



Red Hat Ceph Storage 7

Hardware Guide

Hardware selection recommendations for Red Hat Ceph Storage

Red Hat Ceph Storage 7 Hardware Guide

Hardware selection recommendations for Red Hat Ceph Storage

Legal Notice

Copyright © 2024 Red Hat, Inc.

The text of and illustrations in this document are licensed by Red Hat under a Creative Commons Attribution–Share Alike 3.0 Unported license ("CC-BY-SA"). An explanation of CC-BY-SA is available at

<http://creativecommons.org/licenses/by-sa/3.0/>

. In accordance with CC-BY-SA, if you distribute this document or an adaptation of it, you must provide the URL for the original version.

Red Hat, as the licensor of this document, waives the right to enforce, and agrees not to assert, Section 4d of CC-BY-SA to the fullest extent permitted by applicable law.

Red Hat, Red Hat Enterprise Linux, the Shadowman logo, the Red Hat logo, JBoss, OpenShift, Fedora, the Infinity logo, and RHCE are trademarks of Red Hat, Inc., registered in the United States and other countries.

Linux[®] is the registered trademark of Linus Torvalds in the United States and other countries.

Java[®] is a registered trademark of Oracle and/or its affiliates.

XFS[®] is a trademark of Silicon Graphics International Corp. or its subsidiaries in the United States and/or other countries.

MySQL[®] is a registered trademark of MySQL AB in the United States, the European Union and other countries.

Node.js[®] is an official trademark of Joyent. Red Hat is not formally related to or endorsed by the official Joyent Node.js open source or commercial project.

The OpenStack[®] Word Mark and OpenStack logo are either registered trademarks/service marks or trademarks/service marks of the OpenStack Foundation, in the United States and other countries and are used with the OpenStack Foundation's permission. We are not affiliated with, endorsed or sponsored by the OpenStack Foundation, or the OpenStack community.

All other trademarks are the property of their respective owners.

Abstract

This document provides high level guidance on selecting hardware for use with Red Hat Ceph Storage. Red Hat is committed to replacing problematic language in our code, documentation, and web properties. We are beginning with these four terms: master, slave, blacklist, and whitelist. Because of the enormity of this endeavor, these changes will be implemented gradually over several upcoming releases. For more details, see our CTO Chris Wright's message

Table of Contents

CHAPTER 1. EXECUTIVE SUMMARY	3
CHAPTER 2. GENERAL PRINCIPLES FOR SELECTING HARDWARE	4
2.1. IDENTIFY PERFORMANCE USE CASE	4
2.2. CONSIDER STORAGE DENSITY	4
2.3. IDENTICAL HARDWARE CONFIGURATION	5
2.4. NETWORK CONSIDERATIONS FOR RED HAT CEPH STORAGE	5
2.5. AVOID USING RAID SOLUTIONS	6
2.6. SUMMARY OF COMMON MISTAKES WHEN SELECTING HARDWARE	6
CHAPTER 3. OPTIMIZE WORKLOAD PERFORMANCE DOMAINS	8
CHAPTER 4. SERVER AND RACK SOLUTIONS	10
CHAPTER 5. MINIMUM HARDWARE RECOMMENDATIONS FOR CONTAINERIZED CEPH	14
CHAPTER 6. RECOMMENDED MINIMUM HARDWARE REQUIREMENTS FOR THE RED HAT CEPH STORAGE DASHBOARD	16

CHAPTER 1. EXECUTIVE SUMMARY

Many hardware vendors now offer both Ceph-optimized servers and rack-level solutions designed for distinct workload profiles. To simplify the hardware selection process and reduce risk for organizations, Red Hat has worked with multiple storage server vendors to test and evaluate specific cluster options for different cluster sizes and workload profiles. Red Hat's exacting methodology combines performance testing with proven guidance for a broad range of cluster capabilities and sizes. With appropriate storage servers and rack-level solutions, Red Hat Ceph Storage can provide storage pools serving a variety of workloads—from throughput-sensitive and cost and capacity-focused workloads to emerging IOPS-intensive workloads.

Red Hat Ceph Storage significantly lowers the cost of storing enterprise data and helps organizations manage exponential data growth. The software is a robust and modern petabyte-scale storage platform for public or private cloud deployments. Red Hat Ceph Storage offers mature interfaces for enterprise block and object storage, making it an optimal solution for active archive, rich media, and cloud infrastructure workloads characterized by tenant-agnostic OpenStack® environments ^[1]. Delivered as a unified, software-defined, scale-out storage platform, Red Hat Ceph Storage lets businesses focus on improving application innovation and availability by offering capabilities such as:

- Scaling to hundreds of petabytes ^[2].
- No single point of failure in the cluster.
- Lower capital expenses (CapEx) by running on commodity server hardware.
- Lower operational expenses (OpEx) with self-managing and self-healing properties.

Red Hat Ceph Storage can run on myriad industry-standard hardware configurations to satisfy diverse needs. To simplify and accelerate the cluster design process, Red Hat conducts extensive performance and suitability testing with participating hardware vendors. This testing allows evaluation of selected hardware under load and generates essential performance and sizing data for diverse workloads—ultimately simplifying Ceph storage cluster hardware selection. As discussed in this guide, multiple hardware vendors now provide server and rack-level solutions optimized for Red Hat Ceph Storage deployments with IOPS-, throughput-, and cost and capacity-optimized solutions as available options.

Software-defined storage presents many advantages to organizations seeking scale-out solutions to meet demanding applications and escalating storage needs. With a proven methodology and extensive testing performed with multiple vendors, Red Hat simplifies the process of selecting hardware to meet the demands of any environment. Importantly, the guidelines and example systems listed in this document are not a substitute for quantifying the impact of production workloads on sample systems.

[1] Ceph is and has been the leading storage for OpenStack according to several semi-annual OpenStack user surveys.

[2] See [Yahoo Cloud Object Store - Object Storage at Exabyte Scale](#) for details.

CHAPTER 2. GENERAL PRINCIPLES FOR SELECTING HARDWARE

As a storage administrator, you must select the appropriate hardware for running a production Red Hat Ceph Storage cluster. When selecting hardware for Red Hat Ceph Storage, review these following general principles. These principles will help save time, avoid common mistakes, save money and achieve a more effective solution.

Prerequisites

- A planned use for Red Hat Ceph Storage.
- Linux System Administration Advance level with Red Hat Enterprise Linux certification.
- Storage administrator with Ceph Certification.

2.1. IDENTIFY PERFORMANCE USE CASE

One of the most important steps in a successful Ceph deployment is identifying a price-to-performance profile suitable for the cluster's use case and workload. It is important to choose the right hardware for the use case. For example, choosing IOPS-optimized hardware for a cloud storage application increases hardware costs unnecessarily. Whereas, choosing capacity-optimized hardware for its more attractive price point in an IOPS-intensive workload will likely lead to unhappy users complaining about slow performance.

The primary use cases for Ceph are:

- **IOPS optimized:** IOPS optimized deployments are suitable for cloud computing operations, such as running MySQL or MariaDB instances as virtual machines on OpenStack. IOPS optimized deployments require higher performance storage such as 15k RPM SAS drives and separate SSD journals to handle frequent write operations. Some high IOPS scenarios use all flash storage to improve IOPS and total throughput.
- **Throughput optimized:** Throughput-optimized deployments are suitable for serving up significant amounts of data, such as graphic, audio and video content. Throughput-optimized deployments require networking hardware, controllers and hard disk drives with acceptable total throughput characteristics. In cases where write performance is a requirement, SSD journals will substantially improve write performance.
- **Capacity optimized:** Capacity-optimized deployments are suitable for storing significant amounts of data as inexpensively as possible. Capacity-optimized deployments typically trade performance for a more attractive price point. For example, capacity-optimized deployments often use slower and less expensive SATA drives and co-locate journals rather than using SSDs for journaling.

This document provides examples of Red Hat tested hardware suitable for these use cases.

2.2. CONSIDER STORAGE DENSITY

Hardware planning should include distributing Ceph daemons and other processes that use Ceph across many hosts to maintain high availability in the event of hardware faults. Balance storage density considerations with the need to rebalance the cluster in the event of hardware faults. A common hardware selection mistake is to use very high storage density in small clusters, which can overload networking during backfill and recovery operations.

2.3. IDENTICAL HARDWARE CONFIGURATION

Create pools and define CRUSH hierarchies such that the OSD hardware within the pool is identical.

- Same controller.
- Same drive size.
- Same RPMs.
- Same seek times.
- Same I/O.
- Same network throughput.
- Same journal configuration.

Using the same hardware within a pool provides a consistent performance profile, simplifies provisioning and streamlines troubleshooting.

2.4. NETWORK CONSIDERATIONS FOR RED HAT CEPH STORAGE

An important aspect of a cloud storage solution is that storage clusters can run out of IOPS due to network latency, and other factors. Also, the storage cluster can run out of throughput due to bandwidth constraints long before the storage clusters run out of storage capacity. This means that the network hardware configuration must support the chosen workloads to meet price versus performance requirements.

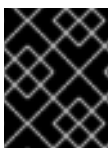
Storage administrators prefer that a storage cluster recovers as quickly as possible. Carefully consider bandwidth requirements for the storage cluster network, be mindful of network link oversubscription, and segregate the intra-cluster traffic from the client-to-cluster traffic. Also consider that network performance is increasingly important when considering the use of Solid State Disks (SSD), flash, NVMe, and other high performing storage devices.

Ceph supports a public network and a storage cluster network. The public network handles client traffic and communication with Ceph Monitors. The storage cluster network handles Ceph OSD heartbeats, replication, backfilling, and recovery traffic. At a **minimum**, a single 10 GB Ethernet link should be used for storage hardware, and you can add additional 10 GB Ethernet links for connectivity and throughput.



IMPORTANT

Red Hat recommends allocating bandwidth to the storage cluster network, such that it is a multiple of the public network using the **osd_pool_default_size** as the basis for the multiple on replicated pools. Red Hat also recommends running the public and storage cluster networks on separate network cards.



IMPORTANT

Red Hat recommends using 10 GB Ethernet for Red Hat Ceph Storage deployments in production. A 1 GB Ethernet network is not suitable for production storage clusters.

In the case of a drive failure, replicating 1 TB of data across a 1 GB Ethernet network takes 3 hours, and 3 TB takes 9 hours. Using 3 TB is the typical drive configuration. By contrast, with a 10 GB Ethernet network, the replication times would be 20 minutes and 1 hour. Remember that when a Ceph OSD fails,

the storage cluster will recover by replicating the data it contained to other Ceph OSDs within the pool.

The failure of a larger domain such as a rack means that the storage cluster utilizes considerably more bandwidth. When building a storage cluster consisting of multiple racks, which is common for large storage implementations, consider utilizing as much network bandwidth between switches in a "fat tree" design for optimal performance. A typical 10 GB Ethernet switch has 48 10 GB ports and four 40 GB ports. Use the 40 GB ports on the spine for maximum throughput. Alternatively, consider aggregating unused 10 GB ports with QSFP+ and SFP+ cables into more 40 GB ports to connect to other rack and spine routers. Also, consider using LACP mode 4 to bond network interfaces. Additionally, use jumbo frames, with a maximum transmission unit (MTU) of 9000, especially on the backend or cluster network.

Before installing and testing a Red Hat Ceph Storage cluster, verify the network throughput. Most performance-related problems in Ceph usually begin with a networking issue. Simple network issues like a kinked or bent Cat-6 cable could result in degraded bandwidth. Use a minimum of 10 GB ethernet for the front side network. For large clusters, consider using 40 GB ethernet for the backend or cluster network.

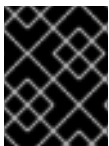


IMPORTANT

For network optimization, Red Hat recommends using jumbo frames for a better CPU per bandwidth ratio, and a non-blocking network switch back-plane. Red Hat Ceph Storage requires the same MTU value throughout all networking devices in the communication path, end-to-end for both public and cluster networks. Verify that the MTU value is the same on all hosts and networking equipment in the environment before using a Red Hat Ceph Storage cluster in production.

2.5. AVOID USING RAID SOLUTIONS

Ceph can replicate or erasure code objects. RAID duplicates this functionality on the block level and reduces available capacity. Consequently, RAID is an unnecessary expense. Additionally, a degraded RAID will have a **negative** impact on performance.



IMPORTANT

Red Hat recommends that each hard drive be exported separately from the RAID controller as a single volume with write-back caching enabled.

This requires a battery-backed, or a non-volatile flash memory device on the storage controller. It is important to make sure the battery is working, as most controllers will disable write-back caching if the memory on the controller can be lost as a result of a power failure. Periodically check the batteries and replace them if necessary, as they do degrade over time. See the storage controller vendor's documentation for details. Typically, the storage controller vendor provides storage management utilities to monitor and adjust the storage controller configuration without any downtime.

Using Just a Bunch of Drives (JBOD) in independent drive mode with Ceph is supported when using all Solid State Drives (SSDs), or for configurations with high numbers of drives per controller. For example, 60 drives attached to one controller. In this scenario, the write-back caching can become a source of I/O contention. Since JBOD disables write-back caching, it is ideal in this scenario. One advantage of using JBOD mode is the ease of adding or replacing drives and then exposing the drive to the operating system immediately after it is physically plugged in.

2.6. SUMMARY OF COMMON MISTAKES WHEN SELECTING HARDWARE

- Repurposing underpowered legacy hardware for use with Ceph.
- Using dissimilar hardware in the same pool.
- Using 1Gbps networks instead of 10Gbps or greater.
- Neglecting to setup both public and cluster networks.
- Using RAID instead of JBOD.
- Selecting drives on a price basis without regard to performance or throughput.
- Journaling on OSD data drives when the use case calls for an SSD journal.
- Having a disk controller with insufficient throughput characteristics.

Use the examples in this document of Red Hat tested configurations for different workloads to avoid some of the foregoing hardware selection mistakes.

Additional Resources

- Supported configurations [article](#) on the Red Hat Customer Portal.

CHAPTER 3. OPTIMIZE WORKLOAD PERFORMANCE DOMAINS

One of the key benefits of Ceph storage is the ability to support different types of workloads within the same cluster using Ceph performance domains. Dramatically different hardware configurations can be associated with each performance domain. Ceph system administrators can deploy storage pools on the appropriate performance domain, providing applications with storage tailored to specific performance and cost profiles. Selecting appropriately sized and optimized servers for these performance domains is an essential aspect of designing a Red Hat Ceph Storage cluster.

The following lists provide the criteria Red Hat uses to identify optimal Red Hat Ceph Storage cluster configurations on storage servers. These categories are provided as general guidelines for hardware purchases and configuration decisions, and can be adjusted to satisfy unique workload blends. Actual hardware configurations chosen will vary depending on specific workload mix and vendor capabilities.

IOPS optimized

An IOPS-optimized storage cluster typically has the following properties:

- Lowest cost per IOPS.
- Highest IOPS per GB.
- 99th percentile latency consistency.

Typically uses for an IOPS-optimized storage cluster are:

- Typically block storage.
- 3x replication for hard disk drives (HDDs) or 2x replication for solid state drives (SSDs).
- MySQL on OpenStack clouds.

Throughput optimized

A throughput-optimized storage cluster typically has the following properties:

- Lowest cost per MBps (throughput).
- Highest MBps per TB.
- Highest MBps per BTU.
- Highest MBps per Watt.
- 97th percentile latency consistency.

Typically uses for an throughput-optimized storage cluster are:

- Block or object storage.
- 3x replication.
- Active performance storage for video, audio, and images.
- Streaming media.

Cost and capacity optimized

A cost- and capacity-optimized storage cluster typically has the following properties:

- Lowest cost per TB.
- Lowest BTU per TB.
- Lowest Watts required per TB.

Typically uses for an cost- and capacity-optimized storage cluster are:

- Typically object storage.
- Erasure coding common for maximizing usable capacity
- Object archive.
- Video, audio, and image object repositories.

How performance domains work

To the Ceph client interface that reads and writes data, a Ceph storage cluster appears as a simple pool where the client stores data. However, the storage cluster performs many complex operations in a manner that is completely transparent to the client interface. Ceph clients and Ceph object storage daemons (Ceph OSDs, or simply OSDs) both use the controlled replication under scalable hashing (CRUSH) algorithm for storage and retrieval of objects. OSDs run on OSD hosts—the storage servers within the cluster.

A CRUSH map describes a topography of cluster resources, and the map exists both on client nodes as well as Ceph Monitor (MON) nodes within the cluster. Ceph clients and Ceph OSDs both use the CRUSH map and the CRUSH algorithm. Ceph clients communicate directly with OSDs, eliminating a centralized object lookup and a potential performance bottleneck. With awareness of the CRUSH map and communication with their peers, OSDs can handle replication, backfilling, and recovery—allowing for dynamic failure recovery.

Ceph uses the CRUSH map to implement failure domains. Ceph also uses the CRUSH map to implement performance domains, which simply take the performance profile of the underlying hardware into consideration. The CRUSH map describes how Ceph stores data, and it is implemented as a simple hierarchy (acyclic graph) and a ruleset. The CRUSH map can support multiple hierarchies to separate one type of hardware performance profile from another.

The following examples describe performance domains.

- Hard disk drives (HDDs) are typically appropriate for cost- and capacity-focused workloads.
- Throughput-sensitive workloads typically use HDDs with Ceph write journals on solid state drives (SSDs).
- IOPS-intensive workloads such as MySQL and MariaDB often use SSDs.

All of these performance domains can coexist in a Ceph storage cluster.

CHAPTER 4. SERVER AND RACK SOLUTIONS

Hardware vendors have responded to the enthusiasm around Ceph by providing both optimized server-level and rack-level solution SKUs. Validated through joint testing with Red Hat, these solutions offer predictable price-to-performance ratios for Ceph deployments, with a convenient modular approach to expand Ceph storage for specific workloads.

Typical rack-level solutions include:

- **Network switching:** Redundant network switching interconnects the cluster and provides access to clients.
- **Ceph MON nodes:** The Ceph monitor is a datastore for the health of the entire cluster, and contains the cluster log. A minimum of three monitor nodes are strongly recommended for a cluster quorum in production.
- **Ceph OSD hosts:** Ceph OSD hosts house the storage capacity for the cluster, with one or more OSDs running per individual storage device. OSD hosts are selected and configured differently depending on both workload optimization and the data devices installed: HDDs, SSDs, or NVMe SSDs.
- **Red Hat Ceph Storage:** Many vendors provide a capacity-based subscription for Red Hat Ceph Storage bundled with both server and rack-level solution SKUs.



NOTE

Red Hat recommends to review the [Red Hat Ceph Storage:Supported Configurations](#) article prior to committing to any server and rack solution. Contact [Red Hat support](#) for any additional assistance.

IOPS-optimized solutions

With the growing use of flash storage, organizations increasingly host IOPS-intensive workloads on Ceph storage clusters to let them emulate high-performance public cloud solutions with private cloud storage. These workloads commonly involve structured data from MySQL-, MariaDB-, or PostgreSQL-based applications.

Typical servers include the following elements:

- **CPU:** 10 cores per NVMe SSD, assuming a 2 GHz CPU.
- **RAM:** 16 GB baseline, plus 5 GB per OSD.
- **Networking:** 10 Gigabit Ethernet (GbE) per 2 OSDs.
- **OSD media:** High-performance, high-endurance enterprise NVMe SSDs.
- **OSDs:** Two per NVMe SSD.
- **Bluestore WAL/DB:** High-performance, high-endurance enterprise NVMe SSD, co-located with OSDs.
- **Controller:** Native PCIe bus.

**NOTE**

For Non-NVMe SSDs, for **CPU**, use two cores per SSD OSD.

Table 4.1. Solutions SKUs for IOPS-optimized Ceph Workloads, by cluster size.

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
SuperMicro ^[a]	SYS-5038MR-OSD006P	N/A	N/A
[a] See Supermicro® Total Solution for Ceph for details.			

Throughput-optimized Solutions

Throughput-optimized Ceph solutions are usually centered around semi-structured or unstructured data. Large-block sequential I/O is typical.

Typical server elements include:

- **CPU:** 0.5 cores per HDD, assuming a 2 GHz CPU.
- **RAM:** 16 GB baseline, plus 5 GB per OSD.
- **Networking:** 10 GbE per 12 OSDs each for client- and cluster-facing networks.
- **OSD media:** 7,200 RPM enterprise HDDs.
- **OSDs:** One per HDD.
- **Bluestore WAL/DB:** High-performance, high-endurance enterprise NVMe SSD, co-located with OSDs.
- **Host bus adapter (HBA):** Just a bunch of disks (JBOD).

Several vendors provide pre-configured server and rack-level solutions for throughput-optimized Ceph workloads. Red Hat has conducted extensive testing and evaluation of servers from Supermicro and Quanta Cloud Technologies (QCT).

Table 4.2. Rack-level SKUs for Ceph OSDs, MONs, and top-of-rack (TOR) switches.

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
SuperMicro ^[a]	SRS-42E112-Ceph-03	SRS-42E136-Ceph-03	SRS-42E136-Ceph-03

Table 4.3. Individual OSD Servers

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
SuperMicro ^[a]	SSG-6028R-OSD072P	SSG-6048-OSD216P	SSG-6048-OSD216P

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
QCT [a]	QxStor RCT-200	QxStor RCT-400	QxStor RCT-400
[a] See QCT: QxStor Red Hat Ceph Storage Edition for details.			

Table 4.4. Additional Servers Configurable for Throughput-optimized Ceph OSD Workloads.

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
Dell	PowerEdge R730XD [a]	DSS 7000 [b], twin node	DSS 7000, twin node
Cisco	UCS C240 M4	UCS C3260 [c]	UCS C3260 [d]
Lenovo	System x3650 M5	System x3650 M5	N/A
[a] See Dell PowerEdge R730xd Performance and Sizing Guide for Red Hat Ceph Storage - A Dell Red Hat Technical White Paper for details. [b] See Dell EMC DSS 7000 Performance & Sizing Guide for Red Hat Ceph Storage for details. [c] See Red Hat Ceph Storage hardware reference architecture for details. [d] See UCS C3260 for details			

Cost and capacity-optimized solutions

Cost- and capacity-optimized solutions typically focus on higher capacity, or longer archival scenarios. Data can be either semi-structured or unstructured. Workloads include media archives, big data analytics archives, and machine image backups. Large-block sequential I/O is typical.

Solutions typically include the following elements:

- **CPU.** 0.5 cores per HDD, assuming a 2 GHz CPU.
- **RAM.** 16 GB baseline, plus 5 GB per OSD.
- **Networking.** 10 GbE per 12 OSDs (each for client- and cluster-facing networks).
- **OSD media.** 7,200 RPM enterprise HDDs.
- **OSDs.** One per HDD.
- **Bluestore WAL/DB** Co-located on the HDD.
- **HBA.** JBOD.

Supermicro and QCT provide pre-configured server and rack-level solution SKUs for cost- and capacity-focused Ceph workloads.

Table 4.5. Pre-configured Rack-level SKUs for Cost- and Capacity-optimized Workloads

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
SuperMicro [a]	N/A	SRS-42E136-Ceph-03	SRS-42E172-Ceph-03

Table 4.6. Pre-configured Server-level SKUs for Cost- and Capacity-optimized Workloads

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
SuperMicro [a]	N/A	SSG-6048R-OSD216P [a]	SSD-6048R-OSD360P
QCT	N/A	QxStor RCC-400 [a]	QxStor RCC-400 [a]
[a] See Supermicro's Total Solution for Ceph for details.			

Table 4.7. Additional Servers Configurable for Cost- and Capacity-optimized Workloads

Vendor	Small (250TB)	Medium (1PB)	Large (2PB+)
Dell	N/A	DSS 7000, twin node	DSS 7000, twin node
Cisco	N/A	UCS C3260	UCS C3260
Lenovo	N/A	System x3650 M5	N/A

Additional Resources

- [Red Hat Ceph Storage on Samsung NVMe SSDs](#)
- [Red Hat Ceph Storage on the InfiniFlash All-Flash Storage System from SanDisk](#)
- [Deploying MySQL Databases on Red Hat Ceph Storage](#)
- [Intel® Data Center Blocks for Cloud – Red Hat OpenStack Platform with Red Hat Ceph Storage](#)
- [Red Hat Ceph Storage on QCT Servers](#)
- [Red Hat Ceph Storage on Servers with Intel Processors and SSDs](#)

CHAPTER 5. MINIMUM HARDWARE RECOMMENDATIONS FOR CONTAINERIZED CEPH

Ceph can run on non-proprietary commodity hardware. Small production clusters and development clusters can run without performance optimization with modest hardware.



IMPORTANT

Hardware accelerated compression in Ceph Object Gateway requires RHEL 9.4 on a Sapphire or Emerald Rapids Xeon CPU (or newer) with QAT devices. For more information, see [Intel Ark](#).

Process	Criteria	Minimum Recommended
ceph-osd-container	Processor	1x AMD64 or Intel 64 CPU CORE per OSD container
	RAM	Minimum of 5 GB of RAM per OSD container
	OS Disk	1x OS disk per host
	OSD Storage	1x storage drive per OSD container. Cannot be shared with OS Disk.
	block.db	Optional, but Red Hat recommended, 1x SSD or NVMe or Optane partition or lvm per daemon. Sizing is 4% of block.data for BlueStore for object, file and mixed workloads and 1% of block.data for the BlueStore for Block Device, Openstack cinder, and Openstack cinder workloads.
	block.wal	Optionally, 1x SSD or NVMe or Optane partition or logical volume per daemon. Use a small size, for example 10 GB, and only if it's faster than the block.db device.
	Network	2x 10 GB Ethernet NICs
ceph-mon-container	Processor	1x AMD64 or Intel 64 CPU CORE per mon-container
	RAM	3 GB per mon-container
	Disk Space	10 GB per mon-container , 50 GB Recommended
	Monitor Disk	Optionally, 1x SSD disk for Monitor rocksdb data
	Network	2x 1GB Ethernet NICs, 10 GB Recommended
ceph-mgr-container	Processor	1x AMD64 or Intel 64 CPU CORE per mgr-container
	RAM	3 GB per mgr-container

Process	Criteria	Minimum Recommended
	Network	2x 1GB Ethernet NICs, 10 GB Recommended
ceph-radosgw-container	Processor	1x AMD64 or Intel 64 CPU CORE per radosgw-container
	RAM	1 GB per daemon
	Disk Space	5 GB per daemon
	Network	1x 1GB Ethernet NICs
ceph-mds-container	Processor	1x AMD64 or Intel 64 CPU CORE per mds-container
	RAM	3 GB per mds-container This number is highly dependent on the configurable MDS cache size. The RAM requirement is typically twice as much as the amount set in the mds_cache_memory_limit configuration setting. Note also that this is the memory for your daemon, not the overall system memory.
	Disk Space	2 GB per mds-container , plus taking into consideration any additional space required for possible debug logging, 20GB is a good start.
	Network	2x 1GB Ethernet NICs, 10 GB Recommended Note that this is the same network as the OSD containers. If you have a 10 GB network on your OSDs you should use the same on your MDS so that the MDS is not disadvantaged when it comes to latency.

CHAPTER 6. RECOMMENDED MINIMUM HARDWARE REQUIREMENTS FOR THE RED HAT CEPH STORAGE DASHBOARD

The Red Hat Ceph Storage Dashboard has minimum hardware requirements.

Minimum requirements

- 4 core processor at 2.5 GHz or higher
- 8 GB RAM
- 50 GB hard disk drive
- 1 Gigabit Ethernet network interface

Additional Resources

- For more information, see [High-level monitoring of a Ceph storage cluster](#) in the [Administration Guide](#).