



Red Hat Ceph Storage 7

Upgrade Guide

Upgrading a Red Hat Ceph Storage Cluster

Red Hat Ceph Storage 7 Upgrade Guide

Upgrading a Red Hat Ceph Storage Cluster

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 instructions on upgrading a Red Hat Ceph Storage cluster running Red Hat Enterprise Linux on AMD64 and Intel 64 architectures. 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. UPGRADE A RED HAT CEPH STORAGE CLUSTER USING CEPHADM	3
1.1. COMPATIBILITY CONSIDERATIONS BETWEEN RHCS AND PODMAN VERSIONS	3
1.2. UPGRADING THE RED HAT CEPH STORAGE CLUSTER	4
1.3. UPGRADING THE RED HAT CEPH STORAGE CLUSTER IN A DISCONNECTED ENVIRONMENT	7
CHAPTER 2. UPGRADING A HOST OPERATING SYSTEM FROM RHEL 8 TO RHEL 9	11
CHAPTER 3. UPGRADING RHCS 5 TO RHCS 7 INVOLVING RHEL 8 TO RHEL 9 UPGRADES WITH STRETCH MODE ENABLED	13
CHAPTER 4. STAGGERED UPGRADE	18
4.1. STAGGERED UPGRADE OPTIONS	18
4.1.1. Performing a staggered upgrade	18
4.1.2. Performing a staggered upgrade from previous releases	21
CHAPTER 5. MONITORING AND MANAGING UPGRADE OF THE STORAGE CLUSTER	24
CHAPTER 6. TROUBLESHOOTING UPGRADE ERROR MESSAGES	26

CHAPTER 1. UPGRADE A RED HAT CEPH STORAGE CLUSTER USING CEPHADM

As a storage administrator, you can use the **cephadm** Orchestrator to Red Hat Ceph Storage 7 with the **ceph orch upgrade** command.



NOTE

Upgrading directly from Red Hat Ceph Storage 5 to Red Hat Ceph Storage 7 is supported.

The automated upgrade process follows Ceph best practices. For example:

- The upgrade order starts with Ceph Managers, Ceph Monitors, then other daemons.
- Each daemon is restarted only after Ceph indicates that the cluster will remain available.

The storage cluster health status is likely to switch to **HEALTH_WARNING** during the upgrade. When the upgrade is complete, the health status should switch back to **HEALTH_OK**.



NOTE

You do not get a message once the upgrade is successful. Run **ceph versions** and **ceph orch ps** commands to verify the new image ID and the version of the storage cluster.

1.1. COMPATIBILITY CONSIDERATIONS BETWEEN RHCS AND PODMAN VERSIONS

podman and Red Hat Ceph Storage have different end-of-life strategies that might make it challenging to find compatible versions.

Red Hat recommends to use the **podman** version shipped with the corresponding Red Hat Enterprise Linux version for Red Hat Ceph Storage. See the [Red Hat Ceph Storage: Supported configurations](#) knowledge base article for more details. See the [Contacting Red Hat support for service](#) section in the *Red Hat Ceph Storage Troubleshooting Guide* for additional assistance.

The following table shows version compatibility between Red Hat Ceph Storage 7 and versions of **podman**.

Ceph	Podman					
	1.9	2.0	2.1	2.2	3.0	>3.0
Red Hat Ceph Storage 7	false	true	true	false	true	true

**WARNING**

You must use a version of Podman that is 2.0.0 or higher.

1.2. UPGRADING THE RED HAT CEPH STORAGE CLUSTER

You can use **ceph orch upgrade** command for upgrading a Red Hat Ceph Storage cluster.

Prerequisites

- Latest version of running Red Hat Ceph Storage cluster.
- Root-level access to all the nodes.
- Ansible user with sudo and passwordless **ssh** access to all nodes in the storage cluster.
- At least two Ceph Manager nodes in the storage cluster: one active and one standby.

NOTE

Red Hat Ceph Storage 5 also includes a health check function that returns a `DAEMON_OLD_VERSION` warning if it detects that any of the daemons in the storage cluster are running multiple versions of Red Hat Ceph Storage. The warning is triggered when the daemons continue to run multiple versions of Red Hat Ceph Storage beyond the time value set in the **mon_warn_older_version_delay** option. By default, the **mon_warn_older_version_delay** option is set to 1 week. This setting allows most upgrades to proceed without falsely seeing the warning. If the upgrade process is paused for an extended time period, you can mute the health warning:

```
ceph health mute DAEMON_OLD_VERSION --sticky
```

After the upgrade has finished, unmute the health warning:

```
ceph health unmute DAEMON_OLD_VERSION
```

Procedure

1. Enable the Ceph Ansible repositories on the Ansible administration node:

Red Hat Enterprise Linux 9

```
subscription-manager repos --enable=rhceph-7-tools-for-rhel-9-x86_64-rpms
```

2. Update the **cephadm** and **cephadm-ansible** package:

Example

```
[root@admin ~]# dnf update cephadm
[root@admin ~]# dnf update cephadm-ansible
```

3. Navigate to the `/usr/share/cephadm-ansible/` directory:

Example

```
[root@admin ~]# cd /usr/share/cephadm-ansible
```

4. Run the preflight playbook with the **upgrade_ceph_packages** parameter set to **true** on the bootstrapped host in the storage cluster:

Syntax

```
ansible-playbook -i INVENTORY_FILE cephadm-preflight.yml --extra-vars "ceph_origin=rhcs
upgrade_ceph_packages=true"
```

Example

```
[ceph-admin@admin cephadm-ansible]$ ansible-playbook -i /etc/ansible/hosts cephadm-
preflight.yml --extra-vars "ceph_origin=rhcs upgrade_ceph_packages=true"
```

This package upgrades **cephadm** on all the nodes.

5. Log into the **cephadm** shell:

Example

```
[root@host01 ~]# cephadm shell
```

6. Ensure all the hosts are online and that the storage cluster is healthy:

Example

```
[ceph: root@host01 /]# ceph -s
```

7. Set the OSD **noout**, **noscrub**, and **nodeep-scrub** flags to prevent OSDs from getting marked out during upgrade and to avoid unnecessary load on the cluster:

Example

```
[ceph: root@host01 /]# ceph osd set noout
[ceph: root@host01 /]# ceph osd set noscrub
[ceph: root@host01 /]# ceph osd set nodeep-scrub
```

8. Check service versions and the available target containers:

Syntax

```
ceph orch upgrade check IMAGE_NAME
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade check registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

**NOTE**

The image name is applicable for both Red Hat Enterprise Linux 8 and Red Hat Enterprise Linux 9.

9. Upgrade the storage cluster:

Syntax

```
ceph orch upgrade start IMAGE_NAME
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade start registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

**NOTE**

To perform a staggered upgrade, see [Performing a staggered upgrade](#).

While the upgrade is underway, a progress bar appears in the **ceph status** output.

Example

```
[ceph: root@host01 /]# ceph status
[...]
progress:
  Upgrade to 18.2.0-128.el9cp (1s)
  [.....]
```

10. Verify the new *IMAGE_ID* and *VERSION* of the Ceph cluster:

Example

```
[ceph: root@host01 /]# ceph versions
[ceph: root@host01 /]# ceph orch ps
```



NOTE

If you are not using the **cephadm-ansible** playbooks, after upgrading your Ceph cluster, you must upgrade the **ceph-common** package and client libraries on your client nodes.

Example

```
[root@client01 ~] dnf update ceph-common
```

Verify you have the latest version:

Example

```
[root@client01 ~] ceph --version
```

11. When the upgrade is complete, unset the **noout**, **noscrub**, and **nodeep-scrub** flags:

Example

```
[ceph: root@host01 /]# ceph osd unset noout
[ceph: root@host01 /]# ceph osd unset noscrub
[ceph: root@host01 /]# ceph osd unset nodeep-scrub
```

1.3. UPGRADING THE RED HAT CEPH STORAGE CLUSTER IN A DISCONNECTED ENVIRONMENT

You can upgrade the storage cluster in a disconnected environment by using the **--image** tag.

You can use **ceph orch upgrade** command for upgrading a Red Hat Ceph Storage cluster.

Prerequisites

- Latest version of running Red Hat Ceph Storage cluster.
- Root-level access to all the nodes.
- Ansible user with sudo and passwordless **ssh** access to all nodes in the storage cluster.
- At least two Ceph Manager nodes in the storage cluster: one active and one standby.
- Register the nodes to CDN and attach subscriptions.
- Check for the customer container images in a disconnected environment and change the configuration, if required. See the [Changing configurations of custom container images for disconnected installations](#) section in the *Red Hat Ceph Storage Installation Guide* for more details.

By default, the monitoring stack components are deployed based on the primary Ceph image. For disconnected environment of the storage cluster, you have to use the latest available monitoring stack component images.

Table 1.1. Custom image details for monitoring stack

Monitoring stack component	Image details for Red Hat Ceph Storage 7.0	Image details for Red Hat Ceph Storage 7.1
Prometheus	registry.redhat.io/openshift4/ose-prometheus:v4.12	registry.redhat.io/openshift4/ose-prometheus:v4.15
Grafana	registry.redhat.io/rhceph/grafana-rhel9:latest	registry.redhat.io/rhceph/grafana-rhel9:latest
Node-exporter	registry.redhat.io/openshift4/ose-prometheus-node-exporter:v4.12	registry.redhat.io/openshift4/ose-prometheus-node-exporter:v4.15
AlertManager	registry.redhat.io/openshift4/ose-prometheus-alertmanager:v4.12	registry.redhat.io/openshift4/ose-prometheus-alertmanager:v4.15
HAProxy	registry.redhat.io/rhceph/rhceph-haproxy-rhel9:latest	registry.redhat.io/rhceph/rhceph-haproxy-rhel9:latest
Keepalived	registry.redhat.io/rhceph/keepalived-rhel9:latest	registry.redhat.io/rhceph/keepalived-rhel9:latest
SNMP Gateway	registry.redhat.io/rhceph/snmp-notifier-rhel9:latest	registry.redhat.io/rhceph/snmp-notifier-rhel9:latest
Loki	registry.redhat.io/openshift-logging/logging-loki-rhel8:v2.6.1	registry.redhat.io/openshift-logging/logging-loki-rhel8:v2.6.1
Promtail	registry.redhat.io/rhceph/rhceph-promtail-rhel9:v2.4.0	registry.redhat.io/rhceph/rhceph-promtail-rhel9:v2.4.0

Find the latest available supported container images on the [Red Hat Ecosystem Catalog](#).

Procedure

1. Enable the Ceph Ansible repositories on the Ansible administration node:

Red Hat Enterprise Linux 9

```
subscription-manager repos --enable=rhceph-7-tools-for-rhel-9-x86_64-rpms
```

2. Update the **cephadm** and **cephadm-ansible** package.

Example

```
[root@admin ~]# dnf update cephadm
[root@admin ~]# dnf update cephadm-ansible
```

3. Navigate to the `/usr/share/cephadm-ansible/` directory:

Example

Example

```
[root@admin ~]# cd /usr/share/cephadm-ansible
```

4. Run the preflight playbook with the **upgrade_ceph_packages** parameter set to **true** and the **ceph_origin** parameter set to **custom** on the bootstrapped host in the storage cluster:

Syntax

```
ansible-playbook -i INVENTORY_FILE cephadm-preflight.yml --extra-vars  
"ceph_origin=custom upgrade_ceph_packages=true"
```

Example

```
[ceph-admin@admin ~]$ ansible-playbook -i /etc/ansible/hosts cephadm-preflight.yml --  
extra-vars "ceph_origin=custom upgrade_ceph_packages=true"
```

This package upgrades **cephadm** on all the nodes.

5. Log into the **cephadm** shell:

Example

```
[root@node0 ~]# cephadm shell
```

6. Ensure all the hosts are online and that the storage cluster is healthy:

Example

```
[ceph: root@node0 /]# ceph -s
```

7. Set the OSD **noout**, **noscrub**, and **nodeep-scrub** flags to prevent OSDs from getting marked out during upgrade and to avoid unnecessary load on the cluster:

Example

```
[ceph: root@host01 /]# ceph osd set noout  
[ceph: root@host01 /]# ceph osd set noscrub  
[ceph: root@host01 /]# ceph osd set nodeep-scrub
```

8. Check service versions and the available target containers:

Syntax

```
ceph orch upgrade check IMAGE_NAME
```

Example

```
[ceph: root@node0 /]# ceph orch upgrade check LOCAL_NODE_FQDN:5000/rhceph/rhceph-  
7-rhel9
```

9. Upgrade the storage cluster:

Syntax

```
ceph orch upgrade start IMAGE_NAME
```

Example

```
[ceph: root@node0 /]# ceph orch upgrade start LOCAL_NODE_FQDN:5000/rhceph/rhceph-7-rhel9
```

While the upgrade is underway, a progress bar appears in the **ceph status** output.

Example

```
[ceph: root@node0 /]# ceph status
[...]
progress:
  Upgrade to 18.2.0-128.el9cp (1s)
  [.....]
```

10. Verify the new *IMAGE_ID* and *VERSION* of the Ceph cluster:

Example

```
[ceph: root@node0 /]# ceph version
[ceph: root@node0 /]# ceph versions
[ceph: root@node0 /]# ceph orch ps
```

11. When the upgrade is complete, unset the **noout**, **noscrub**, and **nodeep-scrub** flags:

Example

```
[ceph: root@host01 /]# ceph osd unset noout
[ceph: root@host01 /]# ceph osd unset noscrub
[ceph: root@host01 /]# ceph osd unset nodeep-scrub
```

Additional Resources

- See the [Registering Red Hat Ceph Storage nodes to the CDN and attaching subscriptions](#) section in the *Red Hat Ceph Storage Installation Guide*.
- See the [Configuring a private registry for a disconnected installation](#) section in the *Red Hat Ceph Storage Installation Guide*.

CHAPTER 2. UPGRADING A HOST OPERATING SYSTEM FROM RHEL 8 TO RHEL 9

You can perform a Red Hat Ceph Storage host operating system upgrade from Red Hat Enterprise Linux 8 to Red Hat Enterprise Linux 9 using the Leapp utility.

Prerequisites

- A running Red Hat Ceph Storage 5.3 cluster.

The following are the supported combinations of containerized Ceph daemons. For more information, see the [How colocation works and its advantages](#) section in the *Red Hat Ceph Storage Installation Guide*.

- Ceph Metadata Server (**ceph-mds**), Ceph OSD (**ceph-osd**), and Ceph Object Gateway (**radosgw**)
- Ceph Monitor (**ceph-mon**) or Ceph Manager (**ceph-mgr**), Ceph OSD (**ceph-osd**), and Ceph Object Gateway (**radosgw**)
- Ceph Monitor (**ceph-mon**), Ceph Manager (**ceph-mgr**), Ceph OSD (**ceph-osd**), and Ceph Object Gateway (**radosgw**)

Procedure

1. Deploy Red Hat Ceph Storage 5.3 on Red Hat Enterprise Linux 8.8 with service.



NOTE

Verify that the cluster contains two admin nodes, so that while performing host upgrade in one *admin* node (with **_admin** label), the second admin can be used for managing clusters.

For full instructions, see [Red Hat Ceph Storage installation](#) in the *Red Hat Ceph Storage Installation Guide* and [Deploying the Ceph daemons using the service specifications](#) in the *Operations guide*.

1. Set the **noout** flag on the Ceph OSD.

Example

```
[ceph: root@host01 /]# ceph osd set noout
```

2. Perform host upgrade one node at a time using the Leapp utility.
 - a. Put respective node maintenance mode before performing host upgrade using Leapp.

Syntax

```
ceph orch host maintenance enter HOST
```

Example

```
ceph orch host maintenance enter host01
```

- b. Enable the ceph tools repo manually when executing the leapp command with the **--enablerepo** parameter.

Example

```
leapp upgrade --enablerepo rhceph-5-tools-for-rhel-9-x86_64-rpms
```

- c. Refer to [Upgrading RHEL 8 to RHEL 9](#) within the Red Hat Enterprise Linux product documentation on the Red Hat Customer Portal.



IMPORTANT

After performing in-place upgrade from Red Hat Enterprise Linux 8 to Red Hat Enterprise Linux 9, you need to manually enable and start the **logrotate.timer** service.

```
# systemctl start logrotate.timer  
# systemctl enable logrotate.timer
```

3. Verify the new *IMAGE_ID* and *VERSION* of the Ceph cluster:

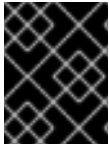
Example

```
[ceph: root@node0 /]# ceph version  
[ceph: root@node0 /]# ceph orch ps
```

4. Continue with the Red Hat Ceph Storage 5.3 to Red Hat Ceph Storage 7 upgrade by following the [Upgrading the Red Hat Ceph Storage cluster](#) steps in the *Red Hat Ceph Storage Upgrade Guide*.

CHAPTER 3. UPGRADING RHCS 5 TO RHCS 7 INVOLVING RHEL 8 TO RHEL 9 UPGRADES WITH STRETCH MODE ENABLED

You can perform an upgrade from Red Hat Ceph Storage 5 to Red Hat Ceph Storage 7 involving Red Hat Enterprise Linux 8 to Red Hat Enterprise Linux 9 with the stretch mode enabled.



IMPORTANT

Upgrade to the latest version of Red Hat Ceph Storage 5 prior to upgrading to the latest version of Red Hat Ceph Storage 7.

Prerequisites

- Red Hat Ceph Storage 5 on Red Hat Enterprise Linux 8 with necessary hosts and daemons running with stretch mode enabled.
- Backup of Ceph binary (**/usr/sbin/cephadm**), **ceph.pub** (**/etc/ceph**), and the Ceph cluster's public SSH keys from the admin node.

Procedure

1. Log into the Cephadm shell:

Example

```
[ceph: root@host01 /]# cephadm shell
```

2. Label a second node as the admin in the cluster to manage the cluster when the admin node is re-provisioned.

Syntax

```
ceph orch host label add HOSTNAME_admin
```

Example

```
[ceph: root@host01 /]# ceph orch host label add host02_admin
```

3. Set the **noout** flag.

Example

```
[ceph: root@host01 /]# ceph osd set noout
```

4. Drain all the daemons from the host:

Syntax

```
ceph orch host drain HOSTNAME --force
```

Example

```
[ceph: root@host01 /]# ceph orch host drain host02 --force
```

The **_no_schedule** label is automatically applied to the host which blocks deployment.

5. Check if all the daemons are removed from the storage cluster:

Syntax

```
ceph orch ps HOSTNAME
```

Example

```
[ceph: root@host01 /]# ceph orch ps host02
```

6. Zap the devices so that if the hosts being drained have OSDs present, then they can be used to re-deploy OSDs when the host is added back.

Syntax

```
ceph orch device zap HOSTNAME DISK --force
```

Example

```
[ceph: root@host01 /]# ceph orch device zap ceph-host02 /dev/vdb --force  
zap successful for /dev/vdb on ceph-host02
```

7. Check the status of OSD removal:

Example

```
[ceph: root@host01 /]# ceph orch osd rm status
```

When no placement groups (PG) are left on the OSD, the OSD is decommissioned and removed from the storage cluster.

8. Remove the host from the cluster:

Syntax

```
ceph orch host rm HOSTNAME --force
```

Example

```
[ceph: root@host01 /]# ceph orch host rm host02 --force
```

9. Re-provision the respective hosts from RHEL 8 to RHEL 9 as described in [Upgrading from RHEL 8 to RHEL 9](#).

10. Run the preflight playbook with the **--limit** option:

Syntax

```
ansible-playbook -i INVENTORY_FILE cephadm-preflight.yml --limit NEWHOST_NAME
```

Example

```
[ceph: root@host01 /]# ansible-playbook -i hosts cephadm-preflight.yml --extra-vars  
"ceph_origin={storage-product}" --limit host02
```

The preflight playbook installs **podman**, **lvm2**, **chronyd**, and **cephadm** on the new host. After installation is complete, **cephadm** resides in the **/usr/sbin/** directory.

11. Extract the cluster's public SSH keys to a folder:

Syntax

```
ceph cephadm get-pub-key ~/PATH
```

Example

```
[ceph: root@host01 /]# ceph cephadm get-pub-key ~/ceph.pub
```

12. Copy Ceph cluster's public SSH keys to the re-provisioned node:

Syntax

```
ssh-copy-id -f -i ~/PATH root@HOST_NAME_2
```

Example

```
[ceph: root@host01 /]# ssh-copy-id -f -i ~/ceph.pub root@host02
```

- a. Optional: If the removed host has a monitor daemon, then, before adding the host to the cluster, add the **--unmanaged** flag to monitor deployment.

Syntax

```
ceph orch apply mon PLACEMENT --unmanaged
```

13. Add the host again to the cluster and add the labels present earlier:

Syntax

```
ceph orch host add HOSTNAME IP_ADDRESS --labels=LABELS
```

- a. Optional: If the removed host had a monitor daemon deployed originally, the monitor daemon needs to be added back manually with the location attributes as described in [Replacing the tiebreaker with a new monitor](#).

Syntax

```
ceph mon add HOSTNAME IP LOCATION
```

Example

```
[ceph: root@host01 /]# ceph mon add ceph-host02 10.0.211.62 datacenter=DC2
```

Syntax

```
ceph orch daemon add mon HOSTNAME
```

Example

```
[ceph: root@host01 /]# ceph orch daemon add mon ceph-host02
```

14. Verify the daemons on the re-provisioned host running successfully with the same ceph version:

Syntax

```
ceph orch ps
```

15. Set back the monitor daemon placement to **managed**.

Syntax

```
ceph orch apply mon PLACEMENT
```

16. Repeat the above steps for all hosts.
 - a. .Arbiter monitor cannot be drained or removed from the host. Hence, the arbiter mon needs to be re-provisioned to another tie-breaker node, and then drained or removed from host as described in [Replacing the tiebreaker with a new monitor](#) .
17. Follow the same approach to re-provision admin nodes and use a second admin node to manage clusters.
18. Add the backup files again to the node.
19. . Add admin nodes again to cluster using the second admin node. Set the **mon** deployment to **unmanaged**.
20. Follow [Replacing the tiebreaker with a new monitor](#) to add back the old arbiter mon and remove the temporary mon created earlier.
21. Unset the **noout** flag.

Syntax

```
ceph osd unset noout
```

22. Verify the Ceph version and the cluster status to ensure that all demons are working as expected after the Red Hat Enterprise Linux upgrade.

23. Follow [Upgrade a Red Hat Ceph Storage cluster using cephadm](#) to perform Red Hat Ceph Storage 5 to Red Hat Ceph Storage 7 Upgrade.

CHAPTER 4. STAGGERED UPGRADE

As a storage administrator, you can upgrade Red Hat Ceph Storage components in phases rather than all at once. The **ceph orch upgrade** command enables you to specify options to limit which daemons are upgraded by a single upgrade command.



NOTE

If you want to upgrade from a version that does not support staggered upgrades, you must first manually upgrade the Ceph Manager (**ceph-mgr**) daemons. For more information on performing a staggered upgrade from previous releases, see [Performing a staggered upgrade from previous releases](#).

4.1. STAGGERED UPGRADE OPTIONS

The **ceph orch upgrade** command supports several options to upgrade cluster components in phases. The staggered upgrade options include:

- **--daemon_types**: The **--daemon_types** option takes a comma-separated list of daemon types and will only upgrade daemons of those types. Valid daemon types for this option include **mgr**, **mon**, **crash**, **osd**, **mds**, **rgw**, **rbd-mirror**, **cephfs-mirror**, and **nfs**.
- **--services**: The **--services** option is mutually exclusive with **--daemon-types**, only takes services of one type at a time, and will only upgrade daemons belonging to those services. For example, you cannot provide an OSD and RGW service simultaneously.
- **--hosts**: You can combine the **--hosts** option with **--daemon_types**, **--services**, or use it on its own. The **--hosts** option parameter follows the same format as the command line options for orchestrator CLI placement specification.
- **--limit**: The **--limit** option takes an integer greater than zero and provides a numerical limit on the number of daemons **cephadm** will upgrade. You can combine the **--limit** option with **--daemon_types**, **--services**, or **--hosts**. For example, if you specify to upgrade daemons of type **osd** on **host01** with a limit set to **3**, **cephadm** will upgrade up to three OSD daemons on host01.

4.1.1. Performing a staggered upgrade

As a storage administrator, you can use the **ceph orch upgrade** options to limit which daemons are upgraded by a single upgrade command.

Cephadm strictly enforces an order for the upgrade of daemons that is still present in staggered upgrade scenarios. The current upgrade order is:

1. Ceph Manager nodes
2. Ceph Monitor nodes
3. Ceph-crash daemons
4. Ceph OSD nodes
5. Ceph Metadata Server (MDS) nodes
6. Ceph Object Gateway (RGW) nodes

7. Ceph RBD-mirror node
8. CephFS-mirror node
9. Ceph NFS nodes



NOTE

If you specify parameters that upgrade daemons out of order, the upgrade command blocks and notes which daemons you need to upgrade before you proceed.

Example

```
[ceph: root@host01 /]# ceph orch upgrade start --image
registry.redhat.io/rhceph/rhceph-7-rhel9:latest --hosts host02
```

Error EINVAL: Cannot start upgrade. Daemons with types earlier in upgrade order than daemons on given host need upgrading.
Please first upgrade mon.ceph-host01



NOTE

There is no required order for restarting the instances. Red Hat recommends restarting the instance pointing to the pool with primary images followed by the instance pointing to the mirrored pool.

Prerequisites

- A cluster running Red Hat Ceph Storage 5.3 or 6.1.
- Root-level access to all the nodes.
- At least two Ceph Manager nodes in the storage cluster: one active and one standby.

Procedure

1. Log into the **cephadm** shell:

Example

```
[root@host01 ~]# cephadm shell
```

2. Ensure all the hosts are online and that the storage cluster is healthy:

Example

```
[ceph: root@host01 /]# ceph -s
```

3. Set the OSD **noout**, **noscrub**, and **nodeep-scrub** flags to prevent OSDs from getting marked out during upgrade and to avoid unnecessary load on the cluster:

Example

```
[ceph: root@host01 /]# ceph osd set noout
[ceph: root@host01 /]# ceph osd set noscrub
[ceph: root@host01 /]# ceph osd set nodeep-scrub
```

4. Check service versions and the available target containers:

Syntax

```
ceph orch upgrade check IMAGE_NAME
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade check registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

5. Upgrade the storage cluster:

- a. To upgrade specific daemon types on specific hosts:

Syntax

```
ceph orch upgrade start --image IMAGE_NAME --daemon-types
DAEMON_TYPE1,DAEMON_TYPE2 --hosts HOST1,HOST2
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade start --image registry.redhat.io/rhceph/rhceph-7-rhel9:latest --daemon-types mgr,mon --hosts host02,host03
```

- b. To specify specific services and limit the number of daemons to upgrade:

Syntax

```
ceph orch upgrade start --image IMAGE_NAME --services SERVICE1,SERVICE2 --limit
LIMIT_NUMBER
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade start --image registry.redhat.io/rhceph/rhceph-7-rhel9:latest --services rgw.example1,rgw1.example2 --limit 2
```



NOTE

In staggered upgrade scenarios, if using a limiting parameter, the monitoring stack daemons, including Prometheus and **node-exporter**, are refreshed after the upgrade of the Ceph Manager daemons. As a result of the limiting parameter, Ceph Manager upgrades take longer to complete. The versions of monitoring stack daemons might not change between Ceph releases, in which case, they are only redeployed.



NOTE

Upgrade commands with limiting parameters validates the options before beginning the upgrade, which can require pulling the new container image. As a result, the **upgrade start** command might take a while to return when you provide limiting parameters.

- To see which daemons you still need to upgrade, run the **ceph orch upgrade check** or **ceph versions** command:

Example

```
[ceph: root@host01 /]# ceph orch upgrade check --image registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

- To complete the staggered upgrade, verify the upgrade of all remaining services:

Syntax

```
ceph orch upgrade start --image IMAGE_NAME
```

Example

```
[ceph: root@host01 /]# ceph orch upgrade start --image registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

Verification

- Verify the new *IMAGE_ID* and *VERSION* of the Ceph cluster:

Example

```
[ceph: root@host01 /]# ceph versions
[ceph: root@host01 /]# ceph orch ps
```

- When the upgrade is complete, unset the **noout**, **noscrub**, and **nodeep-scrub** flags:

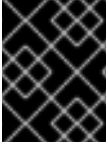
Example

```
[ceph: root@host01 /]# ceph osd unset noout
[ceph: root@host01 /]# ceph osd unset noscrub
[ceph: root@host01 /]# ceph osd unset nodeep-scrub
```

4.1.2. Performing a staggered upgrade from previous releases

You can perform a staggered upgrade on your storage cluster by providing the necessary arguments

You can perform a staggered upgrade on your storage cluster by providing the necessary arguments. If you want to upgrade from a version that does not support staggered upgrades, you must first manually upgrade the Ceph Manager (**ceph-mgr**) daemons. Once you have upgraded the Ceph Manager daemons, you can pass the limiting parameters to complete the staggered upgrade.

**IMPORTANT**

Verify you have at least two running Ceph Manager daemons before attempting this procedure.

Prerequisites

- A cluster running Red Hat Ceph Storage 5.2 or lesser.
- At least two Ceph Manager nodes in the storage cluster: one active and one standby.

Procedure

1. Log into the Cephadm shell:

Example

```
[root@host01 ~]# cephadm shell
```

2. Determine which Ceph Manager is active and which are standby:

Example

```
[ceph: root@host01 /]# ceph -s
cluster:
  id: 266ee7a8-2a05-11eb-b846-5254002d4916
  health: HEALTH_OK

services:
  mon: 2 daemons, quorum host01,host02 (age 92s)
  mgr: host01.ndtpjh(active, since 16h), standbys: host02.pzgrhz
```

3. Manually upgrade each standby Ceph Manager daemon:

Syntax

```
ceph orch daemon redeploy mgr.ceph-HOST.MANAGER_ID --image IMAGE_ID
```

Example

```
[ceph: root@host01 /]# ceph orch daemon redeploy mgr.ceph-host02.pzgrhz --image
registry.redhat.io/rhceph/rhceph-7-rhel9:latest
```

4. Fail over to the upgraded standby Ceph Manager:

Example

```
[ceph: root@host01 /]# ceph mgr fail
```

5. Check that the standby Ceph Manager is now active:

Example

```
[ceph: root@host01 /]# ceph -s
cluster:
  id: 266ee7a8-2a05-11eb-b846-5254002d4916
  health: HEALTH_OK

services:
  mon: 2 daemons, quorum host01,host02 (age 1h)
  mgr: host02.pzgrhz(active, since 25s), standbys: host01.ndtpjh
```

6. Verify that the active Ceph Manager is upgraded to the new version:

Syntax

```
ceph tell mgr.ceph-HOST.MANAGER_ID version
```

Example

```
[ceph: root@host01 /]# ceph tell mgr.host02.pzgrhz version
{
  "version": "18.2.0-128.el8cp",
  "release": "reef",
  "release_type": "stable"
}
```

7. Repeat steps 2 - 6 to upgrade the remaining Ceph Managers to the new version.
8. Check that all Ceph Managers are upgraded to the new version:

Example

```
[ceph: root@host01 /]# ceph mgr versions
{
  "ceph version 18.2.0-128.el8cp (600e227816517e2da53d85f2fab3cd40a7483372) pacific
  (stable)": 2
}
```

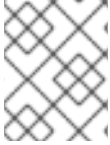
9. Once you upgrade all your Ceph Managers, you can specify the limiting parameters and complete the remainder of the staggered upgrade.

Additional Resources

- For more information about performing a staggered upgrade and staggered upgrade options, see [Performing a staggered upgrade](#).

CHAPTER 5. MONITORING AND MANAGING UPGRADE OF THE STORAGE CLUSTER

After running the **ceph orch upgrade start** command to upgrade the Red Hat Ceph Storage cluster, you can check the status, pause, resume, or stop the upgrade process. The health of the cluster changes to **HEALTH_WARNING** during an upgrade. If the host of the cluster is offline, the upgrade is paused.



NOTE

You have to upgrade one daemon type after the other. If a daemon cannot be upgraded, the upgrade is paused.

Prerequisites

- A running Red Hat Ceph Storage cluster.
- Root-level access to all the nodes.
- At least two Ceph Manager nodes in the storage cluster: one active and one standby.
- Upgrade for the storage cluster initiated.

Procedure

1. Determine whether an upgrade is in process and the version to which the cluster is upgrading:

Example

```
[ceph: root@node0 /]# ceph orch upgrade status
```



NOTE

You do not get a message once the upgrade is successful. Run **ceph versions** and **ceph orch ps** commands to verify the new image ID and the version of the storage cluster.

2. Optional: Pause the upgrade process:

Example

```
[ceph: root@node0 /]# ceph orch upgrade pause
```

3. Optional: Resume a paused upgrade process:

Example

```
[ceph: root@node0 /]# ceph orch upgrade resume
```

4. Optional: Stop the upgrade process:

Example

```
[ceph: root@node0 /]# ceph orch upgrade stop
```

CHAPTER 6. TROUBLESHOOTING UPGRADE ERROR MESSAGES

The following table shows some **cephadm** upgrade error messages. If the **cephadm** upgrade fails for any reason, an error message appears in the storage cluster health status.

Error Message	Description
UPGRADE_NO_STANDBY_MGR	Ceph requires both active and standby manager daemons to proceed, but there is currently no standby.
UPGRADE_FAILED_PULL	Ceph was unable to pull the container image for the target version. This can happen if you specify a version or container image that does not exist (e.g., 1.2.3), or if the container registry is not reachable from one or more hosts in the cluster.