

Ceph/Storage 클러스터에서 Ultra-M 격리 및 실패한 디스크 교체 - vEPC

목차

[소개](#)

[배경 정보](#)

[약어](#)

[MoP의 워크플로](#)

[필수 구성 요소 상태 확인](#)

[클러스터에서 결합이 있는 OSD 디스크 격리 및 제거](#)

[OSD 디스크 교체 및 새 VD 생성](#)

[클러스터에 OSD 다시 추가](#)

소개

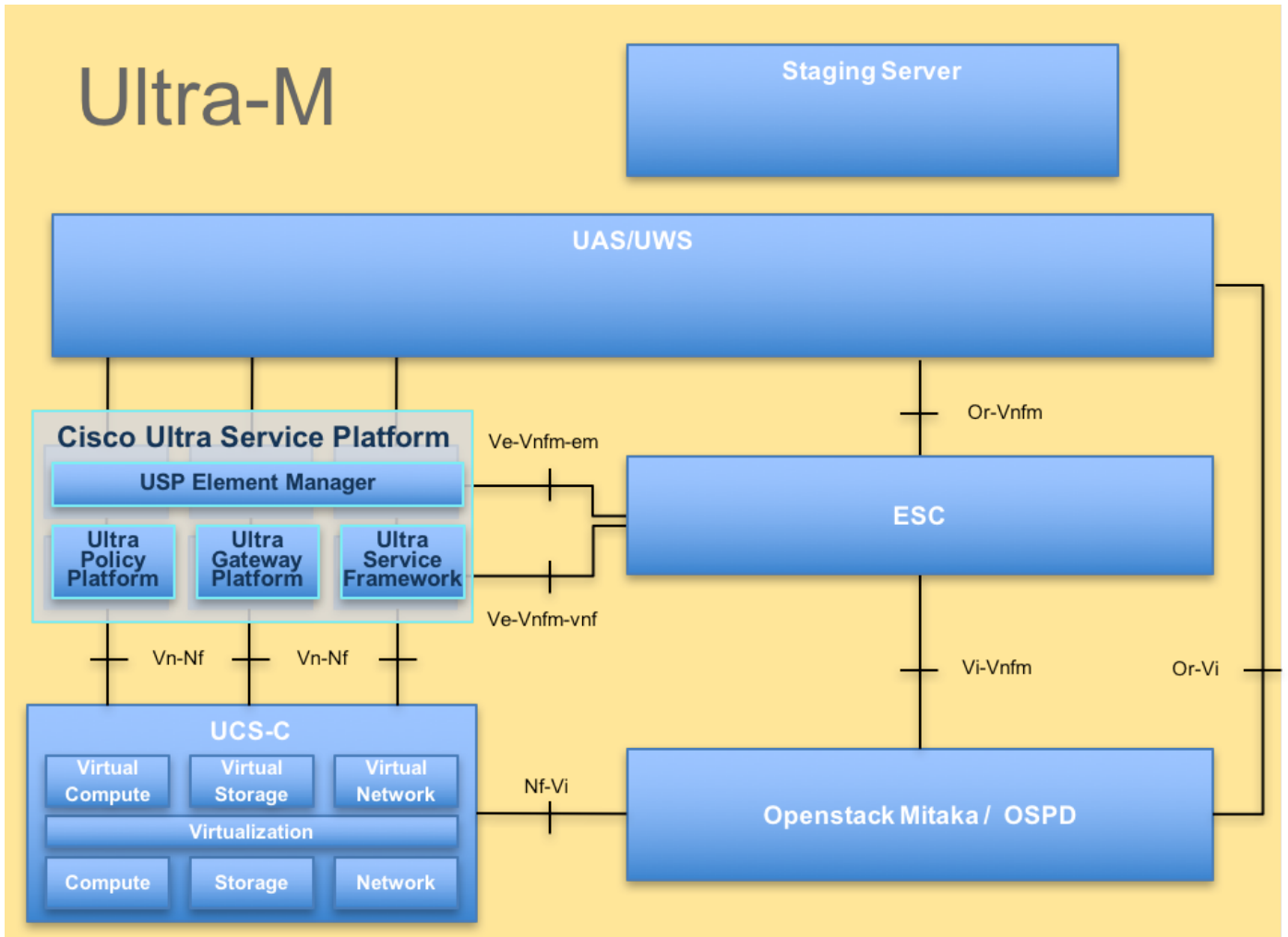
이 문서에서는 Ultra-M 설정에서 OSD(Object Storage Disk)-Compute에 호스팅된 Ceph/Storage 클러스터에서 OSD 디스크를 분리하고 교체하는 데 필요한 단계를 설명합니다.

배경 정보

Ultra-M은 VNF의 구축을 간소화하도록 설계된 사전 패키지 및 검증된 가상 모바일 패킷 코어 솔루션입니다. OpenStack은 Ultra-M용 VIM(Virtualized Infrastructure Manager)이며 다음 노드 유형으로 구성됩니다.

- 컴퓨팅
- OSD - 컴퓨팅
- 컨트롤러
- OpenStack 플랫폼 - 디렉터(OSPD)

이 그림에는 Ultra-M의 고급 아키텍처와 관련 구성 요소가 나와 있습니다.



UltraM 아키텍처이 문서는 Cisco Ultra-M 플랫폼에 대해 잘 알고 있는 Cisco 직원을 대상으로 하며, OSPD 서버 교체 시 OpenStack 레벨에서 수행해야 하는 단계에 대해 자세히 설명합니다.

 참고: 이 문서의 절차를 정의하기 위해 Ultra M 5.1.x 릴리스가 고려됩니다.

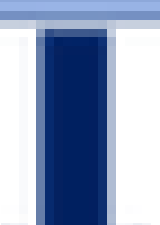
약어

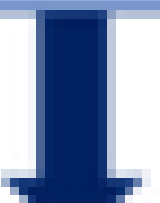
VNF	가상 네트워크 기능
CF	제어 기능
SF	서비스 기능
Esc 키	Elastic Service Controller
자루걸레	절차 방법
OSD	개체 스토리지 디스크
HDD	하드 디스크 드라이브
SSD	SSD(Solid State Drive)
빔	가상 인프라 관리자
VM	가상 머신
엠	요소 관리자

UAS	Ultra Automation 서비스
UUID	보편적으로 고유한 식별자

MoP의 워크플로

- Identify the osd disk to be isolate & replace
- 

- Disable & stop the ceph osd process in identified disk
- 

- Mark the osd out of cluster
 - Verify data rebalance is complete and all the PG's to come back active+clean
- 

```

/dev/sda :
/dev/sda1 other, iso9660
/dev/sda2 other, xfs, mounted on /
/dev/sdb :
/dev/sdb1 ceph journal, for /dev/sdc1
/dev/sdb3 ceph journal, for /dev/sdd1
/dev/sdb2 ceph journal, for /dev/sde1
/dev/sdb4 ceph journal, for /dev/sdf1
/dev/sdc :
/dev/sdc1 ceph data, active, cluster ceph, osd.1, journal /dev/sdb1
/dev/sdd :
/dev/sdd1 ceph data, active, cluster ceph, osd.7, journal /dev/sdb3

/dev/sde :
/dev/sde1 ceph data, active, cluster ceph, osd.4, journal /dev/sdb2
/dev/sdf :
/dev/sdf1 ceph data, active, cluster ceph, osd.10, journal /dev/sdb4

```

2. 식별된 OSD 디스크 격리를 진행하기 전에 Ceph 상태 및 OSD 트리 매핑을 확인합니다.

```

[heat-admin@pod1-osd-compute-3 ~]$ sudo ceph -s
cluster eb2bb192-b1c9-11e6-9205-525400330666
health HEALTH_OK
    1 mons down, quorum 0,1 pod1-controller-0,pod1-controller-1
monmap e1: 3 mons at {pod1-controller-0=11.118.0.10:6789/0,pod1-controller-1=11.118.0.11:6789/0,pod1-controller-2=11.118.0.12:6789/0}
election epoch 28, quorum 0,1 pod1-controller-0,pod1-controller-1
osdmap e709: 12 osds: 12 up, 12 in
    flags sortbitwise,require_jewel_osds
pgmap v941813: 704 pgs, 6 pools, 490 GB data, 163 kobjects
    1470 GB used, 11922 GB / 13393 GB avail
    704 active+clean
client io 58580 B/s wr, 0 op/s rd, 7 op/s wr

```

```

[heat-admin@pod1-osd-compute-3 ~]$ sudo ceph osd tree
ID WEIGHT  TYPE NAME                UP/DOWN REWEIGHT PRIMARY-AFFINITY
-1 13.07996 root default
-2  4.35999  host pod1-osd-compute-0
  0  1.09000    osd.0                up  1.00000        1.00000
  3  1.09000    osd.3                up  1.00000        1.00000
  6  1.09000    osd.6                up  1.00000        1.00000
  9  1.09000    osd.9                up  1.00000        1.00000
-3      0    host pod1-osd-compute-1
-4  4.35999  host pod1-osd-compute-2
  2  1.09000    osd.2                up  1.00000        1.00000
  5  1.09000    osd.5                up  1.00000        1.00000
  8  1.09000    osd.8                up  1.00000        1.00000
 11  1.09000    osd.11               up  1.00000        1.00000
-5  4.35999  host pod1-osd-compute-3
  1  1.09000    osd.1                up  1.00000        1.00000
  4  1.09000    osd.4                up  1.00000        1.00000
  7  1.09000    osd.7                up  1.00000        1.00000
 10  1.09000    osd.10               up  1.00000        1.00000

```

클러스터에서 결함이 있는 OSD 디스크 격리 및 제거

1. OSD 프로세스를 비활성화하고 중지합니다.

```
[heat-admin@pod1-osd-compute-3 ~]$ sudo systemctl disable ceph-osd@7
[heat-admin@pod1-osd-compute-3 ~]$ sudo systemctl stop ceph-osd@7
```

2. OSD를 표시합니다.

```
[heat-admin@pod1-osd-compute-3 ~]$ sudo su

[root@pod1-osd-compute-3 heat-admin]# ceph osd set noout
set noout

[root@pod1-osd-compute-3 heat-admin]# ceph osd set norebalance
set norebalance

[root@pod1-osd-compute-3 heat-admin]# ceph osd out 7
marked out osd.7.
```

 참고: 데이터 리밸런스가 완료되고 모든 PG가 활성+클린으로 돌아올 때까지 기다렸다가 문제를 방지합니다.

3. OSD가 표시되었는지 확인하고 Ceph 리밸런스가 계속 진행될 때까지 기다립니다.

```
[root@pod1-osd-compute-3 heat-admin]# watch -n1 ceph -s
 95 active+undersized+degraded+remapped+wait_backfill
 28 active+recovery_wait+degraded
  2 active+undersized+degraded+remapped+backfilling
  1 active+recovering+degraded
  2 active+undersized+degraded+remapped+backfilling
  1 active+recovering+degraded
  2 active+undersized+degraded+remapped+backfilling
67 active+undersized+degraded+remapped+wait_backfill
  3 active+undersized+degraded+remapped+backfilling
24 active+undersized+degraded+remapped+wait_backfill
22 active+undersized+degraded+remapped+wait_backfill
  1 active+undersized+degraded+remapped+backfilling
  8 active+undersized+degraded+remapped+wait_backfill
```

4. OSD에 대한 인증 키를 제거합니다.

```
[root@pod1-osd-compute-3 heat-admin]# ceph auth del osd.7
updated
```

5. OSD.7의 키가 나열되지 않는지 확인합니다.

```
[root@pod1-osd-compute-3 heat-admin]# ceph auth list
installed auth entries:

osd.0
  key: AQCgpB5b1V9dNhAAzDN1SVdnuJyTN2f7PAdtFw==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.1
  key: AQBdwyBbbuD6IBAAcvG+oQ0z5vk62fa0qv/CEw==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.10
  key: AQCwwyBb7xvHJhAAZKPprXWT7UnvnAXBV9W2rg==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.11
  key: AQDxpB5b9/rGFRAAkCEkpSN1YZVDdeW+Bho7w==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.2
  key: AQCpB5btekoNBAAACoWpDz0VL9bZfyIygDpBQ==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.3
  key: AQC4pB5bBaU10RAAhi3KPzetwvWhYGnerAkAsg==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.4
  key: AQB1wyBbvMIQLRAAXefFVnZxMX61Vt0bQt9KoA==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.5
  key: AQDBpB5buKHq0hAAW1Q861qoYqW6fAYH10xsLg==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.6
  key: AQDQpB5b1BveFxAafCLM3tvDUSnYneutyTmaEg==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.8
  key: AQDZpB5bd4n1GRAAkzbmGPnEDAWV0dUhrhE6w==
  caps: [mon] allow profile osd
  caps: [osd] allow *

osd.9
  key: AQDopB5bKCZPGBAafYtp1GLA7QIi/YxJa801yw==
  caps: [mon] allow profile osd
  caps: [osd] allow *

client.admin
  key: AQDpmx5bAAAAABAA3hLK802tGgaAK+X2L1y5Aw==
  caps: [mds] allow *
  caps: [mon] allow *
  caps: [osd] allow *
```

```

client.bootstrap-mds
    key: AQBDpB5bjR1GJhAAB6CKKxXu1ve9WIiC6ZGXgA==
    caps: [mon] allow profile bootstrap-mds
client.bootstrap-osd
    key: AQDpmx5bAAAAABAA3hLK802tGgaAK+X2L1y5Aw==
    caps: [mon] allow profile bootstrap-osd
client.bootstrap-rgw
    key: AQBDpB5b70WXHBAA1ATmBAOX/QWw+2mLxPq1kQ==
    caps: [mon] allow profile bootstrap-rgw
client.openstack
    key: AQDpmx5bAAAAABAAULxfs9cYG1wkSVTjrtiaDg==
    caps: [mon] allow r
    caps: [osd] allow class-read object_prefix rbd_children, allow rwx pool=volumes, allow rwx pool=

```

7. 클러스터에서 OSD를 제거합니다.

```

[root@pod1-osd-compute-3 heat-admin]# ceph osd rm 7
removed osd.7

```

8. 교체해야 하는 OSD 디스크를 마운트 해제합니다.

```

[root@pod1-osd-compute-3 heat-admin]# umount /var/lib/ceph/osd/ceph-7

```

9. noscrub 및 deepscrub을 설정 취소합니다.

```

[root@pod1-osd-compute-3 heat-admin]# ceph osd unset noscrub
unset noscrub

```

```

[root@pod1-osd-compute-3 heat-admin]# ceph osd unset nodeep-scrub
unset nodeep-scrub

```

10. Ceph 상태를 확인하고 정상 상태 및 모든 PG가 active+clean 상태로 돌아올 때까지 기다립니다

```

.

[root@pod1-osd-compute-3 heat-admin]# ceph -s
cluster eb2bb192-b1c9-11e6-9205-525400330666
health HEALTH_WARN
    28 pgs backfill_wait
    4 pgs backfilling
    5 pgs degraded
    5 pgs recovery_wait
    83 pgs stuck_unclean
    recovery 1697/516881 objects degraded (0.328%)
    recovery 76428/516881 objects misplaced (14.786%)

```



```

noout,norebalance,sortbitwise,require_jewel_osds flag(s) set
1 mons down, quorum 0,1 pod1-controller-0,pod1-controller-1
monmap e1: 3 mons at {pod1-controller-0=11.118.0.10:6789/0,pod1-controller-1=11.118.0.11:6789/0,pod1-controller-2=11.118.0.12:6789/0}
election epoch 28, quorum 0,1 pod1-controller-0,pod1-controller-1
osdmap e877: 11 osds: 11 up, 11 in; 193 remapped pgs
flags noout,norebalance,sortbitwise,require_jewel_osds
pgmap v942974: 704 pgs, 6 pools, 490 GB data, 163 kobjects
1470 GB used, 10806 GB / 12277 GB avail
1697/516881 objects degraded (0.328%)
76428/516881 objects misplaced (14.786%)
511 active+clean
156 active+remapped
28 active+remapped+wait_backfill
5 active+recovery_wait+degraded+remapped
4 active+remapped+backfilling
client io 331 kB/s wr, 0 op/s rd, 56 op/s wr

```

OSD 디스크 교체 및 새 VD 생성

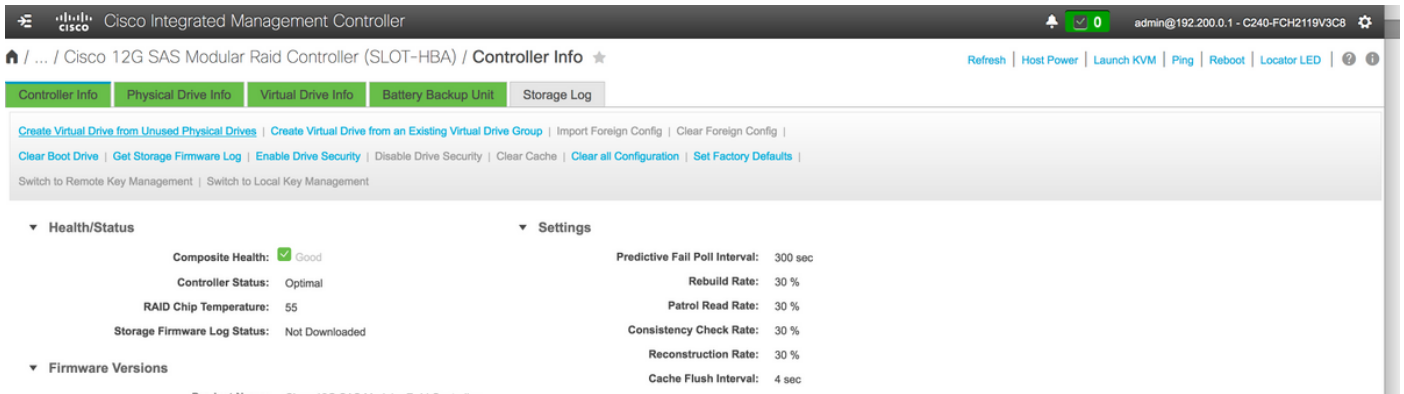
1. 결함이 있는 드라이브를 제거하고 새 드라이브로 교체합니다. [Cisco UCS C240 M4 Server 설치 및 서비스 가이드](#).
2. OSD-Compute의 CIMC에 로그인했는지 확인하고 OSD가 교체되어 상태가 양호한 슬롯을 확인합니다.
3. 새 HDD에 대한 가상 드라이브를 만듭니다. 메타데이터가 없는 새 HDD여야 합니다.
4. 새로 추가된 디스크가 Unconimaged Good 상태인지 확인합니다.

The screenshot shows the Cisco Integrated Management Controller (CIMC) interface for a Cisco 12G SAS Modular Raid Controller (SLOT-HBA). The 'Physical Drive Info' tab is selected, displaying a table of physical drives. The table has columns for Controller, Physical Drive Number, Status, Health, Boot Drive, Drive Firmware, Coerced Size, Model, and Type. The drives are currently in 'Online' or 'Unconfigured Good' state.

Controller	Physical Drive Number	Status	Health	Boot Drive	Drive Firmware	Coerced Size	Model	Type
☐ SLOT-HBA	1	Online	Good	false	5704	1143455 MB	TOSHIBA	HDD
☐ SLOT-HBA	2	Online	Good	false	5704	1143455 MB	TOSHIBA	HDD
☐ SLOT-HBA	3	Online	Good	false	CS01	456809 MB	ATA	SSD
☐ SLOT-HBA	7	Online	Good	false	N004	1143455 MB	SEAGATE	HDD
☐ SLOT-HBA	8	Online	Good	false	N004	1143455 MB	SEAGATE	HDD
☐ SLOT-HBA	9	Unconfigured Good	Good	false	N004	1143455 MB	SEAGATE	HDD
☐ SLOT-HBA	10	Online	Good	false	N004	1143455 MB	SEAGATE	HDD

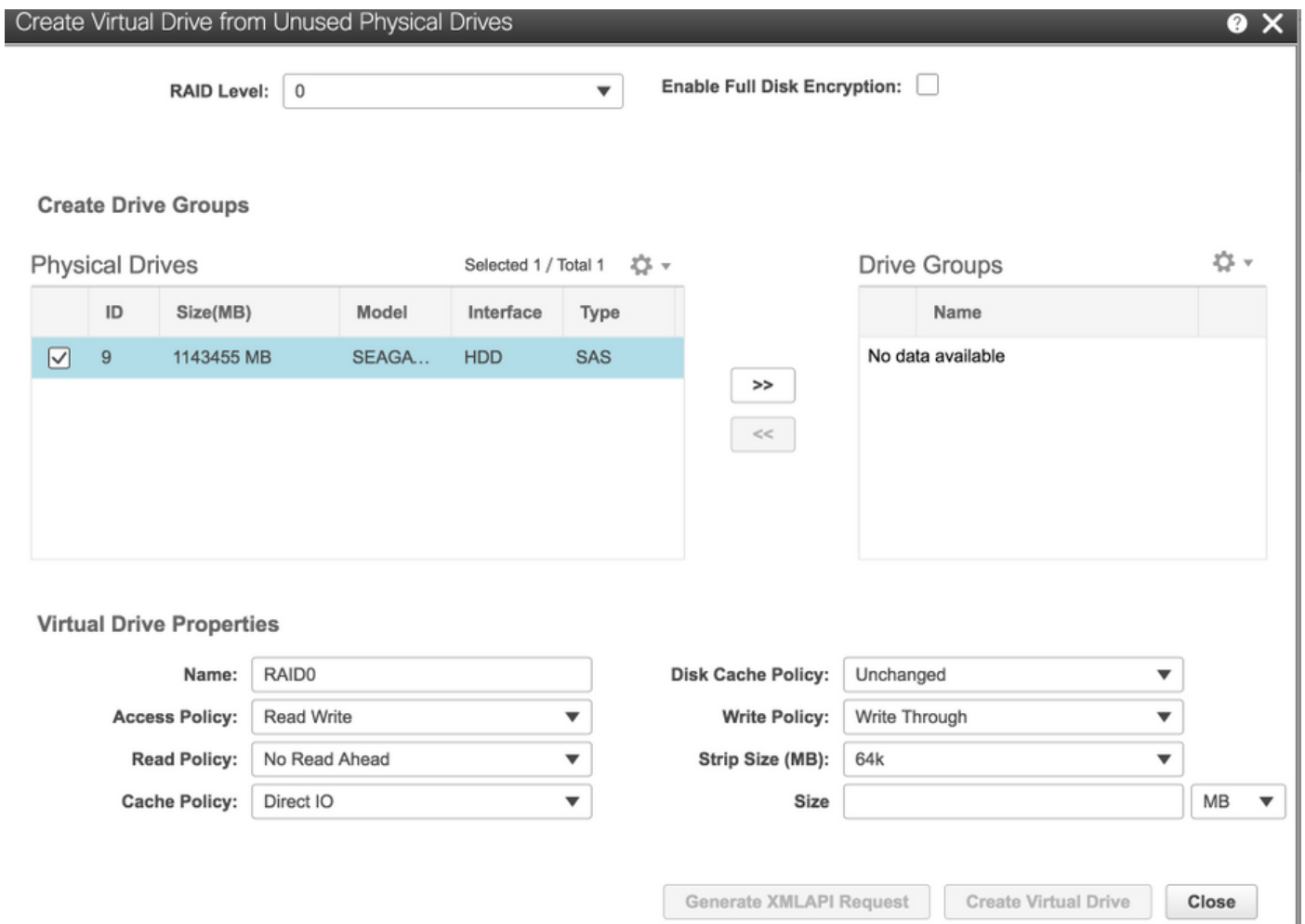
스토리지 > Cisco 12G SAS 모듈형 Raid 컨트롤러(SLOT-HBA) > 물리적 드라이브 정보

5. VD를 생성하려면 Create Virtual Drive from Unused Physical Drives(사용되지 않은 물리적 드라이브에서 가상 드라이브 생성) 옵션을 선택합니다.



스토리지 > Cisco 12G SAS 모듈형 Raid 컨트롤러(SLOT-HBA)

6. 새 VD를 생성하고 이름을 OSD3으로 지정하려면 Physical Drive 9를 사용합니다.



스토리지 > Cisco 12G SAS 모듈형 Raid 컨트롤러(SLOT-HBA) > 컨트롤러 정보 > 사용되지 않은 물리적 드라이브에서 가상 드라이브 생성

Create Virtual Drive from Unused Physical Drives

RAID Level: 0 Enable Full Disk Encryption: ☐

Create Drive Groups

Physical Drives
Selected 0 / Total 0

ID	Size(MB)	Model	Interface	Type
No data available				

>><<

Drive Groups

Name
☐ DG [9]

Virtual Drive Properties

Name: OSD3
Access Policy: Read Write
Read Policy: No Read Ahead
Cache Policy: Direct IO

Disk Cache Policy: Unchanged
Write Policy: Write Through
Strip Size (MB): 64k
Size: 1143455 MB

Generate XMLAPI Request
Create Virtual Drive
Close

스토리지 > Cisco 12G SAS 모듈형 Raid 컨트롤러(SLOT-HBA) > 컨트롤러 정보 > 사용되지 않은 물리적 드라이브에서 가상 드라이브 생성

7. IPMI over LAN을 활성화합니다. Admin > Communication Services > Communication Services.

Cisco Integrated Management Controller

admin@10.65.33.67 - C240-FCH2141V113

Refresh Host Power Launch KVM Ping Reboot Locator LED

Communications Services SNMP Mail Alert

HTTP Properties

IPMI over LAN Properties

HTTP/S Enabled: ☒
Redirect HTTP to HTTPS Enabled: ☒
HTTP Port: 80
HTTPS Port: 443

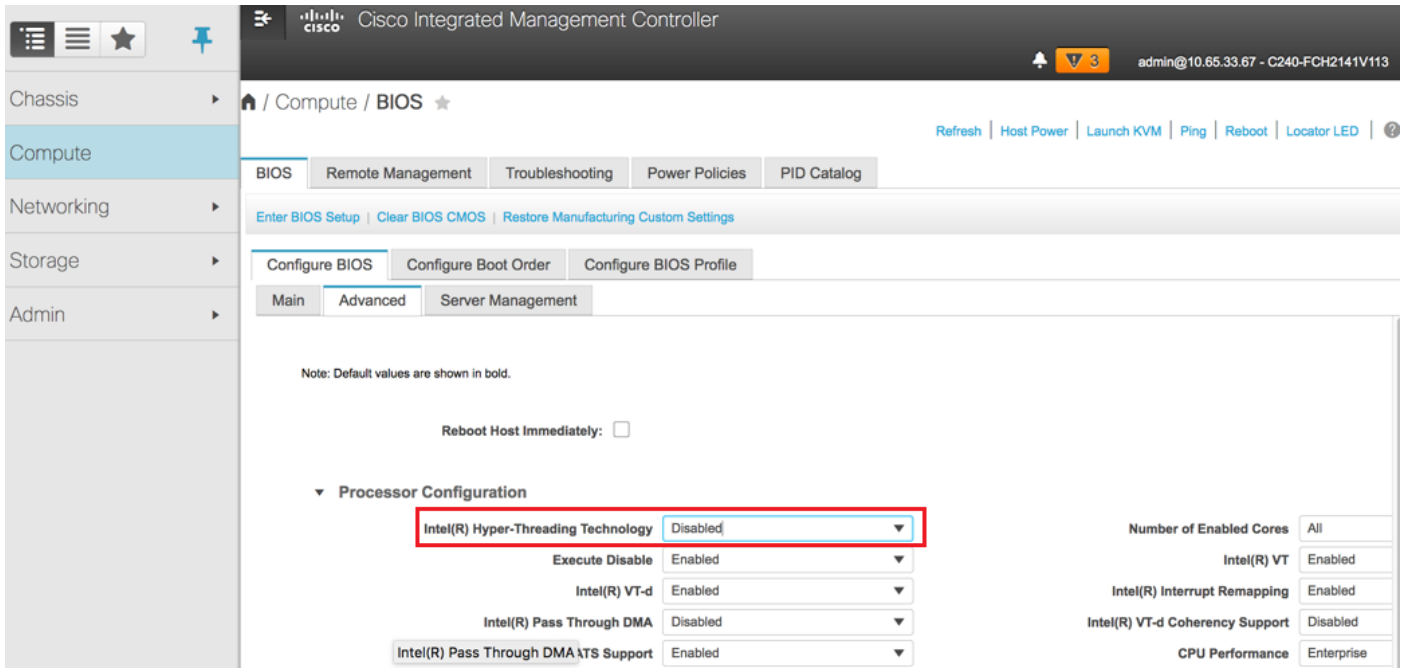
Session Timeout(seconds): 1800
Max Sessions: 4
Active Sessions: 1

Enabled: ☒
Privilege Level Limit: admin
Encryption Key: 00000000000000000000000000000000
Randomize

XML API Properties
XML API Enabled: ☒

IPMI over LAN 활성화: Admin(관리) > Communication Services(통신 서비스) > Communication Services(통신 서비스)

8. 하이퍼스레딩을 비활성화합니다. Compute(컴퓨팅) > BIOS > Conimage BIOS > Advanced(고급) > Processor Configuration(프로세서 컨피그레이션)



하이퍼스레딩을 비활성화합니다. Compute(컴퓨팅) > BIOS > Configure BIOS > Advanced(고급) > Processor Configuration(프로세서 컨피그레이션)



참고: 여기에 표시된 이미지 및 이 섹션에서 설명한 컨피그레이션 단계는 펌웨어 버전 3.0(3e)을 참조하며, 다른 버전에서 작업하는 경우 약간의 차이가 있을 수 있습니다.

클러스터에 OSD 다시 추가

1. 새 디스크를 교체한 후 새 장치를 검색하기 위해 partprobe를 실행합니다.

```
[root@pod1-osd-compute-3 heat-admin]# partprobe
[root@pod1-osd-compute-3 heat-admin]# lsblk
NAME        MAJ:MIN RM  SIZE RO TYPE MOUNTPOINT
sda           8:0    0 278.5G  0 disk
|
-sda1        8:1     0    1M  0 part
-sda2        8:2     0 278.5G  0 part /
sdb           8:16    0 446.1G  0 disk
|
-sdb1        8:17    0   107G  0 part
-sdb2        8:18    0   107G  0 part
-sdb3        8:19    0   107G  0 part
-sdb4        8:20    0   107G  0 part
sdc           8:32    0   1.1T  0 disk
|
-sdc1        8:33    0   1.1T  0 part /var/lib/ceph/osd/ceph-1
sdd 8:48     0   1.1T  0 disk
|
-sdd1        8:49    0   1.1T  0 part
sde           8:64    0   1.1T  0 disk
|
-sde1        8:65    0   1.1T  0 part /var/lib/ceph/osd/ceph-4
sdf           8:80    0   1.1T  0 disk
|
```

```
-sdf1 8:81 0 1.1T 0 part /var/lib/ceph/osd/ceph-10
```

2. 서버에서 사용할 수 있는 장치를 찾습니다.

```
[root@pod1-osd-compute-3 heat-admin]# fdisk -l
```

```
Disk /dev/sda: 299.0 GB, 298999349248 bytes, 583983104 sectors
Units = sectors of 1 * 512 = 512 bytes
Sector size (logical/physical): 512 bytes / 512 bytes
I/O size (minimum/optimal): 512 bytes / 512 bytes
Disk label type: dos
Disk identifier: 0x000b5e87
```

Device	Boot	Start	End	Blocks	Id	System
/dev/sda1		2048	4095	1024	83	Linux
/dev/sda2	*	4096	583983070	291989487+	83	Linux

WARNING: fdisk GPT support is currently new, and therefore in an experimental phase. Use at your own di

```
Disk /dev/sdb: 479.0 GB, 478998953984 bytes, 935544832 sectors
Units = sectors of 1 * 512 = 512 bytes
Sector size (logical/physical): 512 bytes / 4096 bytes
I/O size (minimum/optimal): 4096 bytes / 4096 bytes
Disk label type: gpt
```

#	Start	End	Size	Type	Name
1	2048	224462847	107G	unknown	ceph journal
2	224462848	448923647	107G	unknown	ceph journal
3	448923648	673384447	107G	unknown	ceph journal
4	673384448	897845247	107G	unknown	ceph journal

WARNING: fdisk GPT support is currently new, and therefore in an experimental phase. Use at your own di

```
Disk /dev/sdd: 1199.0 GB, 1198999470080 bytes, 2341795840 sectors
Units = sectors of 1 * 512 = 512 bytes
Sector size (logical/physical): 512 bytes / 512 bytes
I/O size (minimum/optimal): 512 bytes / 512 bytes
Disk label type: gpt
```

#	Start	End	Size	Type	Name
1	2048	2341795806	1.1T	unknown	ceph data

WARNING: fdisk GPT support is currently new, and therefore in an experimental phase. Use at your own di

```
Disk /dev/sdc: 1199.0 GB, 1198999470080 bytes, 2341795840 sectors
Units = sectors of 1 * 512 = 512 bytes
Sector size (logical/physical): 512 bytes / 512 bytes
I/O size (minimum/optimal): 512 bytes / 512 bytes
Disk label type: gpt
```

#	Start	End	Size	Type	Name
1	2048	2341795806	1.1T	unknown	ceph data

WARNING: fdisk GPT support is currently new, and therefore in an experimental phase. Use at your own di

```
Disk /dev/sde: 1199.0 GB, 1198999470080 bytes, 2341795840 sectors
Units = sectors of 1 * 512 = 512 bytes
Sector size (logical/physical): 512 bytes / 512 bytes
I/O size (minimum/optimal): 512 bytes / 512 bytes
```

Disk label type: gpt

#	Start	End	Size	Type	Name
1	2048	2341795806	1.1T	unknown	ceph data

WARNING: fdisk GPT support is currently new, and therefore in an experimental phase. Use at your own di

Disk /dev/sdf: 1199.0 GB, 1198999470080 bytes, 2341795840 sectors

Units = sectors of 1 * 512 = 512 bytes

Sector size (logical/physical): 512 bytes / 512 bytes

I/O size (minimum/optimal): 512 bytes / 512 bytes

Disk label type: gpt

#	Start	End	Size	Type	Name
1	2048	2341795806	1.1T	unknown	ceph data

[root@pod1-osd-compute-3 heat-admin]#

3. 저널 디스크 파티션 맵을 식별하려면 Ceph-disk 목록을 사용합니다.

<#root>

[root@pod1-osd-compute-3 heat-admin]# ceph-disk list

/dev/sda :

/dev/sda1 other, iso9660

/dev/sda2 other, xfs, mounted on /

/dev/sdb :

/dev/sdb1 ceph journal, for /dev/sdc1

/dev/sdb3 ceph journal

/dev/sdb2 ceph journal, for /dev/sde1

/dev/sdb4 ceph journal, for /dev/sdf1

/dev/sdc :

/dev/sdc1 ceph data, active, cluster ceph, osd.1, journal /dev/sdb1

/dev/sdd :

/dev/sdd1 other, xfs

/dev/sde :

/dev/sde1 ceph data, active, cluster ceph, osd.4, journal /dev/sdb2

/dev/sdf :

/dev/sdf1 ceph data, active, cluster ceph, osd.10, journal /dev/sdb4

 참고: ceph-disk 목록에서 강조 표시된 sde1 출력은 sdb2에 대한 저널 파티션입니다. Ceph-disk 목록의 출력을 확인하고 Ceph 준비에 대한 명령에서 저널 디스크 파티션을 매핑하십시오. 아래 명령 OSD.7을 실행하는 즉시 시작/입력되며 데이터 리밸런싱(백필/복구)이 시작됩니다.

4. Ceph-disk를 생성하여 클러스터에 다시 추가합니다.

[root@pod1-osd-compute-3 heat-admin]# ceph-disk --setuser ceph --setgroup ceph prepare --fs-type xfs /

```

prepare_device: OSD will not be hot-swappable if journal is not the same device as the osd data
Creating new GPT entries.
The operation has completed successfully.
meta-data=/dev/sdd1      isize=2048    agcount=4, agsize=73181055 blks
              =          sectsz=512    attr=2, projid32bit=1
              =          crc=1         finobt=0, sparse=0
data       =            bsize=4096    blocks=292724219, imaxpct=5
              =            sunit=0     swidth=0 blks
naming     =version 2       bsize=4096    ascii-ci=0 ftype=1
log        =internal log    bsize=4096    blocks=142931, version=2
              =            sectsz=512   sunit=0 blks, lazy-count=1
realtime   =none           extsz=4096   blocks=0, rtextents=0
Warning: The kernel is still using the old partition table.
The new table will be used at the next reboot.
The operation has completed successfully.

```

#####Hint###

where - sdd is new drive added as OSD

where - sdb3 is journal disk partition number

mapping is sdc1 for sdc, sdd1 for sdd, sde1 for sde

sdf1 for sdf (and so on)

5. Ceph-disks를 활성화하고 noscrub 및 nodeep-scrub 플래그를 설정 취소합니다.

```

[root@pod1-osd-compute-3 heat-admin]# ceph-disk activate-all
[root@pod1-osd-compute-3 heat-admin]# ceph osd unset noout
unset noout
[root@pod1-osd-compute-3 heat-admin]# ceph osd unset norebalance
unset norebalance
[root@pod1-osd-compute-3 heat-admin]# ceph osd unset noscrub
unset noscrub
[root@pod1-osd-compute-3 heat-admin]# ceph osd unset nodeep-scrub
unset nodeep-scrub

```

6. 리밸런스가 끝날 때까지 기다린 후 Ceph 및 OSD 트리의 상태가 정상인지 확인합니다.

```

[root@pod1-osd-compute-3 heat-admin]# watch -n 3 ceph -s

```

```

[heat-admin@pod1-osd-compute-3 ~]$ sudo ceph -s
cluster eb2bb192-b1c9-11e6-9205-525400330666
health HEALTH_OK
1 mons down, quorum 0,1 pod1-controller-0,pod1-controller-1
monmap e1: 3 mons at {pod1-controller-0=11.118.0.10:6789/0,pod1-controller-1=11.118.0.11:6789/0,pod1-controller-2=11.118.0.12:6789/0}
election epoch 28, quorum 0,1 pod1-controller-0,pod1-controller-1

```

```

osdmap e709: 12 osds: 12 up, 12 in
      flags sortbitwise,require_jewel_osds
pgmap v941813: 704 pgs, 6 pools, 490 GB data, 163 kobjects
      1470 GB used, 11922 GB / 13393 GB avail
      704 active+clean
client io 58580 B/s wr, 0 op/s rd, 7 op/s wr

```

```
[heat-admin@pod1-osd-compute-3 ~]$ sudo ceph osd tree
```

ID	WEIGHT	TYPE	NAME	UP/DOWN	REWEIGHT	PRIMARY-AFFINITY
-1	13.07996	root	default			
-2	4.35999	host	pod1-osd-compute-0			
0	1.09000		osd.0	up	1.00000	1.00000
3	1.09000		osd.3	up	1.00000	1.00000
6	1.09000		osd.6	up	1.00000	1.00000
9	1.09000		osd.9	up	1.00000	1.00000
-4	4.35999	host	pod1-osd-compute-2			
2	1.09000		osd.2	up	1.00000	1.00000
5	1.09000		osd.5	up	1.00000	1.00000
8	1.09000		osd.8	up	1.00000	1.00000
11	1.09000		osd.11	up	1.00000	1.00000
-5	4.35999	host	pod1-osd-compute-3			
1	1.09000		osd.1	up	1.00000	1.00000
4	1.09000		osd.4	up	1.00000	1.00000
7	1.09000		osd.7	up	1.00000	1.00000
10	1.09000		osd.10	up	1.00000	1.00000

이 번역에 관하여

Cisco는 전 세계 사용자에게 다양한 언어로 지원 콘텐츠를 제공하기 위해 기계 번역 기술과 수작업 번역을 병행하여 이 문서를 번역했습니다. 아무리 품질이 높은 기계 번역이라도 전문 번역가의 번역 결과물만큼 정확하지는 않습니다. Cisco Systems, Inc.는 이 같은 번역에 대해 어떠한 책임도 지지 않으며 항상 원본 영문 문서(링크 제공됨)를 참조할 것을 권장합니다.