



Red Hat Ceph Storage 5.1

Release Notes

Release notes for Red Hat Ceph Storage 5.1z2

Red Hat Ceph Storage 5.1 Release Notes

Release notes for Red Hat Ceph Storage 5.1z2

Legal Notice

Copyright © 2022 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

The release notes describe the major features, enhancements, known issues, and bug fixes implemented for the Red Hat Ceph Storage 5 product release. This includes previous release notes of the Red Hat Ceph Storage 5.1 release up to the current release.

Table of Contents

MAKING OPEN SOURCE MORE INCLUSIVE	3
PROVIDING FEEDBACK ON RED HAT CEPH STORAGE DOCUMENTATION	4
CHAPTER 1. INTRODUCTION	5
CHAPTER 2. ACKNOWLEDGMENTS	6
CHAPTER 3. NEW FEATURES	7
3.1. THE CEPHADM UTILITY	8
3.2. CEPH DASHBOARD	10
3.3. CEPH MANAGER PLUGINS	11
3.4. THE CEPH VOLUME UTILITY	11
3.5. CEPH OBJECT GATEWAY	11
3.6. MULTI-SITE CEPH OBJECT GATEWAY	11
3.7. RADOS	12
CHAPTER 4. TECHNOLOGY PREVIEWS	13
CHAPTER 5. DEPRECATED FUNCTIONALITY	14
CHAPTER 6. BUG FIXES	15
6.1. THE CEPHADM UTILITY	15
6.2. CEPH FILE SYSTEM	15
6.3. CEPH OBJECT GATEWAY	16
6.4. RADOS	16
6.5. RBD MIRRORING	16
6.6. THE CEPH ANSIBLE UTILITY	17
CHAPTER 7. KNOWN ISSUES	18
7.1. CEPH OBJECT GATEWAY	18
CHAPTER 8. DOCUMENTATION	19
8.1. THE UPGRADE GUIDE	19
CHAPTER 9. SOURCES	20

MAKING OPEN SOURCE MORE INCLUSIVE

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](#).

PROVIDING FEEDBACK ON RED HAT CEPH STORAGE DOCUMENTATION

We appreciate your input on our documentation. Please let us know how we could make it better. To do so:

- For simple comments on specific passages, make sure you are viewing the documentation in the multi-page HTML format. Highlight the part of text that you want to comment on. Then, click the **Add Feedback** pop-up that appears below the highlighted text, and follow the displayed instructions.
- For submitting more complex feedback, create a Bugzilla ticket:
 1. Go to the [Bugzilla](#) website.
 2. In the Component drop-down, select **Documentation**
 3. Select the appropriate version of the document.
 4. Fill in the **Summary** and **Description** field with your suggestion for improvement. Include a link to the relevant part(s) of documentation.
 5. Optional: Add an attachment, if any.
 6. Click **Submit Bug**.

CHAPTER 1. INTRODUCTION

Red Hat Ceph Storage is a massively scalable, open, software-defined storage platform that combines the most stable version of the Ceph storage system with a Ceph management platform, deployment utilities, and support services.

The Red Hat Ceph Storage documentation is available at
<https://access.redhat.com/documentation/en/red-hat-ceph-storage/>.

CHAPTER 2. ACKNOWLEDGMENTS

Red Hat Ceph Storage 5 project is seeing amazing growth in the quality and quantity of contributions from individuals and organizations in the Ceph community. We would like to thank all members of the Red Hat Ceph Storage team, all of the individual contributors in the Ceph community, and additionally, but not limited to, the contributions from organizations such as:

- Intel[®]
- Fujitsu[®]
- UnitedStack
- Yahoo[™]
- Ubuntu Kylin
- Mellanox[®]
- CERN[™]
- Deutsche Telekom
- Mirantis[®]
- SanDisk[™]
- SUSE

CHAPTER 3. NEW FEATURES

This section lists all major updates, enhancements, and new features introduced in this release of Red Hat Ceph Storage.

The main features added by this release are:

- **Containerized Cluster**

Red Hat Ceph Storage 5 supports only containerized daemons. It does not support non-containerized storage clusters. If you are upgrading a non-containerized storage cluster from Red Hat Ceph Storage 4 to Red Hat Ceph Storage 5, the upgrade process includes the conversion to a containerized deployment.

For more information, see the [Upgrading a Red Hat Ceph Storage cluster from RHCS 4 to RHCS 5](#) section in the *Red Hat Ceph Storage Installation Guide* for more details.

- **Cephadm**

Cephadm is a new containerized deployment tool that deploys and manages a Red Hat Ceph Storage 5 cluster by connecting to hosts from the manager daemon. The **cephadm** utility replaces **ceph-ansible** for Red Hat Ceph Storage deployment. The goal of Cephadm is to provide a fully-featured, robust, and well installed management layer for running Red Hat Ceph Storage.

The **cephadm** command manages the full lifecycle of a Red Hat Ceph Storage cluster.

The **cephadm** command can perform the following operations:

- Bootstrap a new Ceph storage cluster.
- Launch a containerized shell that works with the Ceph command-line interface (CLI).
- Aid in debugging containerized daemons.

The **cephadm** command uses **ssh** to communicate with the nodes in the storage cluster and add, remove, or update Ceph daemon containers. This allows you to add, remove, or update Red Hat Ceph Storage containers without using external tools.

The **cephadm** command has two main components:

- The **cephadm** shell launches a **bash** shell within a container. This enables you to run storage cluster installation and setup tasks, as well as to run **ceph** commands in the container.
- The **cephadm** orchestrator commands enable you to provision Ceph daemons and services, and to expand the storage cluster.

For more information, see the [Red Hat Ceph Storage Installation Guide](#).

- **Management API**

The management API creates management scripts that are applicable for Red Hat Ceph Storage 5 and continues to operate unchanged for the version lifecycle. The incompatible versioning of the API would only happen across major release lines.

For more information, see the [Red Hat Ceph Storage Developer Guide](#).

- **Disconnected installation of Red Hat Ceph Storage**

Red Hat Ceph Storage 5 supports the disconnected installation and bootstrapping of storage clusters on private networks. A disconnected installation uses custom images and configuration files and local hosts, instead of downloading files from the network.

You can install container images that you have downloaded from a proxy host that has access to the Red Hat registry, or by copying a container image to your local registry. The bootstrapping process requires a specification file that identifies the hosts to be added by name and IP address. Once the initial monitor host has been bootstrapped, you can use Ceph Orchestrator commands to expand and configure the storage cluster.

See the [Red Hat Ceph Storage Installation Guide](#) for more details.

- **Ceph File System geo-replication**

Starting with the Red Hat Ceph Storage 5 release, you can replicate Ceph File Systems (CephFS) across geographical locations or between different sites. The new **cephfs-mirror** daemon does asynchronous replication of snapshots to a remote CephFS.

See the [Ceph File System mirrors](#) section in the *Red Hat Ceph Storage File System Guide* for more details.

- **A new Ceph File System client performance tool**

Starting with the Red Hat Ceph Storage 5 release, the Ceph File System (CephFS) provides a **top**-like utility to display metrics on Ceph File Systems in realtime. The **cephfs-top** utility is a **curses**-based Python script that uses the Ceph Manager **stats** module to fetch and display client performance metrics.

See the [Using the cephfs-top utility](#) section in the *Red Hat Ceph Storage File System Guide* for more details.

- **Monitoring the Ceph object gateway multisite using the Red Hat Ceph Storage Dashboard**

The Red Hat Ceph Storage dashboard can now be used to monitor an Ceph object gateway multisite configuration.

After the multi-zones are set-up using the **cephadm** utility, the buckets of one zone is visible to other zones and other sites. You can also create, edit, delete buckets on the dashboard.

See the [Management of buckets of a multisite object configuration on the Ceph dashboard](#) chapter in the *Red Hat Ceph Storage Dashboard Guide* for more details.

- **Improved BlueStore space utilization**

The Ceph Object Gateway and the Ceph file system (CephFS) stores small objects and files as individual objects in RADOS. With this release, the default value of BlueStore's **min_alloc_size** for SSDs and HDDs is 4 KB. This enables better use of space with no impact on performance.

See the [OSD BlueStore](#) chapter in the *Red Hat Ceph Storage Administration Guide* for more details.

3.1. THE CEPHADM UTILITY

cephadm supports colocating multiple daemons on the same host

With this release, multiple daemons, such as Ceph Object Gateway and Ceph Metadata Server (MDS), can be deployed on the same host thereby providing an additional performance benefit.

Example

```
service_type: rgw
placement:
  label: rgw
  count-per-host: 2
```

For single node deployments, **cephadm** requires at least two running Ceph Manager daemons in upgrade scenarios. It is still highly recommended even outside of upgrade scenarios but the storage cluster will function without it.

Configuration of NFS-RGW using Cephadm is now supported

In Red Hat Ceph Storage 5.0 configuration, use of NFS-RGW required the use of dashboard as a workaround and it was recommended for such users to delay upgrade until Red Hat Ceph Storage 5.1

With this release, NFS-RGW configuration is supported and the users with this configuration can upgrade their storage cluster and it works as expected.

Users can now bootstrap their storage clusters with custom monitoring stack images

Previously, users had to adjust the image used for their monitoring stack daemons manually after bootstrapping the cluster

With this release, you can specify custom images for monitoring stack daemons during bootstrap by passing a configuration file formatted as follows:

Syntax

```
[mgr]
mgr/cephadm/container_image_grafana = GRAFANA_IMAGE_NAME
mgr/cephadm/container_image_alertmanager = ALERTMANAGER_IMAGE_NAME
mgr/cephadm/container_image_prometheus = PROMETHEUS_IMAGE_NAME
mgr/cephadm/container_image_node_exporter = NODE_EXPORTER_IMAGE_NAME
```

You can run bootstrap with the **--config CONFIGURATION_FILE_NAME** option in the command. If you have other configuration options, you can simply add the lines above in your configuration file before bootstrapping the storage cluster.

cephadm enables automated adjustment of **osd_memory_target**

With this release, **cephadm** enables automated adjustment of **osd_memory_target** configuration parameter by default.

Users can now specify CPU limits for the daemons by service

With this release, you can customize the CPU limits for all daemons within any given service by adding the CPU limit to the service specification file via the **extra_container_args** field.

Example

```
service_type: mon
service_name: mon
placement:
hosts:
- host01
- host02
- host03
extra_container_args:
- "--cpus=2"

service_type: osd
service_id: osd_example
placement:
```

```
hosts:
  - host01
extra_container_args:
  - "--cpus=2"
spec:
  data_devices:
    paths:
      - /dev/sdb
```

cephadm now supports IPv6 networks for Ceph Object Gateway deployment

With this release, **cephadm** supports specifying an IPv6 network for Ceph Object Gateway specifications. An example of a service configuration file for deploying Ceph Object Gateway is:

Example

```
service_type: rgw
service_id: rgw
placement:
  count: 3
networks:
  - fd00:fd00:3000::/64
```

The `ceph nfs export create rgw` command now supports exporting Ceph Object Gateway users

Previously, the **ceph nfs export create rgw** command would only create Ceph Object Gateway exports at the bucket level.

With this release, the command creates the Ceph Object Gateway exports at both the user and bucket level.

Syntax

```
ceph nfs export create rgw --cluster-id CLUSTER_ID --pseudo-path PSEUDO_PATH --user-id USER_ID [--readonly] [--client_addr VALUE...] [--squash VALUE]
```

Example

```
[ceph: root@host01 /]# ceph nfs export create rgw --cluster-id mynfs --pseudo-path /bucketdata --user-id myuser --client_addr 192.168.10.0/24
```

3.2. CEPH DASHBOARD

Users can now view the HAProxy metrics on the Red Hat Ceph Storage Dashboard

With this release, Red Hat introduces the new Grafana dashboard for ingress service used for Ceph Object Gateway endpoints. You can now view the four HAProxy metrics under Ceph Object Gateway Daemons Overall Performance such as Total responses by HTTP code, Total requests/responses, Total number of connections, and Current total of incoming/outgoing bytes.

User can view `mfa_ids` on the Red Hat Ceph Storage Dashboard

With this release, you can view the **mfa_ids** for a user configured with multi-factor authentication (MFA) for the Ceph Object Gateway user in the *User Details* section on the Red Hat Ceph Storage Dashboard.

3.3. CEPH MANAGER PLUGINS

The global recovery event in the progress module is now optimized

With this release, computing for the progress of global recovery events is optimized for a large number of placement groups in a large storage cluster by using the C++ code instead of the python module thereby reducing the CPU utilization.

3.4. THE CEPH VOLUME UTILITY

The lvm commands does not cause metadata corruption when run within the containers

Previously, when the **lvm** commands were run directly within the containers, it would cause LVM metadata corruption.

With this release, **ceph-volume** uses the host namespace to run the **lvm** commands and avoids metadata corruption.

3.5. CEPH OBJECT GATEWAY

Lock contention messages from the Ceph Object Gateway reshard queue are marked as informational

Previously, when the Ceph Object Gateway failed to get a lock on a reshard queue, the output log entry would appear to be an error causing concern to customers.

With this release, the entries in the output log appear as informational and are tagged as "INFO:".

The support for OIDC JWT token validation using modulus and exponent is available

With this release, OIDC JSON web token (JWT) validation supports the usage of modulus and exponent for signature calculation. It also extends the support for available methods to validate OIDC JWT validation.

The role name and role session fields are now available in ops log for temporary credentials

Previously, the role name and role session were not available and it was difficult for the administrator to know which role was being assumed and which session was active for the temporary credentials being used.

With this release, role name and role session to ops log are available for temporary credentials, returned by AssumeRole* APIs, to perform S3 operations.

Users can now use the --bucket argument to process bucket lifecycles

With this release, you can provide a **--bucket=BUCKET_NAME** argument to the **radosgw-admin lc process** command to process the lifecycle for the corresponding bucket. This is convenient for debugging lifecycle problems that affect specific buckets and for backfilling lifecycle processing for specific buckets that have fallen behind.

3.6. MULTI-SITE CEPH OBJECT GATEWAY

Multi-site configuration supports dynamic bucket index resharding

Previously, only manual resharding of the buckets for multi-site configurations was supported.

With this release, dynamic bucket resharding is supported in multi-site configurations. Once the storage clusters are upgraded, enable the **resharding** feature and reshard the buckets either manually with **radogw-admin bucket reshard** command or automatically with dynamic resharding, independently of other zones in the storage cluster.

3.7. RADOS

Use the **noautoscale** flag to manage the PG autoscaler

With this release, **pg_autoscaler** can be turned **on** or **off** globally using the **noautoscale** flag. This flag is set to **off** by default. When this flag is set, then all the pools have **pg_autoscale_mode** as **off**

For more information, see the [Manually updating autoscale profile](#) section in the *Red Hat Ceph Storage Storage Strategies Guide*.

Users can now create pools with the **--bulk** flag

With this release, you can create pools with the **--bulk** flag. It uses a profile of the **pg_autoscaler** and provides better performance from the start and has a full complement of placement groups (PGs) and scales down only when the usage ratio across the pool is not even.

If the pool does not have the **--bulk** flag, the pool starts out with minimal PGs.

To create a pool with the bulk flag:

Syntax

```
ceph osd pool create POOL_NAME --bulk
```

To set/unset bulk flag of existing pool:

Syntax

```
ceph osd pool set POOL_NAME bulk TRUE/FALSE/1/0  
ceph osd pool unset POOL_NAME bulk TRUE/FALSE/1/0
```

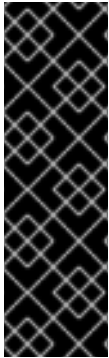
To get bulk flag of existing pool:

Syntax

```
ceph osd pool get POOL_NAME --bulk
```

CHAPTER 4. TECHNOLOGY PREVIEWS

This section provides an overview of Technology Preview features introduced or updated in this release of Red Hat Ceph Storage.



IMPORTANT

Technology Preview features are not supported with Red Hat production service level agreements (SLAs), might not be functionally complete, and Red Hat does not recommend using them for production. These features provide early access to upcoming product features, enabling customers to test functionality and provide feedback during the development process.

For more information on Red Hat Technology Preview features support scope, see the [Technology Preview Features Support Scope](#).

cephadm bootstrapping process supports Red Hat Enterprise Linux 9

With this release, the Red Hat Enterprise Linux 8 based storage cluster **cephadm** bootstrapping process supports Red Hat Enterprise Linux 9. The containers are Red Hat Enterprise Linux 8 that requires Red Hat Enterprise Linux 9 server.

Ceph object gateway in multi-site replication setup now supports a subset of AWS bucket replication API functionality

With this release, Ceph Object Gateway now supports a subset of AWS bucket replication API functionality including {Put, Get, Delete} Replication operations. This feature enables bucket-granularity replication and additionally provides end-user replication control with the caveat that currently, buckets can be replicated within zones in an existing Ceph Object Gateway multi-site replication setup.

SQLite virtual file system (VFS) on top of RADOS is now available

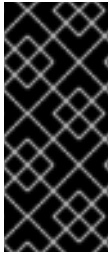
With this release, the new **libcephsqlite** RADOS client library provides SQLite virtual file system (VFS) on top of RADOS. The database and journals are striped over RADOS across multiple objects for virtually unlimited scaling and throughput only limited by the SQLite client. Applications using SQLite might change to the Ceph VFS with minimal changes, usually just by specifying the alternate VFS.

ingress flag to deploy NFS Ganesha daemons

With this release, the **--ingress** flag can be used to deploy one or more NFS Ganesha daemons. The **ingress** service provides load balancing and high-availability for the NFS servers.

CHAPTER 5. DEPRECATED FUNCTIONALITY

This section provides an overview of functionality that has been deprecated in all minor releases up to this release of Red Hat Ceph Storage.



IMPORTANT

Deprecated functionality continues to be supported until the end of life of Red Hat Ceph Storage 5. Deprecated functionality will likely not be supported in future major releases of this product and is not recommended for new deployments. For the most recent list of deprecated functionality within a particular major release, refer to the latest version of release documentation.

NFS support for CephFS is now deprecated

NFS support for CephFS is now deprecated in favor of upcoming NFS availability in OpenShift Data Foundation. Red Hat Ceph Storage support for NFS in OpenStack Manila is not affected. Deprecated functionality will receive only bug fixes for the lifetime of the current release, and may be removed in future releases. Relevant documentation around this technology is identified as "Limited Availability".

iSCSi support is now deprecated

iSCSi support is now deprecated in favor of future NVMeoF support. Deprecated functionality will receive only bug fixes for the lifetime of the current release, and may be removed in future releases. Relevant documentation around this technology is identified as "Limited Availability".

Ceph configuration file is now deprecated

The Ceph configuration file (**ceph.conf**) is now deprecated in favor of new centralized configuration stored in Ceph Monitors. For details, see the [The Ceph configuration database](#) section in the *Red Hat Ceph Storage Configuration Guide*.

The **min_compat_client** parameter for Ceph File System (CephFS) is now deprecated

The **min_compat_client** parameter is deprecated for Red Hat Ceph Storage 5.0 and new client features are added for setting-up the Ceph File Systems (CephFS). For details, see the [Client features](#) section in the *Red Hat Ceph Storage File System Guide*.

The snapshot of Ceph File System subvolume group is now deprecated

The snapshot feature of Ceph File System (CephFS) subvolume group is deprecated for Red Hat Ceph Storage 5.0. The existing snapshots can be listed and deleted, whenever needed. For details, see the [Listing snapshots of a file system subvolume group](#) and [Removing snapshots of a file system subvolume group](#) sections in the *Red Hat Ceph Storage Ceph File System guide*.

The Cockpit Ceph Installer is now deprecated

Installing a Red Hat Ceph Storage cluster 5 using Cockpit Ceph Installer is not supported. Use Cephadm to install a Red Hat Ceph Storage cluster. For details, see the [Red Hat Ceph Storage Installation guide](#).

CHAPTER 6. BUG FIXES

This section describes bugs with significant user impact, which were fixed in this release of Red Hat Ceph Storage. In addition, the section includes descriptions of fixed known issues found in previous versions.

6.1. THE CEPHADM UTILITY

.rgw.root pool is not automatically created during Ceph Object Gateway multisite check

Previously, a check for Ceph Object Gateway multisite was performed by **cephadm** to help with signalling regressions in the release causing a **.rgw.root** pool to be created and recreated if deleted, resulting in users being stuck with the **.rgw.root** pool even when not using Ceph Object Gateway.

With this fix, the check that created the pool is removed, and the pool is no longer created. Users who already have the pool in their system but do not want it, can now delete it without the issue of its recreation. Users who do not have this pool would not have the pool automatically created for them.

([BZ#2090395](#))

6.2. CEPH FILE SYSTEM

Reintegrating stray entries does not fail when the destination directory is full

Previously, Ceph Metadata Servers reintegrated an unlinked file if it had a reference to it, that is, if the deleted file had hard links or was a part of a snapshot. The reintegration, which is essentially an internal rename operation, failed if the destination directory was full. This resulted in Ceph Metadata Server failing to reintegrate stray or deleted entries.

With this release, full space checks during reintegrating stray entries are ignored and these stray entries are reintegrated even when the destination directory is full.

([BZ#2041572](#))

MDS daemons no longer crash when receiving metrics from new clients

Previously, in certain scenarios, newer clients were being used for old CephFS clusters. While upgrading old CephFS, **cephadm** or **mgr** used newer clients to perform checks, tests or configurations with an old Ceph cluster. Due to this, the MDS daemons crashed when receiving unknown metrics from newer clients.

With this fix, the **libceph** clients send only those metrics that are supported by MDS daemons to MDS as default. An additional option to force enable all the metrics when users think it's safe is also added.

([BZ#2081914](#))

Ceph Metadata Server no longer crashes during concurrent lookup and unlink operations

Previously, an incorrect assumption of an assert placed in the code, which gets hit on concurrent lookup and unlink operations from a Ceph client, caused Ceph Metadata Server crash.

The latest fix moves the assertion to the relevant place where the assumption, during concurrent lookup and unlink operation, is valid, resulting in the continuation of Ceph Metadata Server serving the Ceph client operations without crashing.

([BZ#2093064](#))

6.3. CEPH OBJECT GATEWAY

Usage of MD5 for non-cryptographic purposes in a FIPS environment is allowed

Previously, in a FIPS enabled environment, the usage of MD5 digest was not allowed by default, unless explicitly excluded for non-cryptographic purposes. Due to this, a segfault occurred during the S3 complete multipart upload operation.

With this fix, the usage of MD5 for non-cryptographic purposes in a FIPS environment for S3 complete multipart **PUT** operations is explicitly allowed and the S3 multipart operations can be completed.

([BZ#2088601](#))

6.4. RADOS

ceph-objectstore-tool command allows manual trimming of the accumulated PG log dups entries

Previously, trimming of PG log dups entries was prevented during the low-level PG split operation which is used by the PG autoscaler with far higher frequency than by a human operator. Stalling the trimming of dups resulted in significant memory growth of PG log, leading to OSD crashes as it ran out of memory. Restarting an OSD did not solve the problem as the PG log is stored on disk and reloaded to RAM on startup.

With this fix, the **ceph-objectstore-tool** command allows manual trimming of the accumulated PG log dups entries, to unblock the automatic trimming machinery. A debug improvement is implemented that prints the number of dups entries to the OSD's log to help future investigations.

([BZ#2094069](#))

6.5. RBD MIRRORING

Snapshot-based mirroring process no longer gets cancelled

Previously, as a result of an internal race condition, the **rbd mirror snapshot schedule add** command would be cancelled out. The snapshot-based mirroring process for the affected image would not start, if no other existing schedules were applicable.

With this release, the race condition is fixed and the snapshot-based mirroring process starts as expected.

([BZ#2099799](#))

Existing schedules take effect when an image is promoted to primary

Previously, due to an ill-considered optimization, existing schedules would not take effect following an image's promotion to primary resulting in the snapshot-based mirroring process to not start for a recently promoted image.

With this release, the optimization causing this issue is removed and the existing schedules now take effect when an image is promoted to primary and the snapshot-based mirroring process starts as expected.

([BZ#2100519](#))

rbd-mirror daemon no longer acquires exclusive lock

Previously, due to a logic error, **rbd-mirror** daemon could acquire the exclusive lock on a de-facto primary image. Due to this, the snapshot-based mirroring process for the affected image would stop, reporting a "failed to unlink local peer from remote image" error.

With this release, the logic error is fixed resulting in the **rbd-mirror** daemon not acquiring the exclusive lock on a de-facto primary image and the snapshot-based mirroring process not stopping and working as expected.

([BZ#2100520](#))

Mirror snapshot queue used by **rbd-mirror** is extended and is no longer removed

Previously, as a result of an internal race condition, the mirror snapshot used by the **rbd-mirror** daemon on the secondary cluster would be removed causing the snapshot-based mirroring process for the affected image to stop, reporting a "split-brain" error.

With this release, the mirror snapshot queue is extended in length and the mirror snapshot cleanup procedure is amended accordingly, fixing the automatic removal of the mirror snapshots that are still in-use by **rbd-mirror** daemon on the secondary cluster and the snapshot-based mirroring process does not stop.

([BZ#2092843](#))

6.6. THE CEPH ANSIBLE UTILITY

Adoption playbook can now install **cephadm** on OSD nodes

Previously, due to the tools repository being disabled on OSD nodes, you could not install **cephadm** on OSD nodes resulting in the failure of the adoption playbook.

With this fix, the tools repository is enabled on OSD nodes and the adoption playbook can now install **cephadm** on OSD nodes.

([BZ#2073480](#))

Removal of legacy directory ensures error-free cluster post adoption

Previously, **cephadm** displayed an unexpected behaviour with its 'config inferring' function whenever a legacy directory, such as **/var/lib/ceph/mon** was found. Due to this behaviour, post adoption, the cluster was left with the following error "CEPHADM_REFRESH_FAILED: failed to probe daemons or devices".

With this release, the adoption playbook ensures the removal of this directory and the cluster is not left in an error state after the adoption.

([BZ#2075510](#))

CHAPTER 7. KNOWN ISSUES

This section documents known issues found in this release of Red Hat Ceph Storage.

7.1. CEPH OBJECT GATEWAY

Use the **yes-i-know** flag to bootstrap the Red Hat Ceph Storage cluster 5.1

Users are warned against configuring multi-site in a Red Hat Ceph Storage cluster 5.1.

To work around this issue, you must pass the **yes-i-know** flag during bootstrapping.

Syntax

```
sudo --image IMAGE_NAME bootstrap --mon-ip IP_ADDRESS --yes-i-know
```

Users are warned while upgrading to Red Hat Ceph Storage 5.1

Previously, buckets resharded in Red Hat Ceph Storage 5.0, might not have been understandable by a Red Hat Ceph Storage 5.1 Ceph Object Gateway daemon, due to which an upgrade warning or blocker was added to make sure all users upgrading to Red Hat Ceph Storage 5.1 are aware of the issue and can downgrade if they were previously using Red Hat Ceph Storage 5.1.

As a workaround, run **ceph config set mgr mgr/cephadm/yes_i_know true --force** command to remove the warning or blocker and return all operations to normal. Users have to acknowledge that they are aware of the Ceph Object Gateway issue before they upgrade to Red Hat Ceph Storage 5.1.

If the flag is not passed, you get an error message about the multi-site regression and the link to the knowledge base article on the issue.

([BZ#2111679](#))

CHAPTER 8. DOCUMENTATION

This section lists documentation enhancements in this release of Red Hat Ceph Storage.

8.1. THE UPGRADE GUIDE

- **Upgrade instructions are now available in a separate guide.**

Upgrade instructions are no longer delivered within the *Installation Guide* and are now available as a separate guide. See the [Red Hat Ceph Storage Upgrade Guide](#) for details regarding how to upgrade a Red Hat Ceph Storage cluster.

CHAPTER 9. SOURCES

The updated Red Hat Ceph Storage source code packages are available at the following location:

- For Red Hat Enterprise Linux 8:
<http://ftp.redhat.com/redhat/linux/enterprise/8Base/en/RHCEPH/SRPMS/>