience of obtaining high elasticity, including pushing the upper bounds and achieving fine granularity.

#### 3.1 Latency

The latency of a VD is determined by the architecture, namely the path a request has traveled. For example, the latency of an EBS2-backed VD is bounded by the latency of the two-hop network (from BlockClient to BlockServer and then to ChunkServer), the software stack processing (i.e., BlockClient, BlockServer and Pangu) and the SSD I/O. Hence, the elasticity of latency is inherently coarse-grained, namely the different levels of time overhead under various architectures (e.g., EBS2 and EBS3). Next, from the perspectives of average and tail latency, we further analyze the status quo.

**Average latency.** In Figure 10(a), we measure the 8 KiB random read/write average latency breakdowns across different generations of EBS in their corresponding top 10% busiest production clusters. We choose to not include EBS1 in the comparison as it is no longer deployed and many of its hardware (e.g., HDD and 10 Gbps network) are obsolete. From the comparison, we first observe that the hardware processing—including 1st/2nd hop network (marked as orange and pink) and disk I/O (yellow)—accounts for the majority of the total latency in both EBS2 and EBS3. In addition, while EBS3 requires more time to process the data due to the frontend EC/compression, the reduced data volume in return spends less time traveling the network (i.e., lower 2nd hop latency in EBS3), yielding similar overall latency between EBS2 and EBS3. Third, the major difference between read and write lies in the disk I/O latency. Note that EBS2 is backed by TLC-based SSDs while EBS3 is backed by QLC-based SSDs.

Clearly, the key to improving average latency is reducing the hardware processing overhead. Therefore, we built EBSX, targeting latency-sensitive scenarios. EBSX installs the persistent memory (PMem) inside the BlockServers and directly stores the data in PMem with three-way replication. Compared to EBS2 and EBS3, EBSX skips the 2nd hop and drastically speeds up the disk I/O with PMem. Figure 10(a) shows that EBSX achieves 30  $\mu$ s-level latency on both read and write. Note

that, for space efficiency, data in PMem would be eventually flushed to Pangu and read statistics in Figure 10(a) is performance under cache hit (i.e., data in PMem).

**Tail Latency.** A user request may not be always served in time due to hardware failures [41, 42], misconfigurations [28, 34], software bugs [23, 31], or simply resource contentions [20]. Figure 10(b) presents the breakdowns of 99.999th percentile latency, a common threshold for defining the tail latency [27], among EBS3 clusters. We have collected millions of slow requests and calculated the average latency of each procedure. We observe that the BlockServer processing, such as non-IO RPC destruction in the IO thread, background periodic scrubbing and index compaction, accounts for the majority (75.4% for write and 64.6% for read) of the tail latency.

This observation may sound rather counter-intuitive as the hardware-related issues are usually blamed as the culprits [40, 43]. However, in EBS2 and EBS3, we have already incorporated a series of simple techniques to improve the quality of service of hardware processing. For example, network multi-path transport allows user requests to automatically shift traffic to other paths to avoid slow IO when suffering network abnormalities [34]. As another example, we employ a backup read/write strategy to trigger a retry to another location if the I/O takes longer than 99.9th percentile latency.

Our analysis indicates that the principal driver of the tail latency is the contention between the IO and the background tasks (e.g., segment status statistics and index compaction) in the BlockServers. In EBS2 and EBS3, the IO and the background tasks are executed on the same thread, leading to the IO hang-ups when the background tasks are triggered. To address this, we segregate the IO flow from other tasks and execute it on independent threads. With these enhancements, the 99.999th percentile write latency of EBS3 has been reduced to 1 ms, and the read latency reduced to 2.5 ms (i.e., OptWrite/Read in Figure 10(b)).

**Summary.** First, the elasticity of latency is coarse-grained—defined by the architectures, along with the hardware used (e.g., from EBS2 to EBSX). Second, optimizing hardware-induced latency is often straightforward. One can shorten the path (e.g., skip a network hop), use faster devices (e.g., PMem), or simply offset the risks with multi-path or retries. Third, tail latency by software stack has not received enough attention and may be regarded as noise. Our analysis suggests that under the high-speed network and fast SSDs, software-induced tail latency can be the dominant factor.

### 3.2 Throughput and IOPS

In the context of EBS, we often discuss the elasticity of throughput and IOPS together because the two metrics are often constrained by the same set of mechanisms. EBS achieves high elasticity in throughput/IOPS by optimizing two components on the key data path, BlockClient and BlockServer.

**BlockClient.** Every IO issued from the VD is first touched by the BlockClient. Therefore, the throughput and IOPS are bounded by the BlockClient's processing and forwarding capability. In EBS1, the BlockClient is implemented as a kernel module, and all IO requests are processed by the CPU. In EBS2, we move the IO processing to the user space by introducing a user-space TCP stack to handle the IO requests [44]. In EBS3, we further offload the IO processing to the hardware where a general-purpose FPGA, completely bypassing CPU, performs direct data move from VMs, data block CRC calculation, and packet transmission [34].

In Figure 11, we measure the maximum throughput and IOPS of BlockClient under different optimizations. We observe that EBS2 with the  $2 \times 25$  Gbps network, throughput is constrained by network capabilities. For EBS3 with the 2x100G network, the bottleneck shifts to the PCIe

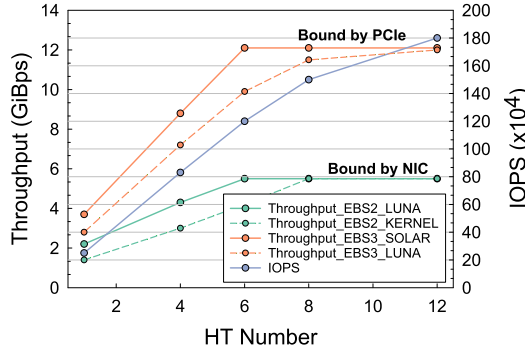


Fig. 11. The maximum throughput and IOPS changes of BlockClient with different HT numbers.

bandwidth. Furthermore, as long as network bandwidth is available, IOPS increases with the number of **hyper-threads (HTs)**.

**BlockServer.** Unlike BlockClient, once the requests reach the BlockServer, the throughput and IOPS are constrained by the levels of parallelism. Obviously, the more a request can be divided and served in parallel, the higher the throughput and the IOPS. Since EBS2, we have introduced three levels (i.e., SegmentGroup, Segment, and Data Sector) of parallelism to enhance virtual disk performance.

Recall that the Data Sector size (initially 2 MiB) is configurable in the segmentation design. Reducing the Data Sector size allows virtual disks to scale one SegmentGroup across more BlockServers, thereby obtaining higher throughput/IOPS. In the field, we further decrease Data Sector size to 128 KiB and EBS2 (and EBS3) are able to deliver 1,000 IOPS for every GiB subscribed. Note that configuring an even smaller Data Sector size may backfire as the write requests can be too fragmented, thereby leading excessive number of sub-I/Os even for small writes and placing prohibitively high pressure on the first-hop network.

**Base+Burst allocation.** With high throughput/IOPS enabled by BlockClient and BlockServer optimizations, efficiently allocating throughput/IOPS to VDs is the next step. Unlike the coarse-grained elasticity in latency, users can subscribe throughput and IOPS of VD on demand without altering the capacity, which is called auto performance level (AutoPL). However, we discover that the throughput/IOPS in practice is often over-provisioned by users to handle the sporadic workload bursts. For better resource efficiency, we have proposed the *Base+Burst* strategy based on the following techniques.

- *Priority-based congestion control.* We categorize IOs into baseIO and burstIO. BaseIO is pre-defined during virtual disk creation, while burstIO is allocated based on the available capability (i.e., not guaranteed). When a BlockServer is unable to meet all IO demands, it prioritizes processing the baseIO to ensure consistent latency. Currently, the maximum baseIO capacity of a VD is 50,000 IOPS, and the maximum burstIO capacity is 1 million IOPS.
- *Server-wise dynamic resource allocation.* Burst workloads can also place a heavy burden on BlockServer processing ability. Therefore, since EBS2, we devise the dynamic resource allocation to allow BlockServer to preempt resources (bandwidth, CPU cores, and memory) from background tasks (e.g., GCWorker) to handle workload spikes.
- *Cluster-wise hot-spot mitigation.* To ensure enough headroom for burstIO, especially under concurrent bursts onto the same BlockServer, we use cluster-wise load-balancing to remove hot spots. Under such scenarios, BlockManager would aggressively check the traffic status and migrate segments more frequently between BlockServers.

**Summary.** The first lesson here is that the upper bound of a VD's throughput/IOPS is determined by the client (i.e., processing/forwarding ability) as the backend can easily scale with parallelism. In addition, the high throughput/IOPS is often desired but not always needed. Therefore, using a *Base+Burst* strategy to cope with workload spikes can be more economically beneficial for both users and vendors.

### 3.3 Capacity

Achieving elasticity in capacity is a fundamental requirement of a cloud block store service. In EBS, we have further included the following features.

- *Flexible space resizing.* The segmentation design enables EBS with seamless support for VD resizing (i.e., adding or removing SegmentGroups). EBS currently supports virtual disk sizes ranging from 1 GiB to 64 TiB.
- *Fast VD cloning.* One outstanding characteristic of serverless applications is a large volume of resources (e.g., VDs) needs to be allocated in a short time. To support this, EBS uses the *Hard Link* of Pangu files, which allows the cloning of multiple disks within a storage cluster via downloading a single snapshot. As a result, EBS2 enables the creation of up to 10,000 virtual disks (each 40 GiB) in 1 minute.

## 4 Availability: The Dark Side of Scaling

Availability has always been a priority of cloud services. In EBS, we especially focus on the *blast radius*—defined as the number of VDs experiencing unavailable services upon failures. Here, we categorize the blast radius as follows.

- *Global.* In the face of such an event, the service availability of an entire cluster is impacted. A simple example is an abnormally operating BlockManager causing the entire cluster to perform in an undesired fashion (e.g., a misconfiguration causing network congestion and subsequent retry storms). Note that EBS service runs on a per-cluster basis and we do not discuss datacenter-level failures here.
- *Regional.* For a regional event, we define it as a failure that incurs the component(s) to deny service for several VDs. For example, when a BlockServer crashes, the hosted VDs would experience an outage until the corresponding segments are migrated or all incoming I/Os are forwarded.
- *Individual.* When an individual event occurs, only one VD is influenced. Representative examples include an uncorrectable error inside the disk (and subsequently a read retry) and a software bug that leads to an unsuccessful and redirected write.

A straightforward solution to minimize the blast radius is setting smaller clusters. In EBS2 and EBS3, we have reduced the cluster size from 700 nodes (in EBS1) to around 100. The benefit is obvious since there are much fewer VDs influenced by a global event now. However, this approach, while straightforward and effective, would not alleviate regional and individual failure events.

Meanwhile, the *regional* events are likely to be more severe due to two trends. First, EBS2 introduces the segmentation to split one VD into 32 GiB segments and hence a VD in EBS2 (or EBS3) is supported by multiple BlockServers instead of one in EBS1. Moreover, the average capacity of VDs only slightly increases from EBS1 to EBS2 and EBS3 (e.g., 197 GiB to 220 GiB). Therefore, we can conclude that in EBS2 and EBS3, a BlockServer hosts much more VDs than EBS1. Consequently, when a BlockServer crashes, more VDs are going to be influenced.

Moreover, *individual* events can cascade into *regional* failures as the segments can be migrated since EBS2. For example, an internal incident occurred as tens of BlockServers in an EBS2 cluster kept crashing and rebooting, degrading the total cluster capacity and the I/O quality of thousands

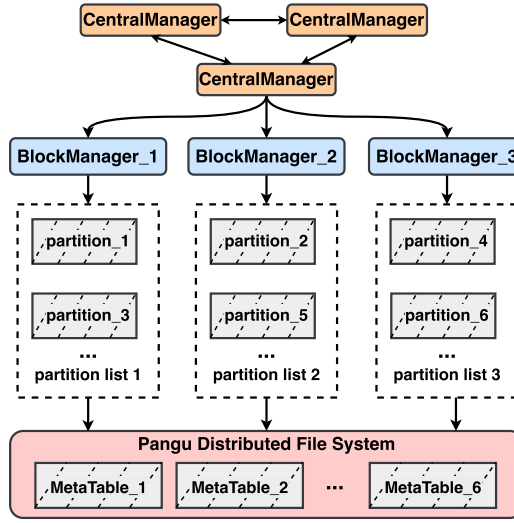


Fig. 12. The architecture of Federated BlockManager.

of VDs. In the beginning, a faulty segment crashes its BlockServer because of a buggy code logic—an individual event. Then, the control plane tries to migrate the segment to other BlockServers. However, as the client keeps retrying the failed requests, every BlockServer that loads the segment crashes as well, turning the individual event into a regional failure. This failure can easily grow to a cluster-wise outage in a short period if not manually intervened. Note that the cascading failures are not unique to EBS (e.g., cases in HBase [4–6]).

To adapt to the trends in regional and individual events, we further come up with techniques in both the control and data plane to improve the availability in EBS.

#### 4.1 Control Plane: Federated BlockManager

In each cluster, the control plane of EBS2 (referred to as BlockManager in Section 2) initially consists of a group of three nodes that leverages the distributed lock service provided by Pangu for leader elections. The leader node in the group serves all the control plane requests to the corresponding cluster and persists any state changes related to VDs to a single metadata table stored in Pangu.

This setup presents two challenges. First, the single leader serves all the VDs in the cluster with one single server. As the VD’s density grows, the chances and scale of its service disruption get higher. Second, the single metadata table hosts the metadata of VDs in the cluster. Once a part of the metadata table becomes corrupted, the BlockManager may not be able to load the metadata in memory, and cannot serve the VDs until the metadata table is repaired. In production, we have seen an increasing VD-per-cluster density over the past several years, urging us to solve both issues to provide high availability. Specifically, since the initial attempt to deploy a 100-node cluster in EBS2, the average number of VDs has increased from 20,000 to 100,000 per cluster. Correspondingly, the CPU and the memory consumption have increased by 4.1× and 4.5×, respectively.

Figure 12 illustrates the architecture of Federated BlockManager—our solution to the availability issue in the control plane. Each cluster now has multiple BlockManagers, with a CentralManager dedicated to managing the BlockManagers. Each BlockManager further manages hundreds of VD-level *partitions*, each of which corresponds to a metadata table that only stores the metadata of a small subset of VDs. The mappings from VDs to partitions are static; given a VD, we use hashing

algorithms to compute its corresponding partition. The Client is not aware of the partition concept. For a given VD, it can query all the BlockManagers in the cluster to find out the BlockManager in charge, then send all its control plane requests to the BlockManager.

Note that instead of having three nodes (one leader and two standby), we now have only one node in each BlockManager. Upon the failure of a BlockManager, CentralManager redistributes its partitions to other BlockManagers, which then load the metadata tables of the partitions from Pangu into memory, and start to serve the VDs. Since the number of VDs in a partition is relatively small, the loading time is only several hundred milliseconds, without the need to have standby nodes continuously fetching the latest metadata updates for a fast leader switch. To ensure the availability of partition scheduling, the CentralManager consists of three nodes based on the lock service of Pangu.

Having multiple BlockManagers effectively reduces the blast radius of single leader service disruptions, since now each BlockManager only processes the requests of the subset of VDs in the partitions it manages. For the single table issue, instead of creating more BlockManagers, we adopt the partition design for two reasons. First, with partitions, we can make the number of VDs in a single metadata table small. For 100,000 VDs in a cluster, we need to have 1,000 tables to make the blast radius of metadata table failures smaller than 100, while it is not resource-efficient to create 1,000 BlockManagers in 100-node clusters. Second, the concept of multiple BlockManagers is mainly for distributing the workloads to multiple servers, so as to reduce the chances of service disruptions due to CPU and memory resource limits. Note that the CentralManager only manages BlockManager registrations and partition scheduling, and thus is less relevant to system availability.

Similar designs for blast radius reduction in the control plane of distributed storage systems can be found in the industry. HDFS Federation [1] is similar to Federated BlockManager, while it does not consider the metadata table failure mode. Each set of NameNode in HDFS Federation persists the metadata to its local disk, and all the requests to a set of NameNodes become unavailable when the data in the local disk is corrupted. AWS Physalia [24] deploys small units called cells, each of which consists of seven nodes deploying Paxos algorithms and serves a group of VDs. Differently, Federated BlockManager adopts a two-level VD management scheme, with each level emphasizing the blast radius of a single node and single table failures, respectively.

#### 4.2 Data Plane: Logical Failure Domain

The data plane of EBS2 constitutes multiple BlockServers, each of which hosts thousands of segments and handles I/O requests of the segments. When a BlockServer crashes, the control plane migrates the segments to other BlockServers in the cluster, such that the I/O services can resume in other BlockServers. However, this mechanism indicates that failures can be cascading among BlockServers. Specifically, if the crash is caused by an error segment (e.g., recall the buggy code case in our production earlier), after the migration, the BlockServer shall resume the requests and crash again. Ironically, optimizing segment scheduling for faster recovery in this case can make more BlockServer crashes, even possibly leading to cluster-wide outages.

We have several observations based on our deployment experiences. First, the failure typically originates from requests of a single VD or segment, and thus we need to monitor the status of segments instead of BlockServers. Second, the root causes of the failures are mostly due to software errors (e.g., bugs or misconfigurations) as hardware or mechanical failures are unlikely to travel along with the migrated segments. Software-induced failures can be time-consuming to pinpoint the culprit, let alone build automatic tools for recovering. Instead, a more practical way is to proactively reduce the impacts. Third, the cascading failures, if not intervened, can propagate quickly to the whole cluster. As a distributed service, it is acceptable to experience a few BlockServers shutdowns but not cluster-wide outage.

Based on these observations, we design the logical failure domain. The core idea is to isolate any suspicious segments by grouping them into a small set of BlockServers, to avoid other BlockServers or even the entire cluster being impacted. We first associate each segment with a unique token bucket with a maximum capacity of three. Each migration to a new BlockServer consumes one token, and the token is refilled every 30 minutes. When the token bucket is empty, any subsequent migrations of the segment can only be selected among the three pre-designated BlockServers, referred to as its logical failure domain. The token bucket design does not limit the number of migrations but the range of migrations. When the segment is successfully loaded in a BlockServer and can readily serve I/O requests for several minutes, we lift the failure domain constraints.

The segment-level failure domain can effectively isolate an error segment. However, if there are multiple error segments (e.g., from one VD) or even VDs, the segment-level failure domain is not sufficient to prevent multiple cascading failures from happening at the same time. Our solution is to merge the failure domains into one. Specifically, when there exists more than one segment forming a failure domain, we pick the first failure domain as the global failure domain, so as to ensure there are at most three BlockServers isolated for the cascading failures, without degrading the cluster capacity. Any subsequent segment with an empty token bucket will share the same global failure domain. Recall that we associate each segment with three tokens. As a result, the chances of a normal segment accidentally sharing the same migration path with an error segment is small (note that a cluster is of 100-node size), which effectively reduces false negatives. As such, we can efficiently identify any error segments with small impacts on other VDs.

After deploying the logical failure domain, we have successfully defended our system against several potential outages due to migration-induced cascading failures. Evidence shows that some open-source storage systems also suffer from the same issue from time to time. For example, the HBase community reports several bugs [4–6] that show similar symptoms and lead to system outages. However, their treatments focus on solving the bugs and reducing the blast radius with physical isolation (similar to our first attempts), while the logical failure domain prevents the cascading failure from causing a cluster-wise outage once and for all.

### 4.3 Lessons Learned

First, with denser SSDs (e.g., QLC NANDs) becoming readily available and the boosting processing ability of CPU or other customized hardware (e.g., DPU), one can expect that a crashed node in a distributed storage system—not just EBS—can impact increasingly more users. As a result, regional failure events would be more frequent and/or severe. The key benefit of Federated Managers in this case is that it enjoys a smaller blast radius without losing the flexibility of a large-scale cluster.

Second, owning a forwarding layer between the users and the underlying service is popular among distributed services, for example, distributed key-value store [2] and cloud storage services [21, 25]. However, this also means the request or certain data structures can be redirected to other destinations upon failures, leading an individual event (e.g., bug or misconfiguration) to become regional or even global. We do not aim at proactively recovering or avoiding these failures because such events, like bugs or human errors, can be unpredictable. Instead, the logical failure domain works in a reactive fashion by confining the suspects among a few controlled servers and does not rely on manual intervention.

## 5 To Whom the EBS Offloads

Offloading the software stack to specialized hardware for better performance or achieving certain features (e.g., bare-metal servers) has been gaining momentum from both the cloud [10, 11] and hardware vendors [8, 12, 13]. Along the evolution of EBS, we have also leveraged FPGA for both frontend BlockClient (i.e., running the customized UDP-based protocol, Solar [34]) and

backend BlockServer for accelerating compression/EC in EBS3. In the following subsections, we first demonstrate the motivations behind the two offloading. More importantly, we discuss why the two eventually both dropped FPGA-based solutions but went on different paths—ASIC for BlockClient and ARM CPU for BlockServer.

### 5.1 Offloading BlockClient

Since EBS2, the frontend BlockClient has become a bottleneck as the backend BlockServers can utilize segmentation for high throughput/IOPS. The BlockClient has been bounded by CPU-heavy tasks, including calculating CRC, encryption and performing per-I/O table lookups. Our stress test reveals it takes 4 CPU cores to saturate a  $2 \times 25$  Gbps NIC in BlockClient. The newly emerged high-speed network would require a doubled or quadrupled number of cores. More importantly, **Elastic Computing Service (ECS**, our VM service)—co-located with BlockClient—requires the servers to be bare-metal-ready (i.e., all CPU cores would be allocated to users). Therefore, offloading BlockClient processing to customized hardware is not only recommended but a must.

We initially built an FPGA-based solution for BlockClient and later decided to directly deploy this version in our production systems. This is because, at the time, directly applying the ASIC-based solution requires much longer development cycles. On the other hand, utilizing the FPGA-based approach is also not desirable due to higher power consumption and increased system complexity.

However, after several years of running in production, we discovered that FPGA is actually not the ideal candidate for BlockClient offloading. The major drawback is the instability. Specifically, 37% of data corruption incidents, as identified by CRC mismatches, are directly caused by FPGA-induced errors such as overheating, signal interference, and timing issues. This is because FPGAs are sensitive to environmental conditions and require precise timing, which can be disrupted by various factors like temperature fluctuations and electrical noise. Moreover, FPGA-related issues account for 22% of BlockClient's operational downtime. Typical reasons include hardware malfunctions, software bugs in the FPGA logic, and incompatibility with updated system components. Apart from reliability concerns, the frequency of FPGA is rather limited (e.g., around 200 MHz to 500 MHz), thereby limiting its potential for adapting to high-speed networks.

Therefore, we later move on to adopt the ASIC-based solution. The preliminary deployment of FPGA-based BlockClient considerably saves our time and effort for transitioning to ASIC (i.e., around 12 months between the release of FPGA and ASIC solutions). Compared to FPGA, ASIC-based offloading incurs approximately 5% of the CapEx and around 1/3 of the power consumption in our EBS scenarios. Also, ASICs are optimized at the hardware level for specific tasks, allowing for more efficient use of resources and higher clock cycles. Another key enabler is that the main functionalities of BlockClient, including data movement, data calculation (e.g., CRC and encryption) and network packet processing, are usually stable. Hence, we do not need to periodically redesign the chips. After deployment, field statistics indicate that the failure rate of ASICs is an order of magnitude lower than that of FPGAs.

### 5.2 Offloading BlockServer

The goal of offloading BlockServer is to reduce costs while maintaining performance. EBS3 introduces data compression in the foreground to reduce traffic and space amplification. However, even with the latency-optimized LZ4 compression algorithm, the compression latency for 16 KiB-sized data blocks remains elevated at  $25 \mu\text{s}$  (25.6% of total write latency) for software-based, and it escalates significantly with larger data blocks. Moreover, to achieve 4,000 MiB/s throughput, at least 8 CPU cores are required, leading to heightened resource contention and diminished performance. Therefore, offloading is necessary to avoid the penalty incurred by software-based compression.

While FPGA-based offloading demonstrates superior performance metrics (see Figure 9), the field deployment has exposed its limitations in terms of reliability and cost-effectiveness. First, BlockServer faces similar FPGA reliability issues. Over the past year, out of every 10,000 deployed production BlockServers, we have documented on average around 150 instances of compression offload failures by FPGA exceptions. Second, we are exploring optimizing compression algorithms tailored to data blocks with varying temperature profiles to minimize storage overhead. For example, using the ZSTD algorithm for separated cold data blocks can further achieve an average 17% space reduction. However, the resource constraints inherent to FPGAs preclude the dynamic adaptation to various compression algorithms. Finally, compared to the scale of BlockClients, the scale of BlockServers is still not large enough to amortize the escalating costs associated with FPGA development.

In light of these considerations, we are reorienting the target of offloading towards server CPUs. This shift is motivated by the advent of multi-core CPUs and specialized computational units integrated within them, which offer superior cost-efficiency while maintaining comparable performance metrics. Noteworthy examples include *Kunpeng 920* ARM CPU [39], and *Yitian 710* ARM CPU [9], all of which are equipped with dedicated units for compression acceleration. The test results show that the average LZ4 compression latency of *Yitian 710* is marginally higher by  $1.3 \mu\text{s}$  in comparison to FPGA-based offload, while 16 ARM cores attain equivalent compression throughput.

Unlike BlockClient, we chose not to use ASIC for BlockServers because of two reasons. First, with no bare-metal requirements, there are no limitations on using CPU cores in BlockServers. Moreover, the BlockServer functionalities (a.k.a. operators) are in an ever-changing fashion. For example, the introduction of new compression and garbage collection algorithms. In this case, applying ASICs may require a complete overhaul from time to time, which can be prohibitively expensive even for the cloud-level scale.

### 5.3 Field Experience and Lessons

First, FPGA is undoubtedly the first choice for many hardware offloading scenarios due to its high flexibility and competitive performance. We, too, adopted FPGA in both BlockClient and BlockServer as proof-of-concept. However, the frequent errors and high CapEx made us realize that FPGA might not be an ideal acceleration option for large scale storage systems.

Second, ASIC and ARM are both suitable for the compute-to-storage architecture but in a different way. Compute end is cost-sensitive due to its massive scale and has stable operators, such as processing (e.g., encryption) and forwarding, matching the characteristics of ASIC. Storage back-end can have frequent upgrades (e.g., improving GC algorithms and optimizing host FTL for ZNS SSDs) and prioritize low interference between tasks. Therefore, the many-core ARM CPU becomes a proper choice.

## 6 Evaluation

In this section, we quantitatively compare the performance of each generation of EBS under various microbenchmarks and macrobenchmarks, including latency, throughput, IOPS, and performance of typical database workloads. Furthermore, we provide data on EBS in field deployment, including failover time, resource consumption, and segment scheduling.

### 6.1 Experiment Setup

We subscribe an ALIBABA CLOUD ECS instance with 128 vCPUs and 512 GiB memory as our testbed. For each generation of EBS, we mount a VD with 20 TiB capacity. Note that the three generations

Table 1. Typical Hardware Specifications for Each Generation of EBS Storage Servers

|               | EBS1                                | EBS2                                     | EBS3                                         |
|---------------|-------------------------------------|------------------------------------------|----------------------------------------------|
| <b>CPU</b>    | Intel(R) Xeon(R) 2.50GHz $\times$ 1 | Intel(R) Xeon(R) Gold 2.30GHz $\times$ 2 | Intel(R) Xeon(R) Platinum 2.50GHz $\times$ 2 |
| <b>Memory</b> | DDR4-2400 16GB $\times$ 4           | DDR4-2666 16GB $\times$ 8                | DDR4-3200 16GB $\times$ 16                   |
| <b>Disk</b>   | SATA3 8TB HDD $\times$ 12           | NVMe 7.8TB TLC SSD $\times$ 12           | NVMe 15.3TB QLC SSD $\times$ 8               |
| <b>NIC</b>    | 2 $\times$ 10Gbps                   | 2 $\times$ 25Gbps                        | 2 $\times$ 100Gbps                           |

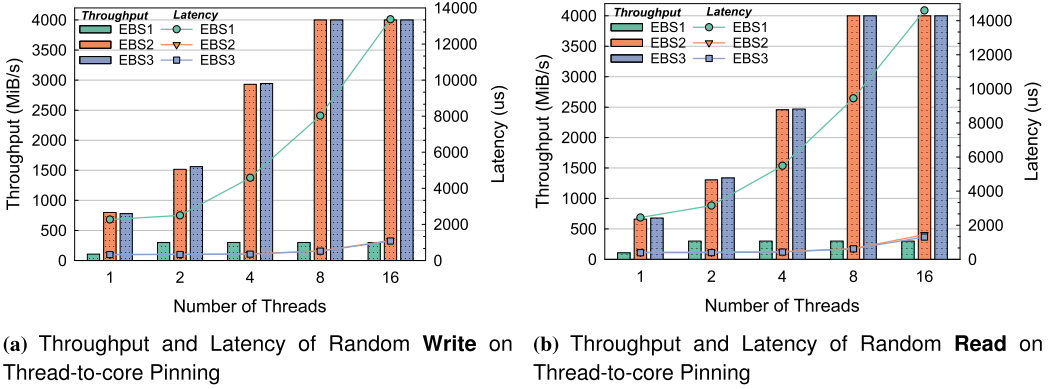


Fig. 13. Random write/read latency of each generation EBS under multiple threads and 4 KiB-sized I/O. Thread-to-core pinning means that each thread occupies one CPU core exclusively.

are backed by different hardware configurations. We list the typical setup in Table 1. We use CentOS 7.9 and enforce direct I/O for all workloads.

## 6.2 Microbenchmark

**Scalability of Throughput/IOPS.** We evaluate the throughput and IOPS of the candidates by stressing random 4KiB write and read. Moreover, to observe the scalability, we increase the number of threads (i.e., FIO jobs) from 1 to 16. Figure 13 shows the overall results.

We can see that the throughput (i.e., bars) of EBS2 and EBS3 increases almost linearly before hitting the peak with 8 threads. With 8 threads, EBS2 and EBS3 can deliver 4,000MB/s throughput per VD (i.e., 1,000K IOPS), which are 13 $\times$  and 50 $\times$  higher than the throughput and IOPS of EBS1.

Moreover, Figure 13 further depicts the latency variation with scaling of FIO jobs. We can observe the latency of EBS2 and EBS3 is similar and remains the same from 1 to 8 threads. Their latencies only experienced a slight increase with 16 jobs due to delays caused by contentions between threads after hitting throughput bottlenecks.

**Single I/O Latency.** Figure 14 presents the distribution of I/O latency and the percentile distribution of 8KiB I/O latency for EBS. The error bars in Figures 14(a) and 14(b) represent the standard deviation, with a larger error bar indicating greater fluctuations in latency.

In terms of write latency, the difference between EBS2 and EBS3 does not exceed 1.6%, and EBS3 has smaller latency variations. For 8KiB write requests, EBS2 and EBS3 exhibit latencies of 110.4us and 110.5us respectively, around 12 $\times$  lower than EBS1. For tail write latency, the P999 latency for EBS1 reaches 4,755us while that for EBS2 and EBS3 are only 655us and 509us. This indicates that it is difficult for EBS1 to achieve superior multi-tenant isolation and low tail latency.

In terms of 8KiB read latency, EBS1 shows an average latency of 1,496us, while EBS2 and EBS3 demonstrate only 203.2us and 208.9us, respectively. The read latency of both EBS2 and EBS3

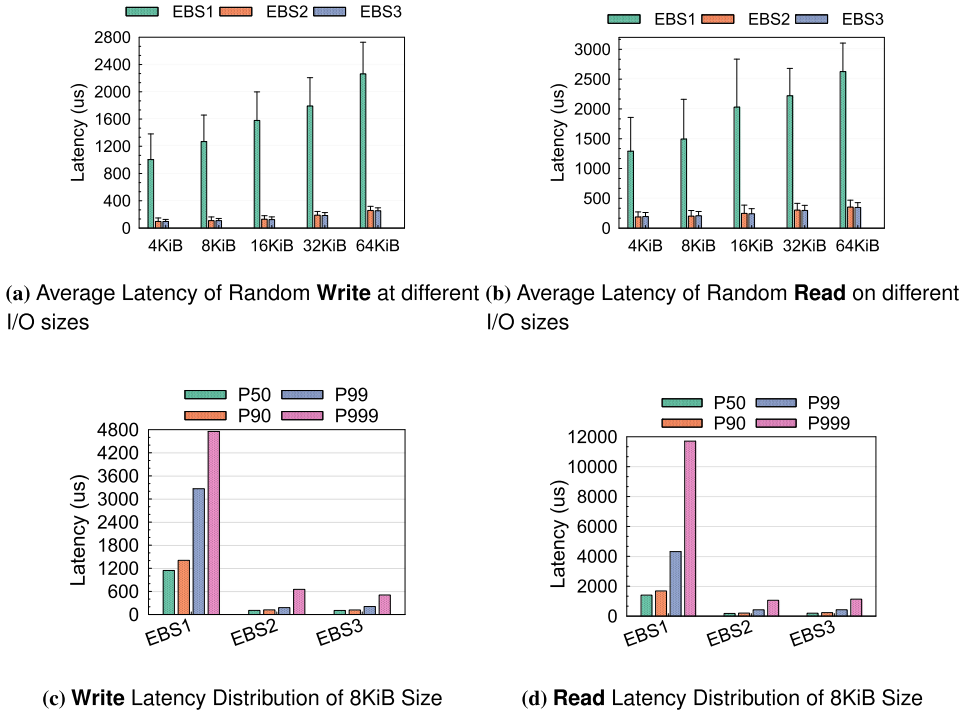


Fig. 14. Latency distribution of random write/read on different I/O sizes.

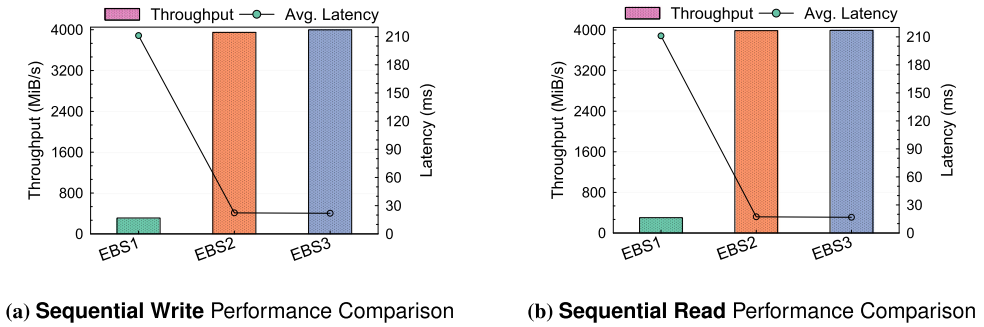


Fig. 15. Sequential write/read performance comparison. Using Fio with iodepth is 64 and blocksize is 1,024 KiB.

depends on retrieving data from DataFile, with their differences mainly arising from changes in physical disk latency. Due to the deployment of more powerful CPUs, hardware decompression offloading, and novel network protocol stacks [28, 34], the read latency of EBS3 is only 2.7% higher than that of EBS2.

**Sequential Performance.** Figure 15 presents the maximum sequential write and read throughput and average latency for different virtual disks. EBS can achieve peak throughput of 4,000 MiB/s at an iodepth of 8 and a blocksize of 1,024 KiB.

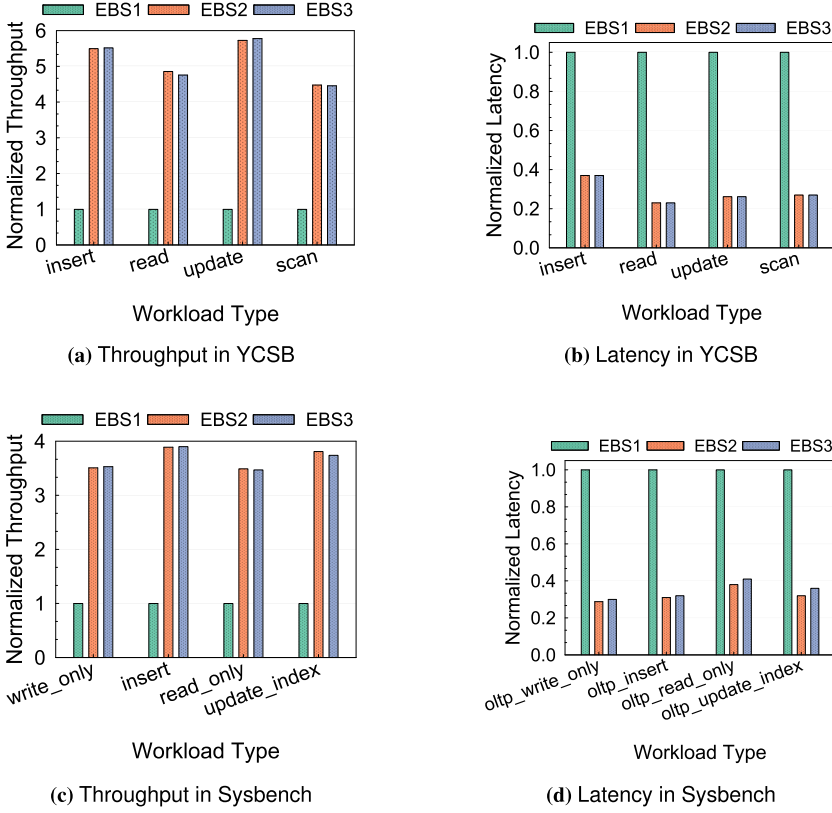


Fig. 16. Throughput and latency comparison of each generation EBS using the **YCSB** and **Sysbench** (normalized to EBS1).

### 6.3 Macrobenchmark

We evaluate the performance of each generation of EBS with two IO-intensive database applications. We use RocksDB with the default configuration parameters. In MySQL, we select the storage engine as InnoDB and configure the user buffer size as 1 GiB, the page size as 16 KiB, the flushing method as direct I/O, the ramp-up time as 180 seconds and the execution time as 20 minutes.

Figure 16 shows the YCSB results. Compared to EBS1, we observe 550% and 573% gains in throughput for workload insert and update respectively, with 63% and 74% decreases in latency. Both these workloads are write-dominated. For read-dominated workloads, the throughput experiences an approximately 470% increase, while the latency is reduced by 77%.

Figures 16(c) and (d) show the Sysbench results. Under oltp\_insert workload, the throughput of EBS2 and EBS3 is increased by 389% compared to EBS1, and the latency is decreased by 69.1%. Under other workloads, the throughput is increased by an average of 350%, while the latency is decreased by 65.0%.

### 6.4 Field Deployment Statistics

**Block-layer Failover Time.** In production deployment, BlockServer and BlockMaster may crash due to various accidents such as **machine check exceptions (MCEs)**, SSD failures, and core switch failures. Therefore, shortening failover time is critical for guaranteeing SLA. We conducted

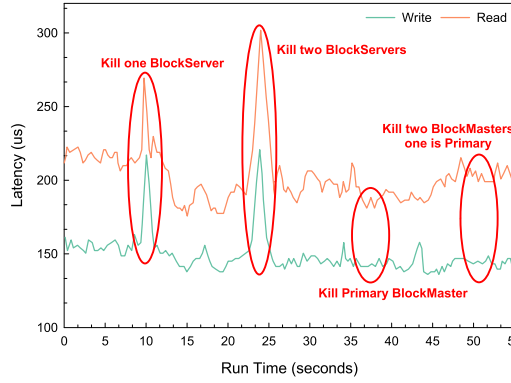


Fig. 17. The variation of I/O latency during the failover of BlockServer and BlockMaster.

Table 2. Operational Data of Typical EBS2 and EBS3 Clusters

| Cluster | Cluster Type | Number of BlockServers | Number of Live Segments | CPU Utilization Distribution of BlockServer (Avg./P90/P99) | Memory Usage Distribution of BlockServer (GiB) (Avg./P90/P99) |
|---------|--------------|------------------------|-------------------------|------------------------------------------------------------|---------------------------------------------------------------|
| A       | EBS2         | 96                     | 152,576                 | 1045% / 1068% / 1130%                                      | 15.1 / 16.7 / 18.5                                            |
| B       | EBS3         | 96                     | 539,420                 | 2069% / 2096% / 2170%                                      | 31.5 / 33.4 / 36.4                                            |

a failover performance test of EBS3 on a storage cluster comprising 96 BlockServers, by reproducing the real traffic pattern of the production cluster. Figure 17 illustrates the I/O latency variations observed during the test. In case of a single BlockServer failover, the average write latency increased by 28.6%, before recovering after 780ms. For two BlockServer failovers, the latency increases by 30.9%, and returns to normal within 860ms.

**Resource Consumption.** We choose two typical EBS2 and EBS3 clusters with 96 BlockServers each to demonstrate the resource consumption of EBS in production. Table 2 shows that EBS2 and EBS3 clusters serve 71,520 and 235,464 VDs, respectively. The average CPU usage of each BlockServer in EBS2 is 1045% and 2069% in EBS3. This difference is mainly due to the higher bandwidth of the network card of the EBS3 cluster, resulting in heavier network and storage processing. In terms of memory usage, since EBS3 deploys extra cache in memory (i.e., SegmentCache), the average memory usage of EBS3 BlockServer is about 16 GiB higher than that of EBS2 (i.e., 31.5 GiB). Since EBS3 cluster serves more segments, its BlockMaster memory reaches 7.1 GiB, which is 2.5× that of the EBS2 cluster.

**Segment Scheduling.** EBS can migrate segments from one BlockServer to another due to load balancing and failure avoidance. However, frequent segment migrations can also cause extra traffic and may even affect the SLA. Table 3 presents the average number of segment schedulings per day across all production clusters over the past six months.

We categorize all segment scheduling into two types: Basic Scheduling and Advanced Scheduling. Basic Scheduling provides the fundamental guarantee of EBS availability and is subject to minimal changes. Advanced Scheduling optimizes availability across various scenarios, and it evolves to meet the changing and complex scheduling requirements. In the field, we observe that, apart from migration due to disk creation and deletion operations, EBS carries out around 1.3 million segment migrations in all production clusters per day, with traffic balancing and segment count balancing accounting for 91%. 6.04% of segment schedulings are attributed to basic scheduling, with Report\_Error contributing to 4.6% and Server\_Crash caused by machine check exceptions (MCE)

Table 3. Average Number of Segment Schedulings per Day across all Production Clusters over the Past Six Months

|                     | Schedule Strategy           | Schedule Reason                                                                            | Schedule Count | Schedule Percentage |
|---------------------|-----------------------------|--------------------------------------------------------------------------------------------|----------------|---------------------|
| Basic Scheduling    | PLAN_Report_Error           | BlockClient reports that I/O is time-out.                                                  | 61,393         | 4.60%               |
|                     | PLAN_Server_Crash           | The heartbeat of the BlockServer is lost.                                                  | 18,798         | 1.41%               |
|                     | PLAN_Report_Abnormal_Server | Multiple BlockClients report that the BlockServer is time-out, while its heartbeat exists. | 431            | 0.03%               |
| Advanced Scheduling | PLAN_Traffic_Balance        | Traffic balancing among BlockServers.                                                      | 695,698        | 52.08%              |
|                     | PLAN_SegmentCount_Balance   | Segment Count balancing among BlockServers.                                                | 521,577        | 39.05%              |
|                     | PLAN_Reload_Device          | Virtual disk migration among storage clusters.                                             | 28,182         | 2.11%               |
| Other               | Other                       | BlockMasters with old version have no attribution.                                         | 9,689          | 0.73%               |

occurring an average of 30 times per day, resulting in 18,798 segment schedulings. Report\_Error is critical to guaranteeing the SLA. When BlockClient detects an I/O timeout, it immediately reports the target BlockServer exception to BlockMaster, and then initiates the segment migration to recover I/O. Since segment migration is lightweight enough (the migration time of a single segment is less than 80ms), we set the timeout threshold to 100ms to trigger Report\_Error as much as possible to avoid I/O hang.

## 7 Discussions and Future Directions

Evolving EBS has been a long journey. Along the way, it is not surprising we have conceived or even tried out promising ideas that are only later to be proved as impractical. Here, we summarize such discussions and further share our views on the future of EBS.

**Fairness.** A common concern in a multi-tenant environment is fairness. In EBS1, it is challenging to ensure **quality-of-service (QoS)** among different VDs because of unpredictable load spikes and heavy VD migration overhead in the storage cluster. Since EBS2, we have introduced the concept of segment and rely on the intra-/inter-cluster segment migration mechanism to balance the load among BlockServers. This mechanism helps us to guarantee the VDs can receive a fair share of resources.

**Driving force behind the evolution.** From EBS1 to EBS2, EBS3, EBSX and beyond, our shifted focus are essentially driven by balancing the performance, space, and cost. For example, developing EBS2 is motivated by then readily available flash drives. However, due to the much higher price tags per GiB, EC and compression become necessary. Similarly for EBS3, when the cost of flash drives drops, the network bandwidth becomes the next bottleneck, which then pushes us to realize frontend compression. One interesting observation we made along the evolution is that, with the next-gen hardware (e.g., Ultra-low latency NVMe SSDs), we are able to provide users with VDs that render current-gen (e.g., SATA SSDs) experience. This is especially true in terms of latency and QoS.

**Not adopting log-structured design.** When developing EBS2, we initially tried to extend the EBS1 with segmentation but later dropped the idea due to high engineering effort. Note that this idea (i.e., in-place update with segmentation), if developed, can meet our requirements, including high performance and space efficiency, for EBS2.

However, this idea still falls short for EBS3. Foreground EC/compression necessitates that a storage system aggregates a sufficient amount of data before the persistence to achieve efficient data reductions. For example, EC(8,3) needs 32 KiB data for one stripe with 4 KiB stripe unit size, and 16 KiB compression units yield higher compression ratios. In Section 2.2 we have shown the domination of small writes (less than 16 KiB) combined with the need for low write latency prevent

us from aggregating 16 KiB data for a single segment. The log-structured design can easily tackle this problem by allowing segments from different VD's to be merged together and flushed to a journal log.

We believe that Ceph [3] faces a similar issue due to the lack of a log-structured layer like BlockServer. To utilize EC in Ceph Block Device and CephFS, with FileStore, users need to set up a cache tier with write-back mode before an erasure-coded pool. With BlueStore, Ceph performs partial writes (i.e., read-modify-write) to an existing stripe without aggregating data, yielding additional network overhead, while EBS3 always writes full EC stripes to Pangu.

**Building EBS with open-source software.** At first glance, one might assume that a cloud-level block store could be constructed by using a series of open-source softwares. For example, we can use HBase [2] (i.e., a log-structured distributed key-value store) and HDFS [7] (i.e., a distributed file system) to replace the block layer (i.e., BlockServer and BlockManager) and the file layer (i.e., Pangu).

However, the two designs are markedly different. EBS is a co-design of the block interface, the software, and the hardware. In the block layer, indexing is specifically tailored for the block service. For each segment, the key space is deterministic. Specifically, each segment represents an address space of a 32 GiB segment consisting of 4 KiB blocks, and thus the key space only includes 8 million numbers. This deterministic feature allows for efficient memory allocation for the indexing structure. To achieve low I/O average and tail latency, EBS deploys hardware offloading (i.e., offload compression algorithms to FPGA and ARM CPU) and customized network protocols [34], and decouples non-I/O activities from the critical I/O path, e.g., using individual GCWorker to perform garbage collection and moving index compaction procedures to background threads. In the file layer, Pangu is a high-performance storage system that builds dedicated user-space file systems for high-speed storage devices (i.e., NVMe SSDs), deploys high-speed networks (i.e., RDMA), and incorporates various hardware-offloading technologies [33].

**Not separating Pangu from EBS.** The short answer is that such integration would have significantly hindered the development of EBS. Recall that in EBS1, the BlockManagers and BlockServers are integrated with the persistence layer (i.e., ChunkServers and ChunkManagers). This organization becomes increasingly unacceptable with the scaling of our engineering team. First, the interfaces (between the block and chunk layers) grew exceedingly complex—at one point there were nearly 10 sets of persistence APIs in EBS1—in order to support various functionalities and performance optimizations. Maintaining such a complex codebase undoubtedly slows down the development schedule. As an example, around that time, it could take up to 10 months for EBS to release a major upgrade due to delays by software bugs or incompatibility between components.

Decoupling the underlying persistence layer (i.e., Pangu) from EBS and adopting a unified log-structured interface have clearly facilitated the development and ease of communication. In addition, the independent block layer enables rapid segment creation and migration across multiple storage nodes, irrespective of data block locations. The separated architecture also localizes the impact radius of single-point failure to each respective layer. Moreover, this disaggregation also allows us to integrate emerging technologies (e.g., FPGA-based accelerator) and extend Pangu as a general-purpose DFS for other services (e.g., object store [17] and file store [15]).

**Future work.** After having deployed EBS2 and EBS3 at scale, we have identified several challenges and potential directions for future work.

- *Dynamic resource allocation.* Burst traffic is prevalent in the EBS. In earlier versions, we statically allocate resources, such as CPU cores, memory, and network bandwidth, to different processes. However, we have discovered several cases in deployment where GCWorker (handling background traffic) and LSB Core (managing foreground VD's requests) both reach

peak pressure with few available resources, leading to high tail latency. To address this, we set a background task to periodically scan the pressure of each process and allocate additional resources to LSBDCore if necessary. This is successful in mitigating the surges in tail latency even during the largest annual online shopping event of ALIBABA.

- *Garbage collection.* Due to the log-structured design, in EBS2 and EBS3, there are GC both in the block layer and inside the SSDs. To further improve I/O efficiency, we plan to build EBS with only one GC from end to end. Our tentative solution is to employ ZNS SSDs [22] and customize host-side FTL to allow EBS GCWorker to directly reclaim garbage at the zone level.
- *Segment scheduling optimization.* EBS can migrate segments from one BlockServer to another due to load balancing and failure avoidance. However, frequent segment migrations can also cause extra traffic and may even affect the SLA. In the field, we observe that, apart from migration due to disk creation and deletion operations, EBS carries out around 1.3 million segment migrations in all production clusters per day, with traffic balancing and segment count balancing accounting for 91%. Currently, we are diagnosing the root causes, analyzing the triggering conditions, and developing potential scheduling algorithms to alleviate this issue.

## 8 Related Work

As a popular service, most cloud vendors provide block storage or similar service to the public, such as Amazon EBS [18], Azure Disk Storage [19, 25], and Google Persistent Disk [14, 29]. Unfortunately, there is little work discussing the detailed architecture and the design tradeoffs in block storage. Our paper reveals the evolution behind the three generations of EBS designs at ALIBABA CLOUD. The evaluations confirm our design choices and demonstrate competitive performance against similar products from other vendors.

Building Virtual Disks prototype systems are also popular in the research community, such as Salus [38], Ursa [32], Blizzard [35], and LSVD [30]. Similar to our EBS design, many of them also utilize the log-structured design. The most fundamental difference between their work and ours is that EBS is a remote block storage system that follows compute-to-storage architecture. Moreover, as a deployed system, we pay attention to not only performance but also CapEx efficiency, such as minimizing storage cost and traffic amplification.

## 9 Conclusion

This paper systematically revisits the technical evolutions behind the three generations of EBS at ALIBABA CLOUD. We analyze the pros and cons of each design. We further discuss our experience and corresponding optimizations from the field deployment and potential directions from four perspectives, including elasticity, availability, hardware offloading, and potential alternative solutions. We hope that this paper not only serves as the first paper on detailed designs of cloud block store but also can provide practical lessons for optimizing such storage systems.

## References

- [1] Apache Hadoop 3.3.6 - HDFS Federation. <https://hadoop.apache.org/docs/stable/hadoop-project-dist/hadoop-hdfs/Federation.html> (n.d.).
- [2] Apache Hbase. <https://hbase.apache.org/> (n.d.).
- [3] Erasure code – Ceph documentation. <https://docs.ceph.com/en/latest/rados/operations/erasure-code> (n.d.).
- [4] HBASE-14598. <https://issues.apache.org/jira/browse/HBASE-14598> (n.d.).
- [5] HBASE-22072. <https://issues.apache.org/jira/browse/HBASE-22072> (n.d.).
- [6] HBASE-22862. <https://issues.apache.org/jira/browse/HBASE-22862> (n.d.).
- [7] HDFS Architecture Guide. [https://hadoop.apache.org/docs/r1.2.1/hdfs\\_design.html](https://hadoop.apache.org/docs/r1.2.1/hdfs_design.html) (n.d.).

- [8] 2020. The fungible DPU™: A new category of microprocessor for the data-centric era: Hot chips 2020. In *Proc. of IEEE HCS*.
- [9] 2021. Alibaba Cloud Unveils New Server Chips to Optimize Cloud Computing Services. <https://www.alibabacloud.com/blog/598159?spm=a3c0i.23458820.2359477120.113.66117d3fm03t9b> (2021).
- [10] 2021. AWS Nitro System. <https://aws.amazon.com/ec2/nitro/> (2021).
- [11] 2021. A Detailed Explanation about Alibaba Cloud CIPU. <https://www.alibabacloud.com/blog/599183?spm=a3c0i.23458820.2359477120.3.76806e9bESi3SD> (2021).
- [12] 2021. Intel Infrastructure Processing Unit. <https://www.intel.com/content/www/us/en/products/details/network-io/ipu/e2000-asic.html> (2021).
- [13] 2021. NVIDIA BlueField Data Processing Units. <https://www.nvidia.com/en-us/networking/products/data-processing-unit/> (2021).
- [14] 2023. About Google Persistent Disk. <https://cloud.google.com/compute/docs/disks> (2023).
- [15] 2023. Alibaba Cloud Apsara File Storage NAS. <https://www.alibabacloud.com/help/en/nas> (2023).
- [16] 2023. Alibaba Cloud Elastic Block Storage Devices. <https://www.alibabacloud.com/help/en/elastic-compute-service/latest/block-storage-overview-elastic-block-storage-devices> (2023).
- [17] 2023. Alibaba Cloud Object Storage Service. <https://www.alibabacloud.com/help/en/object-storage-service> (2023).
- [18] 2023. Amazon Elastic Block Store. <https://docs.aws.amazon.com/AWSEC2/latest/UserGuide/AmazonEBS.html> (2023).
- [19] 2023. Introduction to Azure Managed Disks. <https://learn.microsoft.com/en-us/azure/virtual-machines/managed-disks-overview> (2023).
- [20] Mohammadamin Ajdari, Wonsik Lee, Pyeongsu Park, Joonsung Kim, and Jangwoo Kim. 2019. FIDR: A scalable storage system for fine-grain inline data reduction with efficient memory handling. In *Proc. of IEEE/ACM MICRO*.
- [21] Doug Beaver, Sanjeev Kumar, Harry C. Li, Jason Sobel, and Peter Vajgel. 2010. Finding a needle in haystack: Facebook's photo storage. In *Proc. of USENIX OSDI*.
- [22] Matias Bjorling, Abutalib Aghayev, Hans Holmberg, Aravind Ramesh, Damien Le Moal, Gregory R. Ganger, and George Amvrosiadis. 2021. ZNS: Avoiding the block interface tax for flash-based SSDs. In *Proc. of USENIX ATC*.
- [23] James Bornholt, Rajeev Joshi, Vytautas Astrauskas, Brendan Cully, Bernhard Kragl, Seth Markle, Kyle Sauri, Drew Schleit, Grant Slatton, Serdar Tasiran, et al. 2021. Using lightweight formal methods to validate a key-value storage node in Amazon S3. In *Proc. of ACM SOSP*.
- [24] Marc Brooker, Tao Chen, and Fan Ping. 2020. Millions of tiny databases. In *Proc. of USENIX NSDI*.
- [25] Brad Calder, Ju Wang, Aaron Ogus, Niranjana Nilakantan, Arild Skjolsvold, Sam McKelvie, Yikang Xu, Shashwat Srivastav, Jiesheng Wu, Huseyin Simitci, et al. 2011. Windows Azure storage: A highly available cloud storage service with strong consistency. In *Proc. of ACM SOSP*.
- [26] Sebastian Deorowicz. 2014. Silesia Compression Corpus. <https://sun.aei.polsl.pl/~sdeor/index.php?page=silesia> (2014).
- [27] Peter DeSantis. 2022. Peter DeSantis Keynote Recap - AWS re:Invent 2022. <https://caylent.com/blog/peter-de-santis-keynote-recap-aws-re-invent-2022> (2022).
- [28] Yixiao Gao, Qiang Li, Lingbo Tang, Yongqing Xi, Pengcheng Zhang, Wenwen Peng, Bo Li, Yaohui Wu, Shaozong Liu, Lei Yan, et al. 2021. When cloud storage meets RDMA. In *Proc. of USENIX NSDI*.
- [29] Sanjay Ghemawat, Howard Gobioff, and Shun-Tak Leung. 2003. The Google file system. In *Proc. of ACM SOSP*.
- [30] Mohammad Hossein Hajkazemi, Vojtech Aschenbrenner, Mania Abdi, Emine Ugur Kaynar, Amin Mossayebzadeh, Orran Krieger, and Peter Desnoyers. 2022. Beating the I/O bottleneck: A case for log-structured virtual disks. In *Proc. of ACM EuroSys*.
- [31] Lexiang Huang, Matthew Magnusson, Abishek Bangalore Muralikrishna, Salman Estyak, Rebecca Isaacs, Abutalib Aghayev, Timothy Zhu, and Aleksey Charapko. 2022. Metastable failures in the wild. In *Proc. of USENIX OSDI*.
- [32] Huiba Li, Yiming Zhang, Dongsheng Li, Zhiming Zhang, Shengyun Liu, Peng Huang, Zheng Qin, Kai Chen, and Yongqiang Xiong. 2019. URSA: Hybrid block storage for cloud-scale virtual disks. In *Proc. of ACM EuroSys*.
- [33] Qiang Li, Qiao Xiang, Yuxin Wang, Haohao Song, Ridi Wen, Wenhui Yao, Yuanyuan Dong, Shuqi Zhao, Shuo Huang, Zhaosheng Zhu, et al. 2023. More than capacity: Performance-oriented evolution of Pangu in Alibaba. In *Proc. of USENIX FAST*.
- [34] Rui Miao, Lingjun Zhu, Shu Ma, Kun Qian, Shujun Zhuang, Bo Li, Shuguang Cheng, Jiaqi Gao, Yan Zhuang, Pengcheng Zhang, et al. 2022. From luna to solar: The evolutions of the compute-to-storage networks in Alibaba Cloud. In *Proc. of ACM SIGCOMM*.
- [35] James Mickens, Edmund B. Nightingale, Jeremy Elson, Darren Gehring, Bin Fan, Asim Kadav, Vijay Chidambaram, Osama Khan, and Krishna Nareddy. 2014. Blizzard: Fast, cloud-scale block storage for cloud-oblivious applications. In *Proc. of USENIX NSDI*.
- [36] Mendel Rosenblum and John K. Ousterhout. 1991. The design and implementation of a log-structured file system. In *Proc. of ACM SOSP*.

- [37] Qiuping Wang, Jinhong Li, Patrick P. C. Lee, Tao Ouyang, Chao Shi, and Lilong Huang. 2022. Separating data via block invalidation time inference for write amplification reduction in log-structured storage. In *Proc. of USENIX FAST*.
- [38] Yang Wang, Manos Kapritsos, Zuocheng Ren, Prince Mahajan, Jeevitha Kirubanandam, Lorenzo Alvisi, and Mike Dahlin. 2013. Robustness in the Salus scalable block store. In *Proc. of USENIX NSDI*.
- [39] Jing Xia, Chuanning Cheng, Xiping Zhou, Yuxing Hu, and Peter Chun. 2021. Kunpeng 920: The first 7-nm chiplet-based 64-Core ARM SoC for cloud services. *Proc. of IEEE Micro* (2021). <https://doi.org/10.1109/MM.2021.3085578>
- [40] Jie Zhang, Miryeong Kwon, Donghyun Gouk, Sungjoon Koh, Changlim Lee, Mohammad Alian, Myoungjun Chun, Mahmut Taylan Kandemir, Nam Sung Kim, Jihong Kim, et al. 2018. FlashShare: Punching through server storage stack from kernel to firmware for Ultra-Low latency SSDs. In *Proc. of USENIX OSDI*.
- [41] Teng Zhang, Jianying Wang, Xuntao Cheng, Hao Xu, Nanlong Yu, Gui Huang, Tieying Zhang, Dengcheng He, Feifei Li, Wei Cao, et al. 2020. FPGA-accelerated compactions for LSM-based key-value store. In *Proc. of USENIX FAST*.
- [42] Xiantao Zhang, Xiao Zheng, Zhi Wang, Hang Yang, Yibin Shen, and Xin Long. 2020. High-density multi-tenant bare-metal cloud. In *Proc. of ACM ASPLOS*.
- [43] Bohong Zhu, Youmin Chen, Qing Wang, Youyou Lu, and Jiwu Shu. 2021. Octopus+: An RDMA-enabled distributed persistent memory file system. *ACM Trans. on Storage* 17, 3 (2021), 1–25.
- [44] Lingjun Zhu, Yifan Shen, Erci Xu, Bo Shi, Ting Fu, Shu Ma, Shuguang Chen, Zhongyu Wang, Haonan Wu, Xingyu Liao, et al. 2023. Deploying user-space TCP at cloud scale with LUNA. In *Proc. of USENIX ATC*.

Received 13 May 2024; revised 13 May 2024; accepted 11 September 2024