



Red Hat Ceph Storage 5.2

릴리스 노트

Red Hat Ceph Storage 5.2 릴리스 노트

Red Hat Ceph Storage 5.2 릴리스 노트

Red Hat Ceph Storage 5.2 릴리스 노트

Enter your first name here. Enter your surname here.

Enter your organisation's name here. Enter your organisational division here.

Enter your email address here.

법적 공지

Copyright © 2022 | You need to change the HOLDER entity in the en-US/Release_Notes.ent file |.

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.

초록

릴리스 노트에서는 Red Hat Ceph Storage 5 제품 릴리스에 대해 구현된 주요 기능, 개선 사항, 알려진 문제 및 버그 수정을 설명합니다. 여기에는 Red Hat Ceph Storage 5.2 릴리스의 이전 릴리스 노트가 포함되어 있습니다.

차례

보다 포괄적 수용을 위한 오픈 소스 용어 교체	3
RED HAT CEPH STORAGE 문서에 대한 피드백 제공	4
1장. 소개	5
2장. 감사 인사	6
3장. 새로운 기능	7
3.1. CEPHADM 유틸리티	7
3.2. CEPH 대시보드	8
3.3. CEPH 파일 시스템	11
3.4. CEPH MANAGER 플러그인	13
3.5. CEPH VOLUME 유틸리티	13
3.6. CEPH OBJECT GATEWAY	14
3.7. 다중 사이트 CEPH 개체 게이트웨이	14
3.8. 패키지	15
3.9. RADOS	15
3.10. CEPH ANSIBLE 유틸리티	16
4장. 기술 프리뷰	17
4.1. CEPH OBJECT GATEWAY	17
4.2. RADOS 블록 장치	17
5장. 사용되지 않는 기능	18
6장. 버그 수정	19
6.1. CEPHADM 유틸리티	19
6.2. CEPH 대시보드	25
6.3. CEPH 파일 시스템	26
6.4. CEPH MANAGER 플러그인	30
6.5. CEPH VOLUME 유틸리티	31
6.6. CEPH OBJECT GATEWAY	31
6.7. 다중 사이트 CEPH 개체 게이트웨이	33
6.8. RADOS	35
6.9. RBD 미러링	36
6.10. CEPH ANSIBLE 유틸리티	38
7장. 확인된 문제	39
7.1. CEPHADM 유틸리티	39
7.2. CEPH 대시보드	39
7.3. CEPH 파일 시스템	40
7.4. CEPH OBJECT GATEWAY	40
8장. 소스	42

보다 포괄적 수용을 위한 오픈 소스 용어 교체

Red Hat은 코드, 문서, 웹 속성에서 문제가 있는 용어를 교체하기 위해 최선을 다하고 있습니다. 먼저 마스터(master), 슬레이브(slave), 블랙리스트(blacklist), 화이트리스트(whitelist) 등 네 가지 용어를 교체하고 있습니다. 이러한 변경 작업은 작업 범위가 크므로 향후 여러 릴리스에 걸쳐 점차 구현할 예정입니다. 자세한 내용은 [CTO Chris Wright의 메시지](#)를 참조하십시오.

RED HAT CEPH STORAGE 문서에 대한 피드백 제공

문서 개선을 위한 의견을 보내 주십시오. Red Hat이 어떻게 이를 개선할 수 있는지 알려 주십시오. 이렇게 하려면 다음을 수행합니다.

- 특정 문구에 대한 간단한 설명을 보려면 문서를 다중 페이지 HTML 형식으로 보고 있는지 확인하십시오. 주석하려는 텍스트 부분을 강조 표시합니다. 그런 다음 강조 표시된 텍스트 아래에 표시되는 **피드백 추가** 팝업을 클릭하고 표시된 지침을 따릅니다.
- 보다 상세하게 피드백을 제출하려면 다음과 같이 Bugzilla 티켓을 생성하십시오.
 1. [Bugzilla](#) 웹 사이트로 이동하십시오.
 2. 구성 요소 드롭다운에서 **문서**를 선택합니다.
 3. 문서의 적절한 버전을 선택합니다.
 4. **Summary** (요약) 및 **Description** (설명) 필드에 개선을 위한 제안 사항을 입력합니다. 관련된 문서의 해당 부분 링크를 알려주십시오.
 5. 선택 사항: 첨부 파일(있는 경우)을 추가합니다.
 6. **버그 제출**을 클릭합니다.

1장. 소개

Red Hat Ceph Storage는 대규모 확장이 가능한 오픈 소프트웨어 정의 스토리지 플랫폼으로서 Ceph 스토리지 시스템의 가장 안정적인 버전을 Ceph 관리 플랫폼, 배포 유틸리티 및 지원 서비스와 결합합니다.

Red Hat Ceph Storage 문서는 <https://access.redhat.com/documentation/en/red-hat-ceph-storage/>에서 확인할 수 있습니다.

2장. 감사 인사

Red Hat Ceph Storage 5 프로젝트는 Ceph 커뮤니티에 있는 개인 및 조직의 기여량과 품질에 놀라운 성장을 실현하고 있습니다. Red Hat Ceph Storage 팀의 모든 구성원, Ceph 커뮤니티의 개별 기여자에 대해 감사드리며 다음과 같은 조직의 기여에 국한되지는 않습니다.

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

3장. 새로운 기능

이 섹션에는 이 Red Hat Ceph Storage 릴리스에 도입된 주요 업데이트, 개선 사항 및 새로운 기능이 나열되어 있습니다.

3.1. CEPHADM 유틸리티

Cephadm-Ansible 모듈

Cephadm-Ansible 패키지는 Ansible을 사용하여 전체 데이터 센터를 관리하려는 사용자를 위해 새로운 통합 컨트롤 플레인인 **cephadm**를 래핑하는 여러 모듈을 제공합니다. **Ceph-Ansible**와의 이전 버전과의 호환성을 제공하지는 않지만 고객이 Ansible 통합을 업데이트하는 데 사용할 수 있는 지원되는 플레이북 세트를 제공하는 것을 목표로 합니다.

자세한 내용은 [cephadm-ansible 모듈](#)을 참조하십시오.

부트스트랩 Red Hat Ceph Storage 클러스터가 Red Hat Enterprise Linux 9에서 지원됨

이번 릴리스에서는 Red Hat Enterprise Linux 9 호스트에서 **cephadm** 부트스트랩을 사용하여 Red Hat Enterprise Linux 9에 대해 Red Hat Ceph Storage 5.2 지원을 사용할 수 있습니다. 이제 사용자가 Red Hat Enterprise Linux 9 호스트에 Ceph 클러스터를 부트스트랩할 수 있습니다.

cephadm rm-cluster 명령은 호스트에서 이전 **systemd** 장치 파일을 정리합니다.

이전 버전에서는 **rm-cluster** 명령이 **systemd** 장치 파일을 제거하지 않고 데몬을 제거했습니다.

이번 릴리스에서는 **cephadm rm-cluster** 명령과 데몬 제거와 함께 호스트에서 이전 **systemd** 장치 파일을 정리합니다.

cephadm에서 사양을 적용하지 못하는 경우 상태 경고를 발생

이전에는 사양을 적용하지 못한 경우 사용자가 종종 확인하지 않는 서비스 이벤트로만 보고되었습니다.

이번 릴리스에서는 **iscsi** 사양의 잘못된 풀 이름과 같은 사양을 적용하지 못하는 경우 **cephadm**에서 상태 경고를 생성하여 사용자에게 경고합니다.

Red Hat Ceph Storage 5.2에서 staggered 업그레이드 지원

Red Hat Ceph Storage 5.2부터 **cephadm**에서 대용량 Ceph 클러스터를 여러 작은 단계에서 선택적으로 업그레이드할 수 있습니다.

ceph orch upgrade start 명령은 다음 매개변수를 허용합니다.

- **--daemon-types**
- **--hosts**
- **--services**
- **--limit**

이러한 매개변수는 제공된 값과 일치하는 데몬을 선택적으로 업그레이드합니다.



참고

이러한 매개변수는 **cephadm**이 지원되는 순서로 데몬을 업그레이드하도록 하는 경우 거부됩니다.



참고

활성 Ceph Manager 데몬이 Red Hat Ceph Storage 5.2 빌드에 있는 경우 이러한 업그레이드 매개변수가 허용됩니다. 이전 버전에서 Red Hat Ceph Storage 5.2로 업그레이드해도 이러한 매개변수는 지원되지 않습니다.

OSD 가 있는 호스트에서 **FS.aio-max-nr** 가 1048576로 설정되어 있습니다.

이전에는 **Cephadm** 에서 관리하는 호스트에서 **fs.aio-max-nr** 를 기본값 **65536** 으로 남겨 두면 일부 OSD 가 충돌할 수 있었습니다.

이번 릴리스에서는 OSD 및 OSD가 있는 호스트에서 **fs.aio-max-nr** 가 1048576으로 설정되어 **fs.aio-max-nr** 매개변수 값이 너무 낮기 때문에 더 이상 충돌하지 않습니다.

Ceph orch rm <service-name> 명령은 사용자가 제거하려는 서비스가 존재하는지 여부를 사용자에게 알립니다.

이전에는 서비스를 제거하면 존재하지 않는 서비스에서 사용자에게 혼란을 초래하는 경우에도 항상 성공적인 메시지가 반환되었습니다.

이번 릴리스에서는 **ceph orch rm SERVICE_NAME** 명령을 실행하면 사용자가 해당 서비스가 존재하지 않거나 **cephadm** 에 없는 서비스가 존재하는지 사용자에게 알립니다.

이제 **cephadm -ansible** 에서 **Resharding** 프로시저를 위한 새로운 **Playbook Enablesdb-resharding.yml**을 사용할 수 있습니다.

이전에는 Rip **sdb Res** harding 절차가 번거롭고 번거롭습니다.

이번 릴리스에서는 **cephadm-ansible** 플레이북인 rolls **db-resharding.yml** 이 구현되어 프로세스를 쉽게 수행할 수 있습니다.

cephadm 에서 LVM 계층 없이 OSD 배포 지원

이 릴리스에서는 원시 OSD에 대해 LVM 계층을 사용하지 않으려는 사용자를 지원하기 위해 **cephadm** 또는 **ceph-volume** 지원이 제공됩니다. **Cephadm**에 전달된 OSD 사양 파일에 "method: raw"를 포함하여 LVM 계층 없이 **Cephadm**을 통해 OSD를 원시 모드로 배포할 수 있습니다.

이번 릴리스에서는 **cephadm**에서 OSD 사양 yaml 파일에서 raw를 사용하여 LVM 계층이 없는 원시 모드로 OSD를 배포할 수 있습니다.

자세한 내용은 [특정 장치 및 호스트에 Ceph OSD 배포](#)를 참조하십시오.

3.2. CEPH 대시보드

기본 서비스 데몬에서 시작, 중지, 재시작 및 재배포 작업을 수행할 수 있습니다.

이전 버전에서는 **orchestrator** 서비스를 생성, 편집 및 삭제할 수 있었습니다. 서비스의 기본 데몬에서 작업을 수행할 수 없음

이번 릴리스에서는 오케스트레이터 서비스의 기본 데몬에서 시작, 중지, 재시작 및 재배포와 같은 작업을 수행할 수 있습니다.

Ceph 대시보드의 **OSD** 페이지 및 출시 페이지에는 **OSD**의 사용 표시줄에 다른 색상이 표시됩니다.

이전에는 **OSD**가 전체 또는 전체 상태에 도달할 때마다 클러스터 상태가 **WARN** 또는 **ERROR** 상태로 변경되었지만 랜딩 페이지에서 다른 실패 표시가 없었습니다.

이번 릴리스에서는 **OSD**가 전체 비율 또는 전체로 전환되면 특정 **OSD**의 **OSD** 페이지 및 랜딩 페이지에 다른 색상이 표시됩니다.

ash hit 또는 **miss** 카운터에 대시보드가 표시됩니다.

이번 릴리스에서는 대시보드가 **Bluestore** 통계에서 가져온 세부 정보를 제공하여 **on massive hit** 또는 **miss** 카운터를 표시하여 **OSD**당 **RAM**을 늘리는 것이 클러스터 성능을 향상시키는 데 도움이 될 수 있습니다.

사용자는 특정 데몬의 **CPU** 및 메모리 사용량을 볼 수 있습니다.

이번 릴리스에서는 클러스터 > 호스트> 데몬의 **Red Hat Ceph Storage** 대시보드의 특정 데몬의 **CPU** 및 메모리 사용량을 확인할 수 있습니다.

rbd-mirroring의 **Ceph** 대시보드 기능 개선

이번 릴리스에서는 **CLI**(명령줄 인터페이스)에만 제공된 다음 기능을 사용하여 **Ceph** 대시보드의 **RBD** 미러링 탭이 향상되었습니다.

- 이미지에서 미러링 활성화 또는 비활성화를 지원합니다.
- 승격 및 시연 작업을 지원합니다.
- 이미지 동기화를 지원합니다.
- 사이트 이름 편집에 대한 가시성을 개선하고 부트스트랩 키를 생성합니다.
- 존재하지 않는 경우 **rbd-mirror**를 자동으로 생성하도록 버튼을 구성하는 빈 페이지가 나타납니다.

이제 사용자가 **Red Hat Ceph Storage** 대시보드에서 단순하고 고급 모드로 **OSD**를 생성할 수 있습니다.

이번 릴리스에서는 간단한 배포 시나리오, "**Simple**" 및 "**Advanced**" 모드인 클러스터의 **OSD** 배포를 간소화하기 위해 **OSD Creation**이 도입되었습니다.

이제 다음 세 가지 옵션 중에서 선택할 수 있습니다.

- **비용/Capacity-optimized:** 사용 가능한 모든 **HD**를 사용하여 **OSD**를 배포할 수 있습니다.
- **처리량이 최적화됨:** **DB/WAL** 장치의 데이터 장치 및 **SSD**에 대해 **HD**가 지원됩니다.
- **IOPS-optimized:** 사용 가능한 모든 **NVME**가 데이터 장치로 사용됩니다.

자세한 내용은 [Ceph Orchestrator를 사용한 OSD](#) 관리를 참조하십시오.

Ceph 대시보드 로그인 페이지가 사용자 지정 가능 텍스트 표시

기업 사용자는 시스템에 액세스하는 모든 사람이 승인되도록 하고 법률/보안 조건을 준수하기 위해 최선을 다하고 있습니다.

이번 릴리스에서는 사용자 지정 배너 또는 경고 텍스트를 표시할 수 있도록 **Ceph** 대시보드 로그인 페이지에 자리 표시자가 제공됩니다. **Ceph** 대시보드 관리자는 다음 명령을 사용하여 배너를 설정, 편집 또는 삭제할 수 있습니다.

예제

```
[ceph: root@host01 /]# ceph dashboard set-login-banner -i filename.yaml
[ceph: root@host01 /]# ceph dashboard get-login-banner
[ceph: root@host01 /]# ceph dashboard unset-login-banner
```

활성화된 경우 대시보드 로그인 페이지에 사용자 지정 가능한 텍스트가 표시됩니다.

Ceph 대시 보드에 주요 버전 번호 및 내부 Ceph 버전이 표시됩니다.

이번 릴리스에서는 주요 버전 번호와 함께 사용자가 Red Hat Ceph Storage 다운스트림 릴리스를 Ceph 내부 버전과 관련시킬 수 있도록 내부 Ceph 버전이 Ceph Dashboard에 표시됩니다. 예를 들어 **Version: 16.2.9-98-gccaadd**. 상단 탐색 표시줄을 클릭하고 물음표 메뉴(?)를 클릭하고 **About modal** 상자로 이동하여 Red Hat Ceph Storage 릴리스 번호와 해당 Ceph 버전을 확인합니다.

3.3. CEPH 파일 시스템

ODF에서 새 기능을 외부 모드로 구성한 CephFS 하위 볼륨에 사용할 수 있습니다.

ODF의 CephFS가 외부 모드로 구성된 경우 사용자는 **volume/subvolume** 메타데이터를 사용하여 **volumes/subvolumes**의 **PVC/PV/namespace**와 같은 일부 Openshift 특정 메타데이터 정보를 저장합니다.

이번 릴리스에서는 CephFS 하위 볼륨에서 사용자 정의 메타데이터를 설정, 가져오기, 업데이트, 나열 및 제거하는 다음과 같은 기능이 추가되었습니다.

다음을 사용하여 하위 볼륨에서 키-값 쌍으로 사용자 정의 메타데이터를 설정합니다.

구문

```
ceph fs subvolume metadata set VOLUME_NAME SUBVOLUME_NAME KEY_NAME VALUE [--group-name SUBVOLUME_GROUP_NAME]
```

metadata 키를 사용하여 하위 볼륨에 설정된 사용자 정의 메타데이터를 가져옵니다.

구문

```
ceph fs subvolume metadata get VOLUME_NAME SUBVOLUME_NAME KEY_NAME [--group-name SUBVOLUME_GROUP_NAME]
```

하위 볼륨에 설정된 사용자 정의 메타데이터, 키-값 쌍을 나열합니다.

구문

```
ceph fs subvolume metadata ls VOLUME_NAME SUBVOLUME_NAME [--group-name  
SUBVOLUME_GROUP_NAME]
```

metadata 키를 사용하여 하위 볼륨에서 설정된 사용자 정의 메타데이터를 제거합니다.

구문

```
ceph fs subvolume metadata rm VOLUME_NAME SUBVOLUME_NAME KEY_NAME [--group-name  
SUBVOLUME_GROUP_NAME] [--force]
```

복제 상태 명령을 사용할 때 복제 실패의 이유가 표시됩니다.

이전 버전에서는 복제가 실패할 때마다 실패 이유를 확인하는 유일한 방법은 로그를 확인하는 것입니다.

이번 릴리스에서는 복제본 상태 명령 출력에 복제 실패 이유가 표시됩니다.

예제

```
[ceph: root@host01 /]# ceph fs clone status cephfs clone1  
{  
  "status": {  
    "state": "failed",  
    "source": {
```

```

    "volume": "cephfs",
    "subvolume": "subvol1",
    "snapshot": "snap1"
    "size": "104857600"
  },
  "failure": {
    "errno": "122",
    "errstr": "Disk quota exceeded"
  }
}
}

```

복제 실패 이유는 다음 두 가지 필드에 표시됩니다.

- **errno** : 오류 번호
- **error_msg** : 실패 오류 문자열

3.4. CEPH MANAGER 플러그인

ceph nfs 내보내기 **apply** 명령을 사용하여 **CephFS NFS** 내보내기를 동적으로 업데이트할 수 있습니다.

이전에는 **CephFS NFS** 내보내기를 업데이트할 때 **NFS-Ganesha** 서버가 항상 다시 시작되었습니다. 이로 인해 업데이트되지 않은 내보내기를 포함하여 **ganesha** 서버에서 제공하는 모든 클라이언트 연결에 일시적으로 영향을 미쳤습니다.

이번 릴리스에서는 **ceph nfs export** 명령을 사용하여 **CephFS NFS** 내보내기를 동적으로 업데이트할 수 있습니다. **CephFS NFS** 내보내기가 업데이트될 때마다 **NFS** 서버가 더 이상 재시작되지 않습니다.

3.5. CEPH VOLUME 유틸리티

사용자는 **OSD**를 재배포하기 전에 장치를 수동으로 지우지 않아도 됩니다.

이전에는 **OSD**를 재배포하기 전에 사용자가 장치를 수동으로 제거해야 했습니다.

이번 릴리스에서는 볼륨 그룹 또는 논리 볼륨이 남아 있지 않을 때 장치의 물리 볼륨이 제거되므로 **OSD**를 재배포하기 전에 사용자가 더 이상 장치를 수동으로 지우지 않아도 됩니다.

3.6. CEPH OBJECT GATEWAY

이제 **Ops Log**를 일반 **Unix** 파일로 보내도록 **Ceph Object Gateway**를 구성할 수 있습니다.

이번 릴리스에서는 **Unix** 도메인 소켓과 비교했을 때 파일 기반 로그가 일부 사이트에서 작업하는 것이 더 간단하고 파일 기반 로그이므로 **Ops Log**를 일반 **Unix** 파일로 전달하도록 **Ceph** 개체 게이트웨이를 구성할 수 있습니다. 로그 파일의 내용은 기본 구성에서 **Ops Log** 소켓으로 전송된 것과 동일합니다.

radosgw lc process 명령을 사용하여 단일 버킷의 라이프사이클을 처리

이번 릴리스에서는 단일 버킷의 라이프 사이클을 처리하는 것이 편리하므로 **--bucket** 또는 **ID --bucket -id**를 지정하여 명령 줄 인터페이스에서 하나의 버킷 라이프사이클만 처리하도록 **radosgw-admin lc process** 명령을 사용할 수 있습니다.

사용자 **ID** 정보가 **Ceph Object Gateway Ops Log** 출력에 추가됩니다.

이번 릴리스에서는 사용자 **ID** 정보가 **Ops Log** 출력에 추가되어 고객이 **S3** 액세스 감사를 위해 이 정보에 액세스할 수 있습니다. **Ceph Object Gateway Ops Log**의 모든 버전에서 **S3** 요청에서 사용자 **ID**를 안정적으로 추적할 수 있습니다.

Ceph Object Gateway의 **HTTP** 액세스 로깅의 로그 수준은 **debug_rgw_access** 매개변수를 사용하여 독립적으로 제어할 수 있습니다.

이번 릴리스에서는 **debug_rgw_access** 매개변수와 **Ceph Object Gateway**의 **HTTP** 액세스 로깅의 로그 수준을 제어하여 사용자에게 이러한 **HTTP** 액세스 로그 행을 제외하고 **debug_rgw=0**과 같은 모든 **Ceph Object Gateway** 로깅을 비활성화하는 기능을 제공합니다.

버킷 인덱스를 업데이트할 때 20개의 **Ceph Object Gateway** 로그 메시지가 감소

이번 릴리스에서는 버킷 인덱스를 업데이트하여 값을 추가하고 로그 크기를 줄이기 위해 **Ceph Object Gateway** 수준 20개의 로그 메시지가 감소합니다.

3.7. 다중 사이트 CEPH 개체 게이트웨이

current_time 필드가 여러 **radosgw-admin** 명령의 출력에 추가되었습니다.

이번 릴리스에서는 현재_time 필드가 여러 **radosgw-admin** 명령, 특히 동기화 상태, 버킷 동기화 상태, 메타데이터 동기화 상태, 데이터 동기화 상태 및 **BI log** 상태에 추가됩니다.

HTTP 클라이언트 로깅

이전에는 **Ceph Object Gateway**가 **HTTP** 응답의 오류 본문을 출력하지 않았으며 요청과 일치시킬 방법이 없었습니다.

이번 릴리스에서는 비동기 **HTTP** 클라이언트 및 오류 본문에 대한 **HTTP** 요청과 일치하도록 태그를 유지 관리하여 보다 철저하게 로깅됩니다. **Ceph Object Gateway debug**가 20으로 설정되면 오류 본문 및 기타 세부 정보가 출력됩니다.

이제 **OpenStack Keystone**에 대한 읽기 전용 역할을 사용할 수 있습니다.

OpenStack Keystone 서비스에서는 **admin, member, reader** 의 세 가지 역할을 제공합니다. 이제 역할 기반 액세스 제어(RBAC) 기능을 **OpenStack**으로 확장하려면 새 읽기 전용 **Admin** 역할을 **Keystone** 서비스의 특정 사용자에게 할당할 수 있습니다.

RBAC 지원 범위는 **OpenStack** 릴리스를 기반으로 합니다.

3.8. 패키지

grafana 컨테이너의 새 버전은 보안 수정 및 개선된 기능을 제공합니다.

이번 릴리스에서는 **grafana v8.3.5** 를 사용한 새로운 버전의 **grafana** 컨테이너가 빌드되어 보안 수정 및 향상된 기능을 제공합니다.

3.9. RADOS

pg_autoscale_mode 가 on으로 설정된 경우 **MANY_OBJECTS_PER_PG** 경고가 더 이상 보고되지 않습니다.

이전 버전에서는 **Ceph** 상태 경고 **MANY_OBECTS_PER_PG** 가 상태 경고를 보고하는 다양한 모드를 구분하지 않고 **pg_autoscale_mode** 가 로 설정된 인스턴스에서 보고되었습니다.

이번 릴리스에서는 **pg_autoscale_mode** 가 .으로 설정된 경우 **MANY_OBJECTS_PER_PG** 경고를 보고하지 않으려면 검사가 추가되었습니다.

OSD는 **Ceph Manager** 서비스에 집계된 형식의 느린 작업 세부 정보를 보고합니다.

이전 버전에서는 느린 요청이 너무 많은 세부 정보로 클러스터 로그를 덮어 쓰기하여 모니터 데이터베이스를 채웠었습니다.

이번 릴리스에서는 작업 유형 및 풀 정보를 통해 느린 요청이 클러스터 로그에 기록되며 **OSD**에서 느린 작업 세부 정보를 관리자 서비스에 보고합니다.

이제 사용자가 **CIDR** 범위를 차단 목록에 지정할 수 있습니다.

이번 릴리스에서는 개별 클라이언트 인스턴스 및 IP 외에도 **CIDR** 범위를 차단 목록에 지정할 수 있습니다. 특정 상황에서는 차단 목록에 개별 클라이언트를 지정하는 대신 전체 데이터 센터 또는 랙의 모든 클라이언트를 차단해야 합니다. 예를 들어 워크로드를 다른 머신 세트로 장애 조치하고 이전 워크로드 인스턴스가 부분적으로 작동하지 않도록 합니다. 이제 기존 "**blocklist**" 명령과 유사한 "**blocklist 범위**"를 사용할 수 있습니다.

3.10. CEPH ANSIBLE 유틸리티

이제 **Ceph** 파일을 백업하고 복원하는 새로운 **Ansible Playbook**을 사용할 수 있습니다.

이전에는 OS를 **Red Hat Enterprise Linux 7**에서 **Red Hat Enterprise Linux 8**로 업그레이드하거나 시스템을 재프로비저닝할 때 사용자가 수동으로 백업 및 복원해야 했습니다. 이는 대규모 클러스터 배포의 경우 특히 불편했습니다.

이번 릴리스에서는 백업을 수행하고 **Ceph** 파일(예: **/etc/ceph** 및 **/var/lib/ceph**)에 **backup-and-restore -ceph** 플레이북이 추가되어 사용자가 파일을 수동으로 복원할 필요가 없습니다.

4장. 기술 프리뷰

이 섹션에서는 이 **Red Hat Ceph Storage** 릴리스에서 도입되거나 업데이트된 기술 프리뷰 기능에 대해 설명합니다.

중요

기술 프리뷰 기능은 **Red Hat** 프로덕션 서비스 수준 계약(SLA)에서 지원되지 않으며 기능적으로 완전하지 않을 수 있으므로 프로덕션에 사용하지 않는 것이 좋습니다. 이러한 기능을 사용하면 향후 제품 기능을 조기에 이용할 수 있어 개발 과정에서 고객이 기능을 테스트하고 피드백을 제공할 수 있습니다.

Red Hat 기술 프리뷰 기능 지원 범위에 대한 자세한 내용은 [기술 프리뷰 기능 지원범위](#)를 참조하십시오.

NFS 배포의 가용성 향상을 위해 HA 지원 NFS

이번 릴리스에서는 **NFS**가 **NFS** 배포의 가용성을 개선하기 위해 **HA**에 의해 지원됩니다. **haproxy**에서 지원하는 **NFS** 및 **keepalived**를 배포할 수 있습니다. 배치에서 더 많은 호스트를 지정하지만 **count** 속성과 함께 사용되는 호스트 수를 제한하면 **NFS** 호스트가 오프라인될 때 **NFS** 데몬이 다른 호스트에 배포됩니다.

자세한 내용은 [Ceph 오케스트레이터를 사용한 NFS Ganesha 게이트웨이](#) 관리를 참조하십시오.

4.1. CEPH OBJECT GATEWAY

S3 투명 암호화를 위한 Ceph Object Gateway 기술 프리뷰 지원

이 릴리스에서는 **Ceph Object Gateway**가 **SSE-S3** 및 **S3 PutBucketEncryption API**를 사용하여 **S3** 투명한 암호화를 지원합니다.

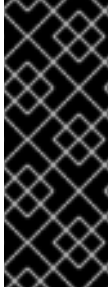
4.2. RADOS 블록 장치

대기 시간을 줄이기 위해 영구 쓰기 로그 캐시라는 **librbd** 플러그인

이번 릴리스에서는 **PWL(Persistent Write Log Cache)**이라는 새로운 **librbd** 플러그인으로 **SSD** 장치를 대상으로 하는 지속적인 내결함성 쓰기 캐시를 제공합니다. 대기 시간을 크게 줄이고 **io_depths**에서 성능을 향상시킵니다. 이 캐시는 내부 검사점을 유지 관리하는 로그 순서의 쓰기를 사용하여 내부적으로 검사점을 유지하므로 클러스터에 다시 플러시되는 쓰기는 항상 일관성을 유지합니다. 클라이언트 캐시가 완전히 손실되더라도 디스크 이미지는 계속 일관되게 유지되지만 데이터는 오래된 것처럼 보입니다.

5장. 사용되지 않는 기능

이 섹션에서는 이 **Red Hat Ceph Storage** 릴리스에 대한 모든 마이너 릴리스에서 더 이상 사용되지 않는 기능에 대해 설명합니다.



중요

더 이상 사용되지 않는 기능은 **Red Hat Ceph Storage 5**의 라이프 사이클이 종료될 때까지 계속 지원됩니다. 사용되지 않는 기능은 이 제품의 향후 주요 릴리스에서 지원되지 않을 가능성이 높으며 새로운 배포에 구현하는 것은 권장되지 않습니다. 특정 주요 릴리스 내에서 더 이상 사용되지 않는 기능의 최신 목록은 최신 릴리스 노트를 참조하십시오.

CephFS에 대한 NFS 지원이 더 이상 사용되지 않음

CephFS에 대한 NFS 지원이 더 이상 사용되지 않습니다. OpenStack Manila의 NFS에 대한 Red Hat Ceph Storage 지원은 영향을 받지 않습니다. 더 이상 사용되지 않는 기능은 현재 릴리스의 수명 동안 버그 수정만 수신하며 향후 릴리스에서 제거될 수 있습니다. 이 기술에 대한 관련 문서는 "**Limited Availability**"로 식별됩니다.

iSCSI 지원이 더 이상 사용되지 않음

iSCSI 지원이 더 이상 사용되지 않으며 향후 NVMeoF 지원이 제공됩니다. 더 이상 사용되지 않는 기능은 현재 릴리스의 수명 동안 버그 수정만 수신하며 향후 릴리스에서 제거될 수 있습니다. 이 기술에 대한 관련 문서는 "**Limited Availability**"로 식별됩니다.

6장. 버그 수정

이 섹션에서는 **Red Hat Ceph Storage** 릴리스에서 수정된 사용자 영향이 상당한 버그에 대해 설명합니다. 또한 섹션에는 이전 버전에서 발견된 수정된 알려진 문제에 대한 설명이 포함되어 있습니다.

6.1. CEPHADM 유틸리티

컨테이너 프로세스 번호 제한을 **max**로 설정

이전에는 컨테이너의 프로세스 번호 제한 **2048**이 제한을 초과하여 새 프로세스를 분기하지 못했습니다.

이번 릴리스에서는 프로세스 번호 제한이 **max**로 설정되어 대상당 필요한 만큼의 **luns**를 만들 수 있습니다. 그러나 이 수는 여전히 서버 리소스에 의해 제한됩니다.

([BZ#1976128](#))

배치에서 **OSD**를 만들 때 사용할 수 없는 장치는 더 이상 전달되지 않습니다.

이전에는 **GPT** 헤더가 있는 장치가 사용할 수 없는 것으로 표시되지 않았습니다. **Cephadm**은 다른 유효한 장치와 함께 다른 유효한 장치와 함께 **OSD**를 만들려고 합니다. 배치에서는 **GPT** 헤더가 있는 장치에서 **OSD**를 만들 수 없기 때문에 배치 **OSD**가 생성되지 않습니다. **OSD**는 생성되지 않습니다.

이번 수정에서는 배치에 **OSD**를 만들 때 사용할 수 없는 장치를 더 이상 전달하지 않고 **GPT** 헤더를 사용하는 장치가 더 이상 유효한 장치에서 **OSD**를 생성하지 않습니다.

([BZ#1962511](#))

지원되지 않는 형식을 사용하여 **--format** 인수를 제공하는 사용자는 역추적이 수신되었습니다.

이전 버전에서는 **orchestrator**가 지원되지 않는 **--format** 인수를 수신할 때마다 예외가 발생하여 지원되지 않는 형식으로 **--format**을 통과한 사용자가 역추적을 받을 수 있었습니다.

이번 수정으로 지원되지 않는 형식이 올바르게 처리되고 지원되지 않는 형식을 제공하는 사용자는 해당 형식이 지원되지 않음을 설명하는 메시지를 받습니다.

(BZ#2006214)

이제 종속성 오류 없이 **ceph-common** 패키지를 설치할 수 있습니다.

이전 버전에서는 **Red Hat Ceph Storage 4**를 **Red Hat Ceph Storage 5**로 업그레이드한 후 몇 가지 패키지가 없어 종속성 오류가 발생했습니다.

이번 수정에서는 남은 **Red Hat Ceph Storage 4** 패키지가 제거되고 사전 빌드 플레이북 실행 중에 **ceph-common** 패키지를 설치할 수 있습니다.

(BZ#2008402)

tcmu-runner 데몬은 더 이상 **stray** 데몬으로 보고되지 않습니다.

이전에는 **tcmu-runner** 데몬이 **iSCSI**의 일부로 간주되므로 **cephadm**에서 적극적으로 추적되지 않았습니다. 이로 인해 **cephadm**이 추적되지 않았기 때문에 **tcmu-runner** 데몬이 **stray** 데몬으로 보고되었습니다.

이번 수정으로 **tcmu-runner** 데몬이 알려진 **iSCSI** 데몬과 일치하면 **stray** 데몬으로 표시되지 않습니다.

(BZ#2018906)

사용자는 명시적 **IP** 없이 활성 관리자를 사용하여 호스트를 다시 추가할 수 있습니다.

이전 버전에서는 **cephadm**에서 컨테이너 내에서 현재 호스트의 **IP** 주소를 확인하려고 할 때마다 루프백 주소로 확인될 가능성이 있었습니다. 사용자가 활성 **Ceph Manager**를 사용하여 호스트를 다시 추가해야 하는 경우 명시적 **IP**가 필요하며 사용자는 이를 제공하지 않으면 오류 메시지가 표시됩니다.

현재 수정으로 **cephadm**은 명시적으로 제공되지 않는 경우 호스트를 다시 추가할 때 이전 **IP**를 재사용하고 이름 확인에서 루프백 주소를 반환합니다. 사용자는 이제 명시적 **IP** 없이 활성 관리자로 호스트를 다시 추가할 수 있습니다.

(BZ#2024301)

cephadm은 데몬의 **fsid**가 예상되는 **fsid**와 일치하는 구성을 유추하는지 확인합니다.

이전에는 **cephadm**에서 데몬의 **fsid**가 예상 **fsid**와 일치했는지 확인할 수 없었습니다. 이로 인해 사

용자에게 예상 이외의 **fsid** 가 있는 **/var/lib/ceph/FSID/ DAE jenkinsfile_NAME** 디렉터리가 있는 경우 해당 데몬 디렉터리의 구성은 여전히 유추됩니다.

이번 수정을 통해 **fsid** 가 예상과 일치하는지 확인하고 사용자는 더 이상 데몬 또는 장치 검색에 "실패한" 오류가 발생하지 않습니다.

(BZ#2024720)

cephadm 에서 다른 이름으로 클라이언트 인증 키 복사 지원

이전에는 클라이언트 인증 키 **ceph.keyring** 을 복사할 때 **cephadm** 에서 파일 이름을 적용했습니다.

현재 수정으로 **cephadm** 은 클라이언트 키링을 다른 이름으로 복사하여 복사 시 자동 이름 변경 문제가 제거됩니다.

(BZ#2028628)

-c ceph.conf 옵션을 사용하여 여러 공용 네트워크가 있는 클러스터를 부트스트랩할 수 있습니다.

이전 버전에서는 **cephadm** 이 **-c ceph.conf** 옵션의 일부로 제공되면 부트스트랩 중에 여러 공용 네트워크를 구문 분석하지 않았습니다. 이로 인해 여러 공용 네트워크가 있는 클러스터를 부트스트랩할 수 없었습니다.

현재 수정으로 제공된 **ceph.conf** 파일에서 공개 네트워크 필드가 올바르게 구문 분석되고 **public_network mon config** 필드를 채우는 데 사용할 수 있으므로 사용자가 **-c ceph.conf** 옵션을 사용하여 여러 공용 네트워크를 제공하는 클러스터를 부트스트랩할 수 있습니다.

(BZ#2035179)

숫자 서비스 ID로 **MDS** 서비스를 설정하면 사용자에게 경고하는 오류가 발생합니다.

이전에는 숫자 서비스 ID를 사용하여 **MDS** 서비스를 설정하면 **MDS** 데몬이 충돌했습니다.

이번 수정을 통해 숫자 서비스 ID를 사용하여 **MDS** 서비스를 생성하려고 하면 오류가 즉시 경고되도록 실패하고 숫자 서비스 ID를 사용하지 않도록 사용자에게 경고 메시지가 표시됩니다.

(BZ#2039669)

ceph orch redeploy mgr 명령은 활성 **Manager** 데몬을 마지막으로 재배포합니다.

이전 버전에서는 **ceph orch redeploy mgr** 명령으로 인해 **Ceph Manager** 데몬이 예약된 재배포 작업을 지우지 않고 지속적으로 재배포되어 **Ceph Manager** 데몬이 무한대로 밀려났습니다.

이번 릴리스에서는 활성 관리자 데몬이 항상 마지막을 다시 배포하고 **ceph** 또는 **ch**를 재배포하는 명령은 이제 각 **Ceph Manager**를 한 번만 재배포 하도록 재배포된 재배포가 변경되었습니다.

(BZ#2042602)

사용자 정의 이름으로 클러스터 채택 지원

이전 버전에서는 **cephadm** 이 **unit.run** 파일의 사용자 정의 클러스터를 전파하지 않았기 때문에 **Ceph** 클러스터에서 **Ceph OSD** 컨테이너를 채택하지 못했습니다.

이번 릴리스에서는 **cephadm** 이 **LVM** 메타데이터를 변경하고 기본 클러스터 이름 "**Ceph**"를 적용하여 사용자 지정 클러스터 이름이 있는 클러스터가 예상대로 작동합니다.

(BZ#2058038)

cephadm 에서 더 이상 **ceph orch upgrade start** 명령에 제공된 이미지 이름에 **docker.io** 를 추가하지 않습니다.

이전에는 **cephadm** 에서 정규화되지 않은 레지스트리의 이미지에 **docker.io** 를 추가했기 때문에 이 이미지를 가져오지 못해 로컬 레지스트리와 같은 정규화되지 않은 레지스트리에서 이미지를 전달할 수 없었습니다.

Red Hat Ceph Storage 5.2부터 **docker.io** 는 **ceph/ceph:v17** 과 같은 업스트림 **ceph** 이미지에 이름이 일치하지 않는 한 이미지 이름에 더 이상 추가되지 않습니다. **ceph orch upgrade** 명령을 실행하면 사용자가 로컬 레지스트리에서 이미지를 전달할 수 있으며 **Cephadm** 은 해당 이미지로 업그레이드할 수 있습니다.



참고

이는 5.2부터 업그레이드에만 적용됩니다. 5.1에서 5.2로 업그레이드하는 것은 여전히 이 문제의 영향을 받습니다.

(BZ#2077843)

Cephadm 에서 더 이상 레거시 데몬의 구성 파일을 유추하지 않습니다.

이전에는 **Cephadm** 이 `/var/lib/ceph/{mon|osd|mgr}` 디렉터리의 존재 여부에 관계없이 기존 데몬의 구성 파일을 유추했습니다. 이로 인해 **Cephadm** 에서 존재하지 않는 구성 파일을 추측하려고 할 때 이러한 디렉토리가 존재하는 노드에서 디스크 새로 고침과 같은 특정 작업이 실패했습니다.

현재 수정으로 **Cephadm** 은 기존 데몬의 구성 파일을 더 이상 유추하지 않고, 대신 기존 구성 파일을 추측하기 전에 확인합니다. **Cephadm** 은 레거시 데몬 디렉토리가 존재하기 때문에 호스트에서 데몬 또는 장치를 새로 고칠 때 더 이상 문제가 발생하지 않습니다.

(BZ#2080242)

.RGW.root 풀이 더 이상 자동으로 생성되지 않음

이전에는 다중 사이트에 대한 **Ceph Object Gateway**에 대한 추가 검사가 존재했습니다. 이로 인해 사용자가 삭제한 경우에도 **.rgw.root** 풀이 자동 생성되었습니다.

Red Hat Ceph Storage 5.2부터 다중 사이트 점검이 제거되고 **.rgw.root** 풀이 더 이상 자동으로 생성되지 않습니다. 사용자가 **Ceph Object Gateway**를 사용하지 않는 한 생성을 초래합니다.

(BZ#2083885)

Ceph Manager 데몬은 **cephadm**의 배치 사양에 더 이상 지정되지 않은 호스트에서 제거됩니다.

이전에는 관리자 서비스 사양에 지정된 배치와 더 이상 일치하지 않는 경우에도 현재 활성 관리자 데몬이 **cephadm** 에서 제거되지 않았습니다. 사용자가 현재 활성 관리자가 된 호스트를 제외하도록 서비스 사양을 변경할 때마다 장애 조치가 발생할 때까지 추가 관리자로 끝납니다.

이번 수정을 통해 **Wait**를 사용할 수 있고 활성 관리자가 더 이상 서비스 사양과 일치하지 않는 호스트에 있는 경우 **cephadm** 이 관리자를 통해 실패합니다. **Ceph Manager** 데몬은 관리자가 활성 상태인 경우에

도 **cephadm** 의 배치 사양에 더 이상 지정되지 않은 호스트에서 제거됩니다.

([BZ#2086438](#))

잘못된 형식의 **URL**로 인해 **404** 오류가 발생하면 로그에 역추적이 발생했습니다.

이전 버전에서는 **cephadm** 이 **prometheus** 수신자의 **URL**을 잘못 형성하여 잘못된 **URL**에 액세스하려고 할 때 발생하는 **404** 오류로 인해 추적이 로그에 출력되었습니다.

이번 수정으로 **URL** 형식이 수정되었으며 **404** 오류가 발생하지 않습니다. 역추적은 더 이상 기록되지 않습니다.

([BZ#2087736](#))

cephadm 에서 더 이상 호스트 수준에서 **osd_memory_target** 구성 설정을 제거하지 않습니다.

이전 버전에서는 **osd_memory_target_autotune** 이 전역적으로 꺼져 있는 경우 **cephadm** 은 호스트 수준에서 **osd_memory_target** 에 설정한 값을 제거합니다. 또한 **ReplicaSet** 맵이 짧은 이름을 사용하지만, **FQDN**이 있는 호스트의 경우, **cephadm** 은 여전히 **FQDN**을 사용하여 **config** 옵션을 설정합니다. 이로 인해 호스트 수준에서 **osd_memory_target** 을 수동으로 설정할 수 없어 **osd_memory_target** 자동 튜닝이 **FQDN** 호스트에서 작동하지 않았습니다.

이번 수정을 통해 **osd_memory_target** 구성 설정이 **osd_memory_target_autotune** 이 **false** 로 설정된 경우 호스트 수준에서 **cephadm** 에서 제거되지 않습니다. 또한 호스트 수준 **osd_memory_target** 을 설정할 때 항상 호스트의 짧은 이름을 사용합니다. 호스트 수준 **osd_memory_target_autotune** 이 **false** 로 설정된 경우 사용자는 **osd_memory_target** 을 수동으로 설정하고 **cephadm** 에서 옵션을 제거하지 않도록 할 수 있습니다. 또한 자동 조정은 이제 **FQDN** 이름이 있는 **cephadm** 에 추가된 호스트에서 작동합니다.

([BZ#2092089](#))

Cephadm 은 **FQDN**을 사용하여 **alertmanager** 웹 후크 **URL**을 빌드합니다.

이전에는 **Cephadm** 이 호스트에 저장된 **IP** 주소를 기반으로 **alertmanager** 웹 후크를 선택했습니다. 이로 인해 이러한 웹 후크 **URL**이 특정 배포에서 작동하지 않기 때문에 문제가 발생했습니다.

이번 수정으로 **Cephadm** 은 **FQDN**을 사용하여 **alertmanager** 웹 후크 **URL**을 빌드하여 이전에 손상된 일부 배포 상황에서 **Webhook URL**을 사용할 수 있습니다.

(BZ#2099348)

6.2. CEPH 대시보드

Ceph 대시보드에서 트레이닝 작업을 사용하면 호스트를 안전하게 삭제할 수 있습니다.

이전 버전에서는 사용자가 모든 데몬을 이동하지 않고 Ceph 대시보드에서 호스트를 제거할 때마다 호스트는 사용할 수 없는 상태 또는 **Recreate** 상태로 전환되었습니다.

이번 수정을 통해 사용자는 대시보드에서 트레이닝 작업을 사용하여 호스트에서 모든 데몬을 이동할 수 있습니다. 트레이닝 작업이 완료되면 호스트를 안전하게 제거할 수 있습니다.

(BZ#1889976)

Ceph 대시보드의 필요한 데이터를 보여주는 성능 세부 그래프

이전 버전에서는 관련 메트릭이 오래되어 있을 때 데몬의 성능 세부 정보 그래프에 **put/get** 작업이 수행된 경우에도 데이터가 표시되지 않았습니다.

이번 수정으로 관련 메트릭이 최신 상태이며 성능 세부 정보 그래프에 필요한 데이터가 표시됩니다.

(BZ#2054967)

Alertmanager 에서 올바른 MTU 불일치 경고를 표시

이전에는 Alertmanager 에서 **down** 상태인 카드에 대한 잘못된 MTU 불일치 경고를 표시했습니다.

이번 수정으로 Alertmanager 는 올바른 MTU 불일치 경고를 보여줍니다.

(BZ#2057307)

PG 상태 차트에서 더 이상 알 수 없는 배치 그룹 상태를 표시하지 않음

이전에는 **snaptrim_wait** 배치 그룹(PG) 상태가 잘못 구문 분석되고 2개의 상태, **snaptrim** 및 **wait** 로 분할되어 PG 상태가 유효하지 않았습니다. 이로 인해 PG 상태 차트가 모두 알려진 상태인 경우에도 PG

상태 차트가 알 수 없는 상태의 몇 개의 **PG**를 잘못 표시했습니다.

이번 수정을 통해 **snaptrim_wait** 및 밑줄이 포함된 모든 상태가 올바르게 구문 분석되고 알 수 없는 **PG** 상태가 더 이상 **PG** 상태 차트에 표시되지 않습니다.

([BZ#2077827](#))

Ceph 대시보드 개선 사용자 인터페이스 개선

이전에는 **Ceph Dashboard** 사용자 인터페이스에서 다음 문제가 확인되어 다중 경로 스토리지 클러스터에서 테스트할 때 사용할 수 없게 되었습니다.

- 클러스터에서 다중 경로 스토리지 장치를 사용하면 *물리적* 디스크 페이지에서 디스크를 선택하면 테이블이 1분 내에 중지될 때까지 테이블의 선택 수가 증가하기 시작합니다.
- **Device Health (Device Health)** 페이지에 **SMART** 데이터를 가져오는 동안 오류가 표시되었습니다.
- **Hosts (호스트)** 페이지의 *서비스 열*에는 많은 항목이 표시되므로 가독성이 줄어듭니다.

이번 릴리스에서는 다음 수정 사항이 구현되어 사용자 인터페이스가 향상됩니다.

- **Physical Disks** 페이지에서 디스크 선택 문제를 수정했습니다.
- **scsi** 장치 **SMART** 데이터를 가져오는 옵션이 추가되었습니다.
- *서비스 열*은 **Service Instances** 로 이름이 변경되고 해당 서비스의 인스턴스 이름 및 인스턴스 수만 배지에 표시됩니다.

([BZ#2097487](#))

6.3. CEPH 파일 시스템

모든 디렉터리에 **ceph.dir.layout** 을 가져오면 가장 가까운 상속된 레이아웃이 반환됩니다.

이전에는 디렉터리 경로에서 루트를 통과하여 가장 가까운 상속된 레이아웃을 찾지 않아 시스템에서 레이아웃이 구체적으로 설정되지 않은 디렉토리에 대해 **"No such attribute"** 메시지를 반환했습니다.

이번 수정으로 디렉터리 경로는 루트를 통과하여 가장 가까운 상속된 레이아웃을 찾고 디렉터리 계층 구조의 모든 디렉터리에 대해 **ceph.dir.layout** 을 가져옵니다.

([BZ#1623330](#))

subvolumegroup ls API는 내부 **ping** 디렉토리 **_deleting**을 필터링합니다.

이전에는 **subvolumegroup ls API**에서 내부 **Blob** 디렉터리 **_deleting** 을 필터링하지 않아 하위 볼륨 그룹으로 나열되었습니다.

이번 수정으로 **subvolumegroup ls API**는 내부 **Blob** 디렉터리 **_deleting** 을 필터링하고 **subvolumegroup ls API**에는 내부 **ping** 디렉터리 **_deleting** 이 표시되지 않습니다.

([BZ#2029307](#))

경쟁 조건으로 인해 더 이상 클러스터의 **MDS**가 혼동되지 않습니다.

이전에는 약자 설정 중에 **MDS**의 경쟁 조건으로 인해 클러스터의 다른 **MDS**가 혼동되어 다른 **MDS**가 통신을 거부했습니다.

이번 수정을 통해 경쟁 조건이 수정되어 **MDS** 간에 통신이 성공적으로 설정됩니다.

([BZ#2030540](#))

MDS에서 온라인 스크립을 사용하여 스레이 재 제공을 트리거할 수 있습니다.

이전에는 **stray reintegrations**가 클라이언트 요청에서만 트리거되어 클라이언트에 의해 비용이 많이 드는 재귀 디렉터리 목록이 필요하도록 **stray inode**를 지우는 프로세스가 발생했습니다.

이번 수정으로 **MDS**는 이제 온라인 스크립을 사용하여 **stray reintegration**을 트리거할 수 있습니다.

(BZ#2041563)

대상 디렉터리가 가득 차면 **MDS 재integrates strays**

이전 버전에서는 링크의 대상 디렉터리가 가득 차면 **MDS**가 다시 의도하지 않아 **degenerate** 상황에서 **stray** 디렉터리가 채워지지 않았습니다.

이번 수정을 통해 크기 변경이 발생하지 않으므로 대상 디렉터리가 가득 차 있는 경우에도 **MDS**는 스프레이 통합을 진행합니다.

(BZ#2041571)

데이터가 복사된 후 복제에 할당량이 적용됩니다.

이전에는 소스 스냅샷에서 데이터를 복사하기 전에 복제의 할당량을 설정하고 소스에서 전체 데이터를 복사하기 전에 할당량이 적용되었습니다. 이로 인해 소스의 할당량을 초과하면 **subvolume** 스냅샷 복제가 실패합니다. 할당량은 바이트 범위에서 엄격하게 적용되지 않으므로 가능합니다.

이번 수정을 통해 데이터가 복사된 후 복제에 할당량이 적용됩니다. 스냅샷 복제본은 할당량에 관계없이 항상 성공합니다.

(BZ#2043602)

ceph-mgr 을 다시 시작한 후 재해 복구 자동화 및 계획이 재개됨

이전 버전에서는 **ceph-mgr** 시작 중에 일정이 시작되지 않아 스냅샷 일정이 **ceph-mgr** 재시작 시 재개된 사용자 재해 복구 계획에 영향을 미쳤습니다.

이번 수정으로 **ceph-mgr** 을 다시 시작할 때 일정이 시작되고 스냅샷 복제와 같은 재해 복구 자동화 및 계획에서는 수동 조작 없이 **ceph-mgr** 이 다시 시작된 후 즉시 다시 시작됩니다.

(BZ#2055173)

읽기 위해 파일을 여는 경우 **mdlog** 는 즉시 플러시됩니다.

이전에는 읽기 위해 파일을 열 때 **MDS**가 다른 클라이언트에서 **Fw** 기능을 취소하고 **Fw** 기능이 릴리스될 때 **MDS**에서 즉시 **mdlog** 를 플러시할 수 없어 **Fr** 기능이 차단되었습니다. 이로 인해 파일에 대해 요청

한 프로세스가 5초마다 주기적으로 **MDS**에서 플러시될 때까지 약 5초 동안 멈춥니다.

이번 릴리스에서는 **Fw** 기능을 해제할 때 원하는 기능이 있으면 **mdlog** 플러시가 즉시 트리거되며 파일을 신속하게 열 수 있습니다.

(BZ#2076850)

특정 복제 상태에 하위 볼륨 복제 삭제가 더 이상 허용되지 않음

이전에는 복제가 **COMPLETED** 또는 **CanCElLED** 상태에 있지 않은 경우 **force** 옵션이 있는 하위 볼륨 복제를 제거하려고 하면 해당 복제본이 현재 복제본을 추적한 인덱스에서 제거되지 않았습니다. 이로 인해 해당 복제 스레드가 무기한 재시도하여 결국 **ENOENT** 오류가 발생했습니다. 기본 복제 스레드 수를 4개로 설정하면 4개의 복제본을 모두 삭제하면 4개의 스레드가 모두 차단된 상태로 전환되어 보류 중인 복제본을 완료할 수 없었습니다.

이 릴리스에서는 복제가 **COMPLETED** 또는 **CanCElLED** 상태에 있지 않으면 제거되지 않습니다. 복제본이 삭제되고 인덱스의 항목이 지속적인 복제본을 추적하는 것과 함께 복제기 때문에 복제 스레드가 더 이상 차단되지 않습니다. 결과적으로 보류 중인 복제본은 예상대로 계속 완료됩니다.

(BZ#2081596)

새 클라이언트는 이전 **Ceph** 클러스터와 호환됩니다.

이전에는 새 클라이언트가 이전 **Ceph** 클러스터와 호환되지 않아 이전 클러스터에서 **abort()** 를 트리거하여 알 수 없는 메트릭을 수신할 때 **MDS** 데몬을 충돌했습니다.

이번 수정으로 클라이언트의 기능 비트를 확인하고 **MDS**에서 지원하는 지표만 수집하고 전송하십시오. 새 클라이언트는 이전 **cephs**와 호환됩니다.

(BZ#2081929)

동시 조회 및 연결 해제 작업 중에 **Ceph Metadata Server**가 더 이상 충돌하지 않음

이전에는 코드에 어설션이 잘못 가정되어 **Ceph** 클라이언트에서 동시 조회 및 연결 해제 작업이 충돌하여 **Ceph Metadata Server**가 충돌했습니다.

최신 수정에서는 동시 조회 및 연결 해제 작업 중에 가정이 유효하여 **Ceph** 클라이언트 작업을 중단하지 않고 **Ceph** 클라이언트 작업을 지속적으로 제공하는 관련 위치로 어설션을 이동합니다.

(BZ#2093065)

연결되지 않은 디렉토리를 가져올 때 **MDS**가 더 이상 충돌하지 않습니다.

이전 버전에서는 연결되지 않은 디렉토리를 가져올 때 예상 버전이 잘못 초기화되어 온전성 검사를 수행할 때 **MDS**가 충돌했습니다.

이번 수정을 통해 연결되지 않은 디렉토리를 가져올 때 예상 버전과 **inode** 버전이 초기화되어 **MDS**가 중단 없이 온전성 검사를 수행할 수 있습니다.

(BZ#2108656)

6.4. CEPH MANAGER 플러그인

우선 순위 **Cache perf** 카운터에 누락된 포인터가 추가되고 **perf** 출력은 **prioritycache** 키 이름을 반환합니다.

이전에는 **PriorityCache perf** 카운터 빌더에 필요한 포인터가 없어지고, **perf** 카운터 출력에 **DAEtekton _TYPE . DAECDHE_ ID perf dump_ID** 를 지정하고 **ceph**는 **DAE jenkinsfile _TYPE. DAE jenkinsfile_ID /f** 스키마 대신 빈 문자열을 반환하도록 합니다. 이 누락된 키로 인해 **collectd-ceph** 플러그인에서 오류가 발생했습니다.

이번 수정을 통해 **PriorityCache perf** 카운터 빌더에 누락된 포인터가 추가됩니다. **perf** 출력은 **prioritycache** 키 이름을 반환합니다.

(BZ#2064627)

Native CephFS 및 외부 Red Hat Ceph Storage 5를 사용하는 OpenStack 16.x Manila의 취약점

이전에는 **Red Hat Ceph Storage 5.0, 5.0.x, 5.1** 또는 **5.1.x**로 업그레이드한 **OpenStack 16.x(Manila)** 및 외부 **Red Hat Ceph Storage 4**를 실행 중인 고객의 영향을 받을 수 있었습니다. 이 취약점으로 인해 **OpenStack Manila** 사용자/테넌트(**Ceph File System** 공유 소유자)가 **CephFS**에서 지원하는 모든 **Manila** 공유에 대한 액세스(**Read/write**) 또는 전체 **CephFS** 파일 시스템에 대한 액세스(**Read/write**)를 악의적으로 받을 수 있었습니다. 이 취약점은 **Ceph Manager**의 "볼륨" 플러그인의 버그로 인해 발생합니다. 이 플러그인은 **Manila** 사용자에게 공유를 제공하는 방법으로 **OpenStack Manila** 서비스에서 사용하는 **Ceph File System** 하위 볼륨을 관리합니다.

이번 릴리스에서는 이 취약점이 수정되었습니다. 외부 **Red Hat Ceph Storage 5.0, 5.0.x, 5.1** 또는

5.1.x로 업그레이드한 **OpenStack 16.x**(기본 **CephFS** 액세스 제공)를 실행하는 고객은 **Red Hat Ceph Storage 5.2**로 업그레이드해야 합니다. **NFS**를 통해서만 액세스 권한을 제공한 고객은 영향을 받지 않습니다.

([BZ#2056108](#))

6.5. CEPH VOLUME 유틸리티

누락된 백포트가 추가되고 **OSD**를 활성화할 수 있습니다.

이전에는 누락된 백포트로 인한 회귀 문제로 인해 **OSD**를 활성화할 수 없었습니다.

이번 수정을 통해 누락된 백포트가 추가되고 **OSD**를 활성화할 수 있습니다.

([BZ#2093022](#))

6.6. CEPH OBJECT GATEWAY

버전이 지정된 버킷에 대한 라이프 사이클 정책은 **reshard** 간에 더 이상 실패하지 않습니다.

이전 버전에서는 내부 논리 오류로 인해 버킷의 라이프사이클 처리가 버킷을 다시 설정하는 동안 비활성화되어 있었기 때문에 영향을 받는 버킷에 대한 라이프사이클 정책이 처리되지 않았습니다.

이번 수정으로 버그 수정으로 버전이 지정된 버킷의 라이프사이클 정책이 더 이상 재스하드 간에 실패하지 않습니다.

([BZ#1962575](#))

삭제된 오브젝트는 더 이상 버킷 인덱스에 나열되지 않습니다.

이전 버전에서는 **delete** 오브젝트 작업이 정상적으로 완료되지 않으면 버킷 인덱스에 오브젝트가 나열되어 있어야 하는 오브젝트가 여전히 나열되었습니다.

이번 릴리스에서는 불완전한 트랜잭션이 수정되고 삭제된 오브젝트가 더 이상 나열되지 않는 내부 **"dir_suggest"**가 있습니다.

([BZ#1996667](#))

Ceph Object Gateway의 영역 그룹이 **awsRegion** 값으로 전송됩니다.

이전에는 **awsRegion** 값이 이벤트 레코드의 **zonegroup**으로 채워지지 않았습니다.

이번 수정에서는 Ceph Object Gateway의 영역 그룹이 **awsRegion** 값으로 전송됩니다.

([BZ#2004171](#))

Ceph Object Gateway가 빈 주제 목록이 제공되면 모든 알림 항목을 삭제합니다.

이전에는 Ceph Object Gateway에서 알림 항목이 이름으로 정확하게 삭제되었지만 AWS 동작을 따라 빈 주제 이름이 지정되면 모든 항목을 삭제하지 않아 Ceph Object Gateway에서 몇 가지 고객 버킷 알림 워크플로를 사용할 수 없게 되었습니다.

이번 수정으로 빈 주제 목록에 대한 명시적 처리가 변경되어 Ceph Object Gateway가 빈 주제 목록을 제공할 때 모든 알림 항목을 삭제합니다.

([BZ#2017389](#))

버킷 목록 충돌, 버킷 통계 및 유사한 작업이 인덱스 없는 버킷에 대해 표시되지 않습니다.

이전에는 일반 버킷 목록을 포함한 여러 작업에서 인덱스 부족 버킷의 인덱스 정보에 잘못 액세스하려고 하면 충돌이 발생했습니다.

이번 수정으로 인덱스 버킷에 대한 새로운 검사가 추가되어 버킷 목록, 버킷 통계 및 유사한 작업이 표시되지 않습니다.

([BZ#2043366](#))

내부 테이블 인덱스가 음수가 되지 않음

이전에는 연속 작업 후 내부 테이블의 인덱스가 음수가 되어 Ceph Object Gateway가 충돌했습니다.

이번 수정으로 인덱스가 음수가 되지 않아 **Ceph Object Gateway**가 더 이상 충돌하지 않습니다.

([BZ#2079089](#))

FIPS 지원 환경에서 **MD5** 사용은 명시적으로 허용되며 **S3** 다중 파트 작업을 완료할 수 있습니다.

이전 버전에서는 **FIPS** 지원 환경에서 비 암호화 목적으로 명시적으로 제외하지 않는 한 **MD5** 다이제스트의 사용이 기본적으로 허용되지 않았습니다. 이로 인해 **S3** 전체 멀티 파트 업로드 작업 중에 분할이 발생했습니다.

이번 수정으로 **S3** 전체 **multipart PUT** 작업을 위한 **FIPS** 지원 환경에서 비암호화된 목적으로 **MD5**를 사용할 수 있으며 **S3** 다중 파트 작업을 완료할 수 있습니다.

([BZ#2088602](#))

결과 코드 **2002** **radosgw-admin** 명령은 **2**로 명시적으로 번역됩니다.

이전에는 **NoSuchBucket**의 **S3** 오류 변환이 변경되면 실수로 **radosgw-admin** 버킷 통계 명령에서 오류 코드가 변경되어 해당 **radosgw-admin** 명령의 셸 결과 코드를 확인하는 프로그램에서 다른 결과 코드를 확인할 수 있었습니다.

이 수정을 통해 결과 코드 **2002**은 **2**로 명시적으로 변환되며 사용자는 원래 동작을 볼 수 있습니다.

([BZ#2100967](#))

FIPS 지원 환경에서 **MD5** 사용은 명시적으로 허용되며 **S3** 다중 파트 작업을 완료할 수 있습니다.

이전 버전에서는 **FIPS** 지원 환경에서 비 암호화 목적으로 명시적으로 제외하지 않는 한 **MD5** 다이제스트의 사용이 기본적으로 허용되지 않았습니다. 이로 인해 **S3** 전체 멀티 파트 업로드 작업 중에 분할이 발생했습니다.

이번 수정으로 **S3** 전체 **multipart PUT** 작업을 위한 **FIPS** 지원 환경에서 비암호화된 목적으로 **MD5**를 사용할 수 있으며 **S3** 다중 파트 작업을 완료할 수 있습니다.

6.7. 다중 사이트 CEPH 개체 게이트웨이

RADOSGW-admin bi purge 명령은 삭제된 버킷에서 작동합니다.

이전에는 **radosgw-admin bi purge** 명령에 버킷 **entrypoint** 오브젝트가 필요했습니다. 이로 인해 삭제된 버킷에 대한 리소스가 없으므로 삭제된 버킷을 삭제한 후 정리할 수 없었습니다.

이번 수정을 통해 버킷 진입점의 필요성을 방지하기 위해 **bi purge -id**가 **--bucket-id**를 허용하고 명령은 삭제된 버킷에서 작동합니다.

(BZ#1910503)

null 포인터 검사로 인해 더 이상 다중 사이트 데이터 동기화가 발생하지 않음

이전에는 **null** 포인터 액세스가 다중 사이트 데이터 동기화를 중단했습니다.

이번 수정으로 **null** 포인터 검사가 성공적으로 구현되어 가능한 모든 충돌이 발생하지 않습니다.

(BZ#1967901)

오류가 발생할 때 메타데이터 동기화가 더 이상 중단되지 않음

이전 버전에서는 메타데이터 동기화의 일부 오류로 인해 **Ceph Object Gateway** 다중 사이트 구성에서 일부 오류가 발생할 때 동기화가 중단되었습니다.

이번 수정을 통해 재시도 동작이 수정되고 오류가 발생할 때 메타데이터 동기화가 중단되지 않습니다.

(BZ#2068039)

rgw_data_notify_interval_msec=0 매개변수에 특수 처리가 추가되었습니다.

이전에는 **rgw_data_notify_interval_msec**에서 **0**을 특별히 처리하지 않아 보조 사이트를 알림으로 플러싱했습니다.

이번 수정으로 **rgw_data_notify_interval_msec=0**에 대한 특수 처리가 추가되어 **async** 데이터 알림을 비활성화할 수 있습니다.

(BZ#2102365)

6.8. RADOS

PGS는 더 이상 스트레인 모드에서 다시 매핑 + 피어링 상태로 잘못 고정되지 않습니다.

이전 버전에서는 논리 오류로 인해 스트레치 모드로 클러스터를 실행할 때 특정 클러스터 조건에서 일부 배치 그룹(**PG**)이 영구적으로 다시 매핑된+피어링 상태로 유지될 수 있었기 때문에 **OSD**가 오프라인 상태가 될 때까지 데이터를 사용할 수 없었습니다.

이번 수정으로 **PG**는 안정적인 **OSD** 세트를 선택하고 더 이상 스트레칭 모드에서 **remapped+peering** 상태로 잘못 고정되지 않습니다.

([BZ#2042417](#))

OSD 배포 톨은 클러스터를 변경하는 동안 모든 **OSD**를 성공적으로 배포합니다.

KVMonitor paxos 서비스는 클러스터에 변경 사항을 수행할 때 추가, 제거 또는 수정되는 키를 관리합니다. 이전에는 **OSD** 배포 톨을 사용하여 새 **OSD**를 추가하는 동안 서비스에서 쓸 수 있는지 확인하지 않고 키가 추가되었습니다. 이로 인해 **paxos** 코드에서 어설션 오류가 발생하여 모니터가 충돌하게 됩니다.

최신 수정을 통해 새 **OSD**를 추가하기 전에 **KVMonitor** 서비스가 쓰기가 가능한지 확인합니다. 그렇지 않으면 명령은 나중에 다시 관련 대기열로 푸시됩니다. **OSD** 배포 톨은 문제 없이 모든 **OSD**를 성공적으로 배포합니다.

([BZ#2086419](#))

PG 로그의 손상된 **dup** 항목은 오프라인 및 온라인 트리밍으로 제거할 수 있습니다.

이전에는 **PG** 로그 **dup** 항목의 트리밍을 낮은 수준의 **PG** 분할 작업 중에 방지하여 작업자보다 더 높은 빈도가 있는 **PG** 자동 스케일러에서 사용할 수 있었습니다. **dups**의 트리밍을 중지하면 **PG** 로그의 상당한 메모리 증가로 인해 메모리가 부족해 **OSD** 충돌이 발생했습니다. **PG** 로그가 디스크에 저장되고 시작 시 **RAM**으로 다시 로드되므로 **OSD**를 다시 시작하지 못했습니다.

이번 수정을 통해 (**ceph-objectstore-tool** 명령을 사용) 및 온-라인 (**Adin OSD**) 트리밍을 통해 온라인 트리밍 머신을 노출했으며 메모리 증가를 담당하는 **PG** 로그의 손상된 **dup** 항목을 제거할 수 있습니다. 향후 조사를 지원하기 위해 **dups** 항목 수를 **OSD**의 로그에 출력하는 디버그 개선 사항이 구현됩니다.

([BZ#2093031](#))

6.9. RBD 미러링

last_copied_object_number 값이 모든 이미지에 대해 올바르게 업데이트됨

이전에는 구현 결함으로 인해 **last_copied_object_number** 값이 완전히 할당된 이미지에 대해서만 올바르게 업데이트되었습니다. 이로 인해 스냅스 이미지에 대해 **last_copied_object_number** 값이 올바르지 않고 **abrupt rbd-mirror** 데몬 재시작 시 이미지 복제 진행률이 손실됩니다.

이번 수정을 통해 **last_copied_object_number** 값이 모든 이미지에 대해 올바르게 업데이트되고 **rbd-mirror** 데몬 재시작 시 이미지 복제가 이전에 중지된 위치에서 재개됩니다.

([BZ#2019909](#))

기존 일정을 사용하면 이미지를 1차로 승격할 때 적용됩니다.

이전 버전에서는 일치하지 않는 최적화로 인해 기존 일정이 이미지 승격에 따라 최근 승격된 이미지를 시작할 수 없었습니다.

이번 릴리스에서는 최적화 **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#2020618](#))

스냅샷 기반 미러링 프로세스가 더 이상 취소되지 않음

이전 버전에서는 내부 경쟁 조건으로 **rbd** 미러 스냅샷 일정 추가 명령이 취소되었습니다. 다른 기존 일정이 적용되지 않은 경우 영향을 받는 이미지의 스냅샷 기반 미러링 프로세스가 시작되지 않습니다.

이번 릴리스에서는 경쟁 조건이 수정되고 스냅샷 기반 미러링 프로세스가 예상대로 시작됩니다.

([BZ#2069720](#))

원격 이미지가 기본이 아닌 경우 재생 또는 다시 동기화가 더 이상 시도되지 않습니다.

이전 버전에서는 구현 결함, 재생 또는 다시 동기화로 인해 원격 이미지가 기본 이미지가 아닌 경우에도 다시 시도되어 해당 이미지를 재생하거나 다시 동기화할 위치가 없었습니다. 이로 인해 스냅샷 기반 미

러링이 라이브 잠금으로 실행되고 "**failed to unlink local peer from remote image**" 오류를 지속적으로 보고합니다.

이번 수정으로 원격 이미지가 기본이 아닌 경우 구현 결함이 수정되고 재생 또는 다시 동기화되지 않으므로 오류가 보고되지 않습니다.

(BZ#2081715)

보조 클러스터에서 **rbd-mirror** 데몬에서 사용 중인 미러 스냅샷은 제거되지 않습니다.

이전 버전에서는 내부 경쟁 조건으로 인해 보조 클러스터의 **rbd-mirror** 데몬에서 사용 중인 미러 스냅샷이 제거되어 영향을 받는 이미지의 스냅샷 기반 미러링 프로세스가 중지되어 "**split-tekton**" 오류를 보고했습니다.

이번 수정을 통해 미러 스냅샷 큐가 길이가 연장되고 미러 스냅샷 정리 절차가 적절하게 수정됩니다. 보조 클러스터에서 **rbd-mirror** 데몬에서 사용 중인 미러 스냅샷은 더 이상 제거되지 않으며 스냅샷 기반 미러링 프로세스가 중지되지 않습니다.

(BZ#2092838)

schedule_request_lock() 중에 소유자가 잠길 때 **RBD** 미러가 더 이상 충돌하지 않습니다.

이전에는 **schedule_request_lock()** 중에 이미 손상된 소유자의 경우 블록 장치 미러가 충돌하고 이미지 동기화가 중지되었습니다.

이번 수정으로 소유자가 이미 잠겼으면 **schedule_request_lock()** 가 정상적으로 중단되고 블록 장치 미러링이 충돌하지 않습니다.

(BZ#2102227)

이미지 복제가 완료되지 않은 로컬 로컬 스냅샷 오류로 인해 더 이상 중지되지 않음

이전에는 구현 결함으로 인해 **abrupt rbd-mirror** 데몬을 다시 시작하면 이미지 복제가 완료되지 않은 로컬 로컬 비-도스 스냅샷 오류로 중지되었습니다.

이번 수정을 통해 이미지 복제가 불완전한 로컬 비 **snapshot** 오류로 더 이상 중지되지 않으며 예상대로 작동합니다.

([BZ#2105454](#))

6.10. CEPH ANSIBLE 유틸리티

cephadm으로 마이그레이션할 때 **autotune_memory_target_ratio**에 대해 올바른 값이 설정됩니다.

이전 버전에서는 **cephadm**으로 마이그레이션할 때 배포, **HCI** 또는 **non_HCI**에 따라 **autotune_memory_target_ratio**에 대한 적절한 값이 설정되지 않았습니다. 이로 인해 백분율이 설정되지 않았으며 두 배포 사이에 차이가 없습니다.

이번 수정으로 **cephadm-adopt** 플레이북은 배포 종류 및 **autotune_memory_target_ratio** 매개변수에 대해 올바른 값이 설정되어 있는 올바른 비율을 설정합니다.

([BZ#2028693](#))

7장. 확인된 문제

이 섹션에서는 이 **Red Hat Ceph Storage** 릴리스에서 발견된 알려진 문제에 대해 설명합니다.

7.1. CEPHADM 유틸리티

크래시 데몬에서는 크래시 보고서를 스토리지 클러스터에 보내지 못할 수 있습니다.

크래시 데몬 구성 관련 문제로 인해 크래시 데몬에서 클러스터에 크래시 보고서를 보낼 수 없습니다.

([BZ#2062989](#))

Red Hat Ceph Storage 5.2로 업그레이드하는 동안 사용자에게 경고가 표시됩니다.

이전 버전에서는 **Red Hat Ceph Storage 5**에서 버킷이 다시 배치된 버킷은 **Red Hat Ceph Storage 5.2**로 업그레이드한 모든 사용자가 문제를 인식하고 이전에 **Red Hat Ceph Storage 5.1**을 사용하여 이전에 **Red Hat Ceph Storage 5.1**을 사용하여 다운그레이드할 수 있도록 업그레이드 경고/블록러 데몬에서 이해할 수 없었습니다.

해결방법은 오브젝트 스토리지를 사용하지 않거나 **5.1** 이외의 버전에서 업그레이드를 사용하지 않는 사용자는 `ceph config set mgr/cephadm/no_five_one_rgw --force`를 실행하여 `warning/blocker`를 제거하고 모든 작업을 정상으로 반환할 수 있습니다. 이 `config` 옵션을 설정하면 **Red Hat Ceph Storage 5.2**로 업그레이드하기 전에 **Ceph Object Gateway** 문제를 인지하고 있습니다.

([BZ#2104780](#))

NFS 데몬이 있는 가상 IP에서 **HA** 지원 I/O 작업은 장애 조치(failover) **HAProxy** 구성에서 유지 관리되지 않는 **NFS** 데몬을 사용하여 업데이트되지 않습니다.

NFS 데몬을 오프라인에서 온라인 호스트로 장애 조치하면 **HAProxy** 구성이 업데이트되지 않습니다. 그 결과 **HA** 지원 I/O 작업은 **NFS** 데몬이 켜져 있고 장애 조치 전반에 걸쳐 유지 관리되지 않는 가상 IP로 전달됩니다.

([BZ#2106849](#))

7.2. CEPH 대시보드

Ceph 대시보드에서 **SSL**을 사용하여 수신 서비스 생성이 작동하지 않음

Ceph 대시보드에서 **SSL**을 사용한 수신 서비스 생성은 사용자가 더 이상 필수 필드가 아닌 개인 키 필드를 채울 것으로 예상하기 때문에 작동하지 않습니다.

이 문제를 해결하기 위해 **Ceph** 오케스트레이터 **CLI**를 사용하여 **Ingress** 서비스가 성공적으로 생성됩니다.

([BZ#2080276](#))

"Perthrough-optimized" 옵션은 **SSD** 및 **NVMe** 장치를 포함하는 클러스터에 권장됩니다.

이전에는 클러스터에 **SSD** 장치 또는 **SSD** 및 **NVMe** 장치만 있을 때마다 사용자 또는 클러스터에 영향을 미치지 않는 경우에도 "Through throughput-optimized" 옵션이 권장되었습니다.

해결 방법으로 사용자는 원하는 사양에 따라 **OSD**를 배포하기 위해 "고급" 모드를 사용할 수 있으며 이 UI 문제와는 별도로 "Simple" 모드의 모든 옵션을 사용할 수 있습니다.

([BZ#2101680](#))

7.3. CEPH 파일 시스템

getpath 명령으로 인해 자동화 실패

getpath 명령에서 반환하는 디렉터리 이름은 스냅샷을 생성하는 디렉토리 디렉토리이므로 자동화 실패와 혼동이 발생합니다.

이 문제를 해결하려면 한 수준 높은 디렉터리 경로를 사용하여 **snap-schedule add** 명령에 추가하는 것이 좋습니다. 스냅샷은 **getpath** 명령에서 반환된 수준보다 높은 하나의 수준을 사용할 수 있습니다.

([BZ#2053706](#))

7.4. CEPH OBJECT GATEWAY

Red Hat Ceph Storage 5.1 with Ceph Object Gateway 구성에서 Red Hat Ceph Storage 5.2로 업그레이드할 수 없습니다.

모든 **Ceph Object Gateway(RGW)** 클러스터(단일 사이트 또는 다중 사이트)의 **Red Hat Ceph Storage 5.2**로 업그레이드하는 것은 알려진 문제인 [BZ#2100602](#) 문제로 인해 지원되지 않습니다.

자세한 내용은 [RGW 업그레이드에 대한 지원 제한 사항](#)을 참조하십시오.



주의

Red Hat Ceph Storage 5.1 및 Ceph Object Gateway(단일 사이트 또는 다중 사이트)에서 실행되는 Red Hat Ceph Storage 클러스터를 Red Hat Ceph Storage 5.2 릴리스로 업그레이드하지 마십시오.

8장. 소스

업데이트된 **Red Hat Ceph Storage** 소스 코드 패키지는 다음 위치에서 사용할 수 있습니다.

- **Red Hat Enterprise Linux 8의 경우:**
<http://ftp.redhat.com/redhat/linux/enterprise/8Base/en/RHCEPH/SRPMS/>
- **Red Hat Enterprise Linux 9의 경우:**
<http://ftp.redhat.com/redhat/linux/enterprise/9Base/en/RHCEPH/SRPMS/>