



从灾难中恢复

ONTAP MetroCluster

NetApp
February 13, 2026

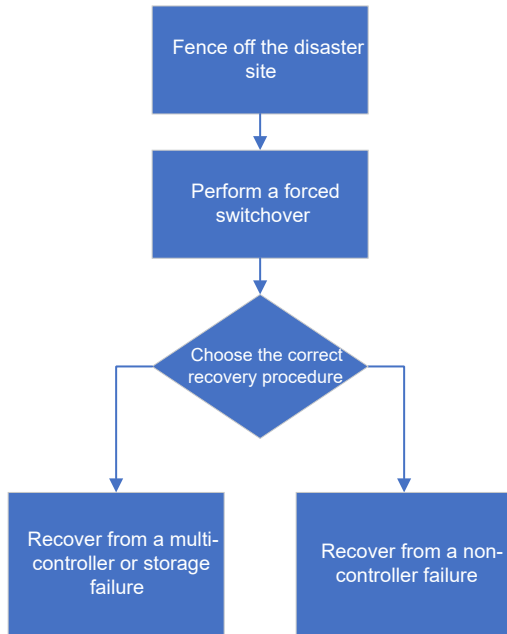
目录

从灾难中恢复	1
灾难恢复工作流	1
在发生灾难后执行强制切换	1
隔离灾难站点	1
执行强制切换	2
在 MetroCluster 切换后， storage aggregate plex show 命令的输出不确定	3
在切换后访问处于 NVFAIL 状态的卷	3
选择正确的恢复操作步骤	3
MetroCluster 安装期间的控制器模块故障场景	4
MetroCluster FC-IP过渡期间的控制器模块故障场景	4
八节点 MetroCluster 配置中的控制器模块故障情形	5
双节点 MetroCluster 配置中的控制器模块故障情形	8
从多控制器或存储故障中恢复	9
从多控制器或存储故障中恢复	9
启用控制台日志记录	10
更换硬件并启动新控制器	10
准备在 MetroCluster IP 配置中切回	21
准备在 MetroCluster FC 配置中切回	65
在混合配置中准备切回（过渡期间恢复）	92
正在完成恢复	95
从非控制器故障中恢复	108
启用控制台日志记录	108
修复MetroCluster配置中的配置	109
验证您的系统是否已做好切回准备	111
执行切回	113
验证切回是否成功	115
在切回后删除陈旧的聚合列表	117

从灾难中恢复

灾难恢复 workflow

使用 workflow 执行灾难恢复。



在发生灾难后执行强制切换

如果发生灾难，则必须在切换后对灾难集群和正常运行的集群执行一些步骤，以确保安全，持续地提供数据服务。

确定是否发生灾难的方法是：

- 管理员
- MetroCluster Tiebreaker 软件（如果已配置）
- ONTAP 调解器软件（如果已配置）

隔离灾难站点

发生灾难后，如果必须更换灾难站点节点，则必须暂停这些节点，以防止站点恢复服务。否则，如果客户端在替换操作步骤完成之前开始访问节点，您将面临数据损坏的风险。

步骤

1. 暂停灾难站点上的节点并使其保持关闭状态或停留在 LOADER 提示符处，直到指示启动 ONTAP 为止：

```
ssystem node halt -node disaster -site-node-name
```

如果灾难站点节点已被销毁或无法暂停，请关闭节点的电源，在恢复操作步骤中指示之前不要启动替代节点。

执行强制切换

除了在测试和维护期间提供无中断操作之外，切换过程还允许您通过一条命令从站点故障中恢复。

开始之前

- 在执行切换之前、必须至少有一个运行正常的站点节点已启动且正在运行。
- 在执行切回操作之前，必须完成所有先前的配置更改。

这是为了避免与协商切换或切回操作产生竞争。



系统会自动删除 SnapMirror 和 SnapVault 配置。

关于此任务

MetroCluster switchover 命令可切换 MetroCluster 配置中所有 DR 组中的节点。例如，在八节点 MetroCluster 配置中，它会切换两个 DR 组中的节点。

步骤

1. 在运行正常的站点上运行以下命令、以执行切换：

```
MetroCluster switchover -forced-on-disaster true
```



此操作可能需要几分钟时间才能完成。您可以使用验证进度 metrocluster operation show 命令：

2. 系统提示您继续切换时，请选择问题解答 y。
3. 运行以验证切换是否已成功完成 metrocluster operation show 命令：

```
mccl1A::> metrocluster operation show
Operation: switchover
Start time: 10/4/2012 19:04:13
State: in-progress
End time: -
Errors:

mccl1A::> metrocluster operation show
Operation: switchover
Start time: 10/4/2012 19:04:13
State: successful
End time: 10/4/2012 19:04:22
Errors: -
```

如果切换被否决，您可以使用 ` -override-vetoes ` 选项重新发出 MetroCluster switchover-forced-

`on-disaster true` 命令。如果您使用此可选参数，则系统将覆盖阻止切换的任何软否决。

完成后
需要在切换后重新建立 SnapMirror 关系。

在 **MetroCluster** 切换后， **storage aggregate plex show** 命令的输出不确定

在 MetroCluster 切换后运行 `storage aggregate plex show` 命令时，切换后的根聚合的 `plex0` 状态不确定，并显示为 `failed`。在此期间，切换后的根不会更新。只有在 MetroCluster 修复阶段之后才能确定此丛的实际状态。

在切换后访问处于 **NVFAIL** 状态的卷

切换后，您必须通过重置 `volume modify` 命令的 ``in-nvfailed-state`` 参数来清除 NVFAIL 状态，以取消客户端访问数据的限制。

开始之前
数据库或文件系统不得正在运行或尝试访问受影响的卷。

关于此任务
设置 ``in-nvfailed-state`` 参数需要高级权限。

- 步骤
1. 使用 `volume modify` 命令并将 ``in-nvfailed-state`` 参数设置为 `false`，以恢复卷。

完成后
有关检查数据库文件有效性的说明，请参见特定数据库软件的文档。

如果数据库使用 LUN，请查看相关步骤，以便在 NVRAM 出现故障后使主机可以访问这些 LUN。

相关信息
["使用 NVFAIL 监控和保护数据库有效性"](#)

选择正确的恢复操作步骤

在 MetroCluster 配置出现故障后，您必须选择正确的恢复操作步骤。使用下表和示例选择相应的恢复操作步骤。

此表中的此信息假定安装或过渡已完成、即 `metrocluster configure` 命令已成功运行。

灾难站点的故障范围	操作步骤
• 无硬件故障（例如电源故障）	"从非控制器故障中恢复"
• 无控制器模块故障 • 其他硬件出现故障	"从非控制器故障中恢复"

• 单控制器模块故障或控制器模块中的 FRU 组件出现故障 • 驱动器未出现故障	如果故障仅限于单个控制器模块，则必须使用适用于此平台型号的控制器模块 FRU 更换操作步骤。在四节点或八节点 MetroCluster 配置中，此类故障将被隔离到本地 HA 对。 • 注：* 如果没有驱动器或其他硬件故障，则可以在双节点 MetroCluster 配置中使用控制器模块 FRU 更换操作步骤。 "ONTAP硬件系统文档"
• 单控制器模块故障或控制器模块中的 FRU 组件出现故障 • 驱动器出现故障	"从多控制器或存储故障中恢复"
• 单控制器模块故障或控制器模块中的 FRU 组件出现故障 • 驱动器未出现故障 • 控制器模块外部的其他硬件出现故障	"从多控制器或存储故障中恢复" 您应跳过驱动器分配的所有步骤。
• 一个 DR 组中的多个控制器模块出现故障（存在或不存在其他故障）	"从多控制器或存储故障中恢复"

MetroCluster 安装期间的控制器模块故障场景

在 MetroCluster 配置操作步骤 期间响应控制器模块故障取决于是否 `metrocluster configure` 命令成功完成。

- 如果 `metrocluster configure` 命令尚未运行或失败、您必须从头使用替代控制器模块重新启动 MetroCluster 软件配置操作步骤。



您必须确保执行中的步骤 ["还原控制器模块上的系统默认值"](#) 在每个控制器(包括替代控制器)上、验证先前的配置是否已删除。

- 如果 `metrocluster configure` 命令成功完成、然后控制器模块出现故障、请使用上表确定正确的恢复操作步骤。

MetroCluster FC-IP过渡期间的控制器模块故障场景

如果过渡期间发生站点故障，可以使用恢复操作步骤。但是，只有在配置为稳定的混合配置且 FC DR 组和 IP DR 组均已完全配置的情况下，才能使用此配置。MetroCluster `node show` 命令的输出应显示两个 DR 组以及全部八个节点。



如果在过渡期间添加或删除节点时发生故障，您必须联系技术支持。

八节点 MetroCluster 配置中的控制器模块故障情形

故障情形：

- 单个 DR 组中的单个控制器模块出现故障
- 一个 DR 组中的两个控制器模块出现故障
- 单独 DR 组中的单个控制器模块出现故障
- 三个控制器模块故障分布在 DR 组中

单个 **DR** 组中的单个控制器模块出现故障

在这种情况下，故障仅限于 HA 对。

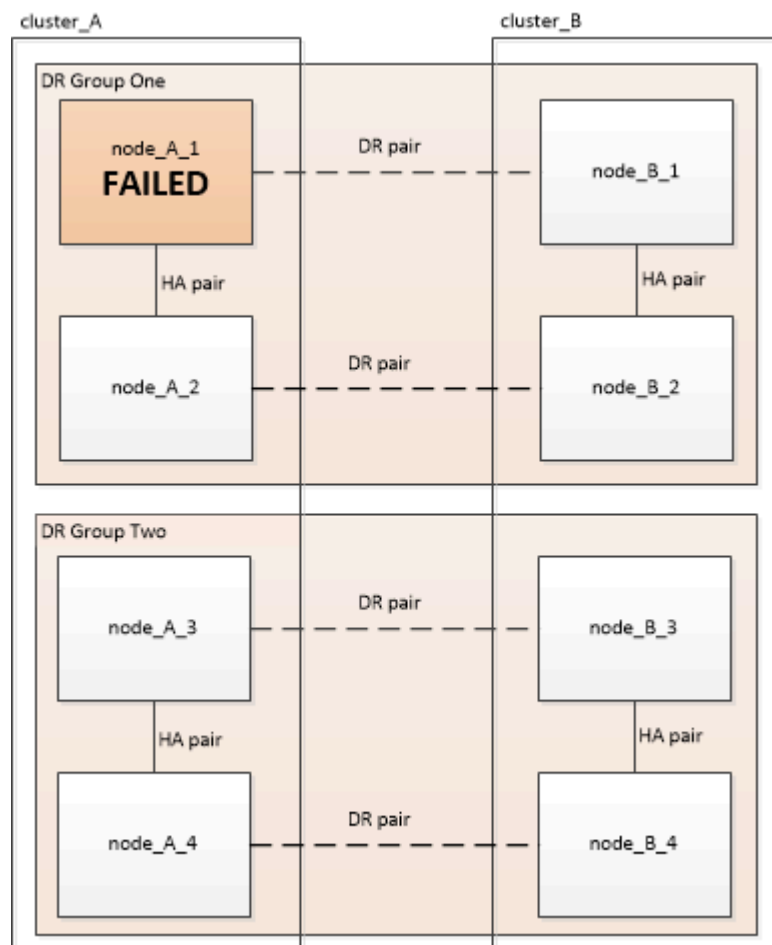
- 如果没有需要更换的存储，您可以使用适用于此平台型号的控制器模块 FRU 更换操作步骤。

["ONTAP 硬件系统文档"](#)

- 如果存储需要更换，您可以使用多控制器模块恢复操作步骤。

["从多控制器或存储故障中恢复"](#)

这种情况下，适用场景四节点 MetroCluster 配置也是如此。

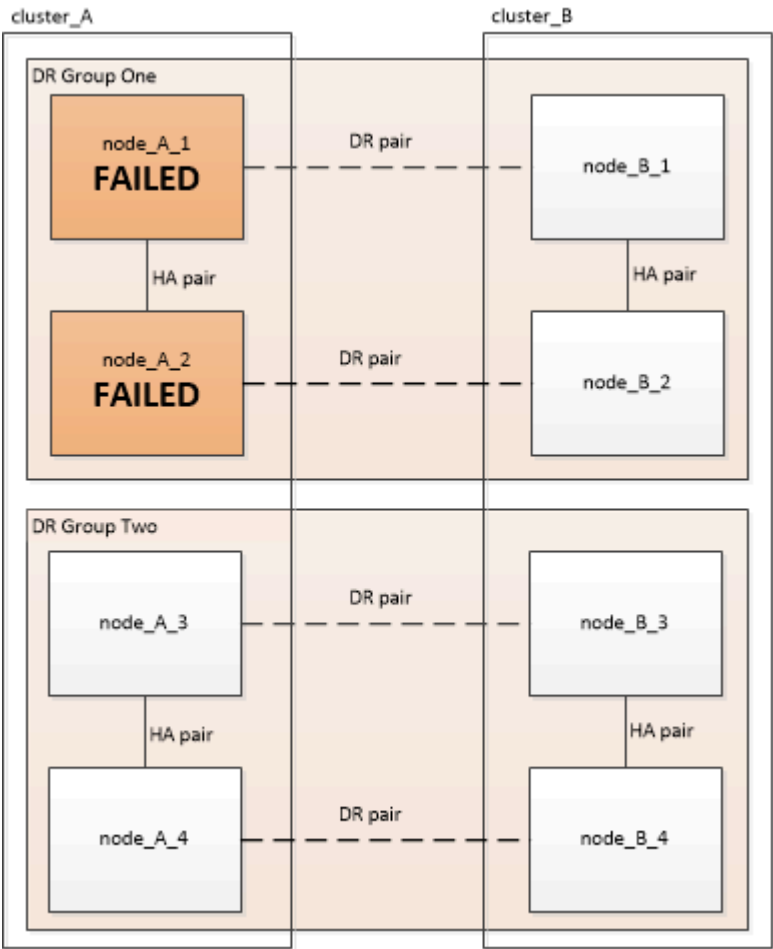


一个 DR 组中的两个控制器模块出现故障

在这种情况下，故障需要切换。您可以使用多控制器模块故障恢复操作步骤。

["从多控制器或存储故障中恢复"](#)

这种情况下，适用场景四节点 MetroCluster 配置也是如此。



单独 DR 组中的单个控制器模块出现故障

在这种情况下，故障仅限于单独的 HA 对。

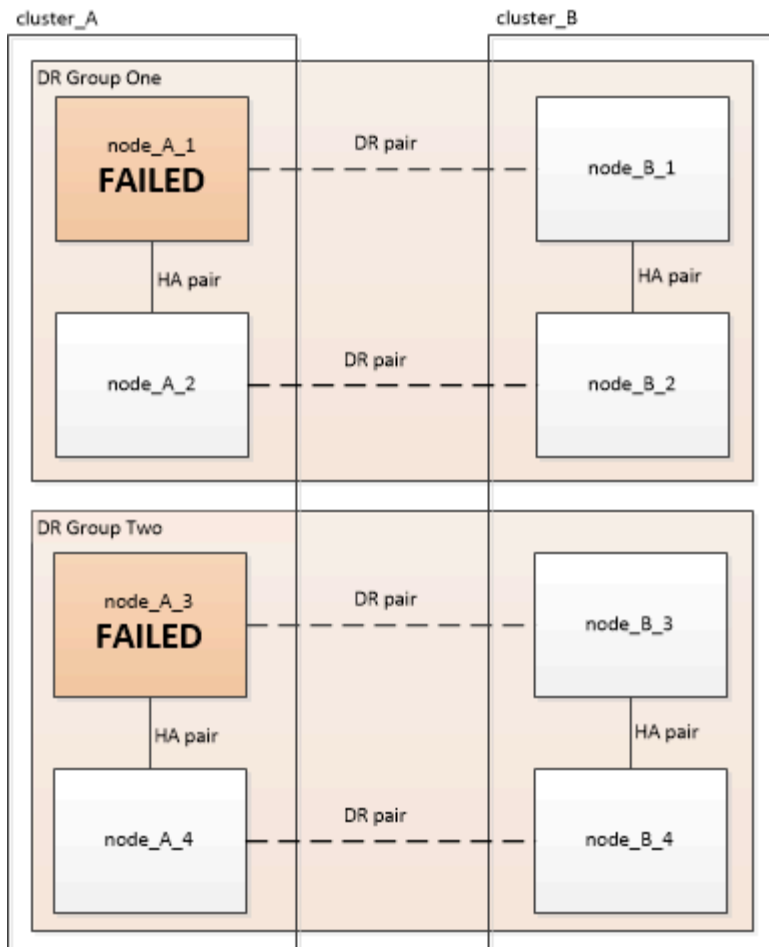
- 如果没有需要更换的存储，您可以使用适用于此平台型号的控制器模块 FRU 更换操作步骤。

FRU 更换操作步骤会执行两次，每个故障控制器模块一次。

["ONTAP硬件系统文档"](#)

- 如果存储需要更换，您可以使用多控制器模块恢复操作步骤。

["从多控制器或存储故障中恢复"](#)



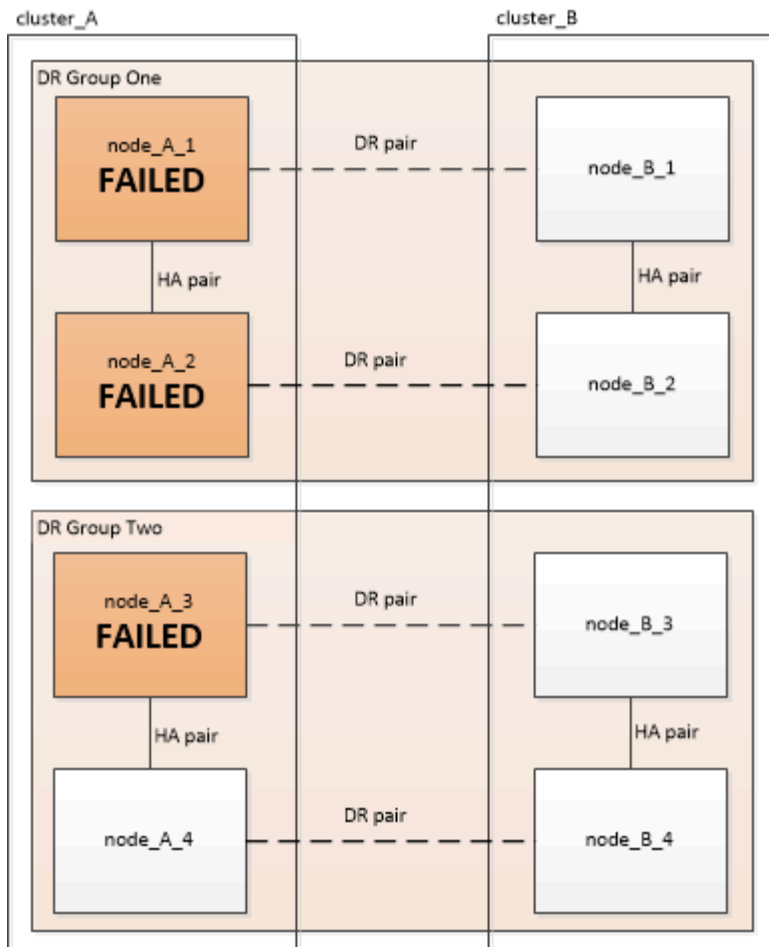
三个控制器模块故障分布在 **DR** 组中

在这种情况下，故障需要切换。您可以对 DR 组 1 使用多控制器模块故障恢复操作步骤。

["从多控制器或存储故障中恢复"](#)

您可以对 DR 组 2 使用平台专用的控制器模块 FRU 更换操作步骤。

["ONTAP硬件系统文档"](#)



双节点 MetroCluster 配置中的控制器模块故障情形

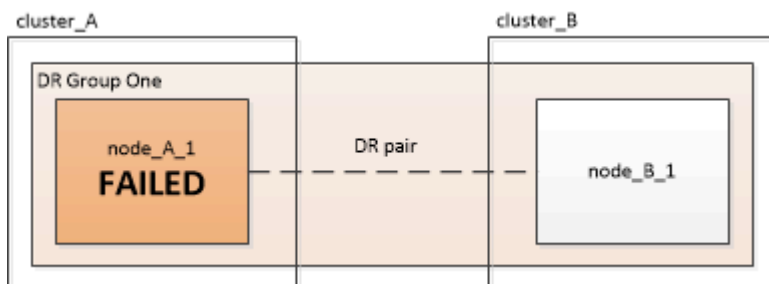
您使用的操作步骤取决于故障的程度。

- 如果没有需要更换的存储，您可以使用适用于此平台型号的控制器模块 FRU 更换操作步骤。

["ONTAP硬件系统文档"](#)

- 如果存储需要更换，您可以使用多控制器模块恢复操作步骤。

["从多控制器或存储故障中恢复"](#)



从多控制器或存储故障中恢复

从多控制器或存储故障中恢复

如果控制器故障扩展到 MetroCluster 配置中 DR 组一侧的所有控制器模块（包括双节点 MetroCluster 配置中的单个控制器），或者存储已被更换，则必须更换设备并重新分配驱动器的所有权，才能从灾难中恢复。

在使用此过程之前、请确认您已检查并执行以下任务：

- 在决定使用此过程之前、请查看可用的恢复过程。

["选择正确的恢复操作步骤"](#)

- 确认已在设备上启用控制台日志记录。

["启用控制台日志记录"](#)

- 确保灾难站点已隔离。

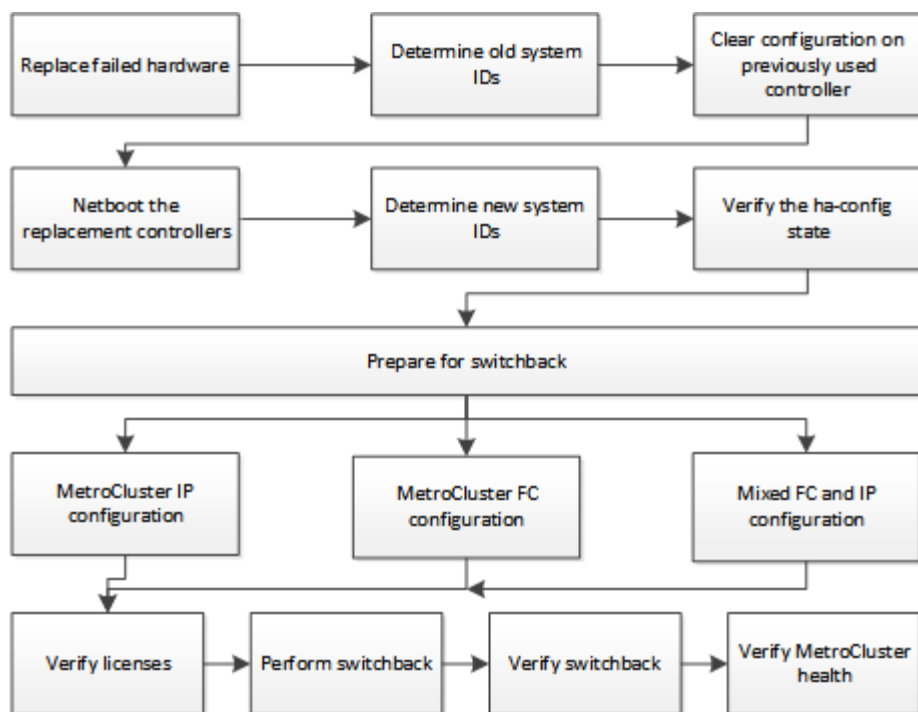
["隔离灾难站点"](#)。

- 验证是否已执行切换。

["执行强制切换"](#)。

- 确认更换的驱动器和控制器模块是新的、并且之前未向其分配所有权。
- 此操作步骤中的示例显示了两个或四节点配置。如果您使用的是八节点配置（两个 DR 组），则必须考虑所有故障，并对其他控制器模块执行所需的恢复任务。

此操作步骤使用以下工作流：



在发生故障时处于过渡中的系统上执行恢复时，可以使用此操作步骤。在这种情况下，您必须按照操作步骤中的说明在准备切回时执行相应的步骤。

启用控制台日志记录

在继续更换硬件和启动新控制器之前、请在设备上启用控制台日志记录。

NetApp强烈建议您在使用的设备上启用控制台日志记录、并在执行此过程时执行以下操作：

- 在维护期间保持AutoSupport处于启用状态。
- 在维护前后触发维护AutoSupport消息、以便在维护活动期间禁用案例创建。

请参阅知识库文章 ["如何在计划的维护时段禁止自动创建案例"](#)。

- 为任何命令行界面会话启用会话日志记录。有关如何启用会话日志记录的说明，请查看知识库文章中的“日志记录会话输出”部分 ["如何配置PuTTY以优化与ONTAP系统的连接"](#)。

更换硬件并启动新控制器

如果需要更换硬件组件，您必须使用其各自的硬件更换和安装指南进行更换。

更换灾难站点上的硬件

开始之前

存储控制器必须关闭电源或保持暂停状态（显示 LOADER 提示符）。

步骤

1. 根据需要更换组件。



在此步骤中，您可以完全按照发生灾难之前的布线方式更换组件并为其布线。您不能打开组件的电源。

如果要更换 ...	执行以下步骤 ...	使用这些指南 ...
MetroCluster FC 配置中的 FC 交换机	- 安装新交换机。 - 为 ISL 链路布线。此时请勿打开 FC 交换机的电源。	"维护 MetroCluster 组件"
MetroCluster IP 配置中的 IP 交换机	- 安装新交换机。 - 为 ISL 链路布线。此时请勿打开 IP 交换机的电源。	"MetroCluster IP 安装和配置：ONTAP MetroCluster 配置之间的差异"
磁盘架	- 安装磁盘架和磁盘。 - 磁盘架堆栈的配置应与正常运行的站点上的配置相同。 - 磁盘大小可以相同或更大，但类型必须相同（SAS 或 SATA）。 - 使用缆线将磁盘架连接到堆栈中的相邻磁盘架以及 FC-SAS 网桥。此时请勿打开磁盘架的电源。	"ONTAP 硬件系统文档"
SAS 缆线	安装新缆线。此时请勿打开磁盘架的电源。	"ONTAP 硬件系统文档"
MetroCluster FC 配置中的 FC-SAS 网桥	- 安装 FC-SAS 网桥。 - 为 FC-SAS 网桥布线。 根据您的 MetroCluster 配置类型，使用缆线将其连接到 FC 交换机或控制器模块。 此时请勿打开 FC-SAS 网桥的电源。	"光纤连接的 MetroCluster 安装和配置" "延伸型 MetroCluster 安装和配置"

控制器模块	a. 安装新控制器模块： ◦ 控制器模块必须与要更换的模块型号相同。 例如，8080 控制器模块必须更换为 8080 控制器模块。 ◦ 控制器模块先前不能属于 MetroCluster 配置中的任一集群，也不能属于先前已有的任何集群配置。 如果是，则必须设置默认值并执行 "wipeconfig" 过程。 ◦ 确保所有网络接口卡（例如以太网或 FC）均位于旧控制器模块上使用的相同插槽中。 b. 使用与旧控制器模块完全相同的缆线连接新控制器模块。 将控制器模块连接到存储（通过连接到 IP 或 FC 交换机，FC-SAS 网桥或直接连接）的端口应与发生灾难之前使用的端口相同。 此时请勿打开控制器模块的电源。	"ONTAP 硬件系统文档"
-------	--	--------------------------------

2. 验证所有组件是否已根据您的配置正确布线。

- ["MetroCluster IP 配置"](#)
- ["MetroCluster 光纤连接配置"](#)

确定旧控制器模块的系统ID和VLAN ID

在更换灾难站点上的所有硬件后，您必须确定已更换的控制器模块的系统 ID。在将磁盘重新分配给新控制器模块时，您需要旧系统 ID。如果系统为 AFF A220，AFF A250，AFF A400，AFF A800，FAS2750，对于 FAS500f，FAS8300 或 FAS8700 型号，您还必须确定 MetroCluster IP 接口使用的 VLAN ID。

开始之前

灾难站点上的所有设备都必须关闭。

关于此任务

本讨论提供了双节点和四节点配置的示例。对于八节点配置，您必须考虑第二个 DR 组上其他节点中的任何故障。

对于双节点 MetroCluster 配置，您可以忽略对每个站点上第二个控制器模块的引用。

此操作步骤中的示例基于以下假设：

- 站点 A 是灾难站点。
- node_A_1 出现故障，正在完全更换。
- node_A_2 出现故障，正在完全更换。

节点 _A_2 仅存在于四节点 MetroCluster 配置中。

- 站点 B 是正常运行的站点。
- node_B_1 运行状况良好。
- node_B_2 运行状况良好。

node_B_2 仅存在于四节点 MetroCluster 配置中。

控制器模块具有以下原始系统 ID：

MetroCluster 配置中的节点数	节点	原始系统 ID
四个	node_A_1	4068741258
node_A_2	4068741260	node_B_1
4068741254	node_B_2	4068741256
两个	node_A_1	4068741258

步骤

1. 在运行正常的站点中，显示 MetroCluster 配置中节点的系统 ID。

MetroCluster 配置中的节点数	使用此命令 ...
四个或八个	<code>MetroCluster node show -fields node-systemID , ha-partner-systemID , dr-partner-systemID , dr-auxiliary-systemID</code>
两个	<code>MetroCluster node show -fields node-systemID , dr-partner-systemID</code>

在此示例中，对于四节点 MetroCluster 配置，将检索以下旧系统 ID：

- node_A_1： 4068741258
- node_A_2： 4068741260

旧控制器模块拥有的磁盘仍归这些系统 ID 所有。

```
metrocluster node show -fields node-systemid,ha-partner-systemid,dr-
partner-systemid,dr-auxiliary-systemid

dr-group-id cluster      node      node-systemid ha-partner-systemid
dr-partner-systemid dr-auxiliary-systemid
-----
1          Cluster_A    Node_A_1  4068741258    4068741260
4068741254          4068741256
1          Cluster_A    Node_A_2  4068741260    4068741258
4068741256          4068741254
1          Cluster_B    Node_B_1  -             -
-
1          Cluster_B    Node_B_2  -             -
-
4 entries were displayed.
```

在此示例中，对于双节点 MetroCluster 配置，将检索以下旧系统 ID：

- node_A_1：4068741258

旧控制器模块拥有的磁盘仍归此系统 ID 所有。

```
metrocluster node show -fields node-systemid,dr-partner-systemid

dr-group-id cluster      node      node-systemid dr-partner-systemid
-----
1          Cluster_A    Node_A_1  4068741258    4068741254
1          Cluster_B    Node_B_1  -             -
2 entries were displayed.
```

2. 对于使用 ONTAP 调解器的 MetroCluster IP 配置，获取 ONTAP 调解器的 IP 地址：

```
storage iscsi-initiator show -node * -label mediator
```

3. 如果系统型号为 AFF A220，AFF A400，FAS2750，FAS8300 或 FAS8700，确定 VLAN ID：

```
MetroCluster interconnect show
```

VLAN ID 包含在输出的适配器列中显示的适配器名称中。

在此示例中，VLAN ID 为 120 和 130：

```
metrocluster interconnect show
```

			Mirror	Mirror				
		Partner	Admin	Oper				
Node	Partner	Name	Type	Status	Status	Adapter	Type	Status
----	-----	-----	-----	-----	-----	-----	-----	-----
Node_A_1	Node_A_2	HA		enabled	online			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up
	Node_B_1	DR		enabled	online			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up
	Node_B_2	AUX		enabled	offline			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up
Node_A_2	Node_A_1	HA		enabled	online			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up
	Node_B_2	DR		enabled	online			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up
	Node_B_1	AUX		enabled	offline			
						e0a-120	iWARP	Up
						e0b-130	iWARP	Up

12 entries were displayed.

将替代驱动器与运行正常的站点隔离(MetroCluster IP配置)

您必须通过从运行正常的节点断开 MetroCluster iSCSI 启动程序连接来隔离任何替代驱动器。

关于此任务

只有 MetroCluster IP 配置才需要此操作步骤。

步骤

1. 在任一正常运行的节点的提示符处，更改为高级权限级别：

```
set -privilege advanced
```

当系统提示您继续进入高级模式并显示高级模式提示符（*>）时，您需要使用 y 进行响应。

2. 断开 DR 组中两个运行正常的节点上的 iSCSI 启动程序：

```
storage iscsi-initiator disconnect -node s幸存节点 -label *
```

此命令必须发出两次，每次针对每个正常运行的节点发出一次。

以下示例显示了用于断开站点 B 上启动程序的命令：

```
site_B::*> storage iscsi-initiator disconnect -node node_B_1 -label *
site_B::*> storage iscsi-initiator disconnect -node node_B_2 -label *
```

3. 返回到管理权限级别：

```
set -privilege admin
```

清除控制器模块上的配置

在 MetroCluster 配置中使用新控制器模块之前，必须清除现有配置。

步骤

1. 如有必要、暂停节点以显示 `LOADER` 提示符：

```
halt
```

2. 在提示符处 LOADER、将环境变量设置为默认值：

```
set-defaults
```

3. 保存环境：

```
saveenv
```

4. 在提示符处 LOADER、启动启动菜单：

```
boot_ontap 菜单
```

5. 在启动菜单提示符处，清除配置：

```
wipeconfig
```

对确认提示回答 `yes`。

节点将重新启动，并再次显示启动菜单。

6. 在启动菜单中，选择选项 * 5* 将系统启动至维护模式。

对确认提示回答 `yes`。

通过网络启动新控制器模块

如果新控制器模块的 ONTAP 版本与正常运行的控制器模块上的版本不同，则必须通过网络启动新控制器模块。

开始之前

- 您必须有权访问 HTTP 服务器。
- 您必须有权访问 NetApp 支持站点，才能下载适用于您的平台及其所运行的 ONTAP 软件版本的必要系统文件。

"NetApp 支持"

步骤

1. 访问 "NetApp 支持站点" 下载用于执行系统网络启动的文件。
2. 从 NetApp 支持站点的软件下载部分下载相应的 ONTAP 软件，并将 ontap-version_image.tgz 文件存储在可通过 Web 访问的目录中。
3. 转到可通过 Web 访问的目录，并验证所需文件是否可用。

平台型号	那么 ...
FAS/AFF8000 系列系统	将 ontap-version_image.tgzfile 的内容提取到目标目录：tar -zxvf ontap-version_image.tgz 注：如果要在 Windows 上提取内容，请使用 7-Zip 或 WinRAR 提取网络启动映像。您的目录列表应包含一个包含内核文件 netboot/kernel 的 netboot 文件夹
所有其他系统	您的目录列表应包含一个包含内核文件的 netboot 文件夹：ontap-version_image.tgz 您无需提取 ontap-version_image.tgz 文件。

4. 在 LOADER 提示符处，为管理 LIF 配置网络启动连接：

- 如果 IP 地址为 DHCP，请配置自动连接：

```
ifconfig e0M -auto
```

- 如果 IP 地址是静态的，请配置手动连接：

```
ifconfig e0M -addr=ip_addr -mask=netmask `gw=gateway`
```

5. 执行网络启动。

- 如果平台是 80xx 系列系统，请使用以下命令：

```
netboot http://web_server_ip/path_to_web-accessible_directory/netboot/kernel
```

- 如果平台是任何其他系统，请使用以下命令：

```
netboot http://web_server_ip/path_to_web-accessible_directory/ontap-version_image.tgz
```

6. 从启动菜单中，选择选项 * (7) Install new software first*，将新软件映像下载并安装到启动设备。

Disregard the following message: "This procedure is not supported for Non-Disruptive Upgrade on an HA pair". It applies to nondisruptive upgrades of software, not to upgrades of controllers.

．如果系统提示您继续运行操作步骤，请输入 `y`
，然后在系统提示您输入软件包时，输入映像文件的 URL：`
\http://web_server_ip/path_to_web-accessible_directory/ontap-version_image.tgz`

```
Enter username/password if applicable, or press Enter to continue.
```

7. 进入 `n` 当您看到类似以下的提示时，请跳过备份恢复：

```
Do you want to restore the backup configuration now? {y|n} n
```

8. 当您看到类似以下内容的提示时，输入 `y` 以重新启动：

```
The node must be rebooted to start using the newly installed software.  
Do you want to reboot now? {y|n} y
```



您必须重启节点才能使用新安装的软件。

9. 从启动菜单中，选择 * 选项 5* 以进入维护模式。

10. 如果您使用的是四节点 MetroCluster 配置，请对另一个新控制器模块重复此操作步骤。

确定更换用的控制器模块的系统ID

更换灾难站点上的所有硬件后，您必须确定新安装的存储控制器模块的系统 ID 。

关于此任务

您必须在更换用的控制器模块处于维护模式时执行此操作步骤。

本节提供了双节点和四节点配置的示例。对于双节点配置，您可以忽略对每个站点上第二个节点的引用。对于八节点配置，您必须考虑第二个 DR 组上的其他节点。这些示例假设以下条件：

- 站点 A 是灾难站点。
- node_A_1 已更换。
- node_A_2 已更换。

仅存在于四节点 MetroCluster 配置中。

- 站点 B 是正常运行的站点。
- node_B_1 运行状况良好。
- node_B_2 运行状况良好。

仅存在于四节点 MetroCluster 配置中。

此操作步骤中的示例使用具有以下系统 ID 的控制器：

MetroCluster 配置中的节点数	节点	原始系统 ID	新系统 ID	将作为 DR 配对节点与此节点配对
----------------------	----	---------	--------	-------------------

四个	node_A_1	4068741258	1574774970	node_B_1
node_A_2	4068741260	1574774991	node_B_2	node_B_1
4068741254	未更改	node_A_1	node_B_2	4068741256
未更改	node_A_2	两个	node_A_1	4068741258
1574774970	node_B_1	node_B_1	4068741254	未更改



在四节点 MetroCluster 配置中，系统通过将 site_A 中系统 ID 最低的节点与 site_B 中系统 ID 最低的节点配对来确定 DR 配对关系。由于系统 ID 会发生变化，因此在完成控制器更换后，DR 对可能会与发生灾难之前不同。

在上述示例中：

- node_A_1 （ 1574774970 ） 将与 node_B_1 （ 4068741254 ） 配对
- node_A_2 （ 1574774991 ） 将与 node_B_2 （ 4068741256 ） 配对

步骤

1. 在节点处于维护模式的情况下，显示每个节点的节点本地系统 ID：`disk show`

在以下示例中，新的本地系统 ID 为 1574774970：

```
*> disk show
Local System ID: 1574774970
...
```

2. 在第二个节点上，重复上一步。



双节点 MetroCluster 配置不需要执行此步骤。

在以下示例中，新的本地系统 ID 为 1574774991：

```
*> disk show
Local System ID: 1574774991
...
```

验证组件的 ha-config 状态

在 MetroCluster 配置中，必须将控制器模块和机箱组件的 ha-config 状态设置为 "mcc" 或 "mcc-2n"，以便它们可以正常启动。

开始之前
系统必须处于维护模式。

关于此任务
必须对每个新控制器模块执行此任务。

步骤

1. 在维护模式下，显示控制器模块和机箱的 HA 状态：

```
ha-config show
```

正确的 HA 状态取决于您的 MetroCluster 配置。

MetroCluster 配置中的控制器数量	所有组件的 HA 状态应为 ...
八节点或四节点 MetroCluster FC 配置	MCC
双节点 MetroCluster FC 配置	MCC-2n
MetroCluster IP 配置	mccip

2. 如果显示的控制器系统状态不正确，请设置控制器模块的 HA 状态：

MetroCluster 配置中的控制器数量	命令
八节点或四节点 MetroCluster FC 配置	ha-config modify controller mcc
双节点 MetroCluster FC 配置	ha-config modify controller mcc-2n
MetroCluster IP 配置	ha-config modify controller mccip

3. 如果显示的机箱系统状态不正确，请设置机箱的 HA 状态：

MetroCluster 配置中的控制器数量	命令
八节点或四节点 MetroCluster FC 配置	ha-config modify chassis mcc
双节点 MetroCluster FC 配置	ha-config modify chassis mcc-2n
MetroCluster IP 配置	ha-config modify chassis mccip

4. 在另一个替代节点上重复上述步骤。

确定原始系统上是否启用了端到端加密
您应验证原始系统是否配置了端到端加密。

步骤

1. 从运行正常的站点运行以下命令：

```
metrocluster node show -fields is-encryption-enabled
```

如果启用了加密、则会显示以下输出：

```
1 cluster_A node_A_1 true
1 cluster_A node_A_2 true
1 cluster_B node_B_1 true
1 cluster_B node_B_2 true
4 entries were displayed.
```



请参见 ["配置端到端加密"](#) 适用于受支持的系统。

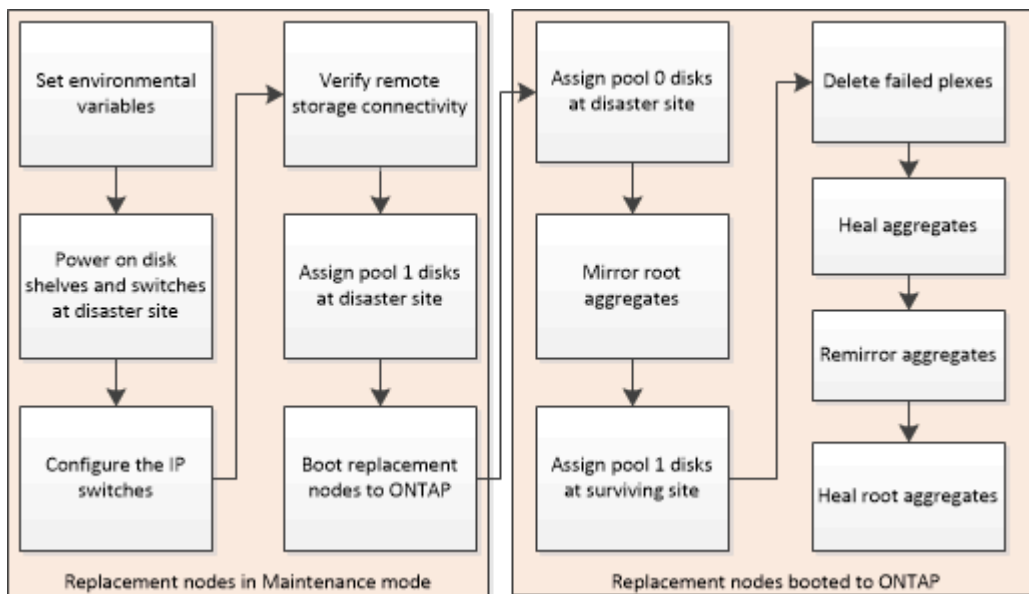
准备在 MetroCluster IP 配置中切回

准备在 **MetroCluster IP** 配置中切回

要为切回操作准备 MetroCluster IP 配置，您必须执行某些任务。

关于此任务

nbsp ;



在 **MetroCluster IP** 配置中设置所需的环境变量

在 MetroCluster IP 配置中，您必须检索以太网端口上 MetroCluster 接口的 IP 地址，然后使用它们配置替代控制器模块上的接口。

关于此任务

- 只有在 MetroCluster IP 配置中才需要执行此任务。
- 此任务中的命令可从运行正常的站点的集群提示符和灾难站点节点的 LOADER 提示符执行。
- 某些平台使用 VLAN 作为 MetroCluster IP 接口。默认情况下，这两个端口中的每个端口都使用不同的 VLAN：10 和 20。

如果支持、您还可以使用参数指定一个高于100 (介于101和4095之间)的其他(非默认) VLAN `vlan-id`。

以下平台*不*支持 `vlan-id` 参数：

- FAS8200 和 AFF A300
- AFF A320
- FAS9000和AFF A700
- AFF C800、ASA C800、AFF A800和ASA A800

所有其他平台均支持 `vlan-id` 参数。

- 这些示例中的节点的 MetroCluster IP 连接具有以下 IP 地址：



这些示例适用于 AFF A700 或 FAS9000 系统。接口因平台型号而异。

节点	端口	IP 地址
node_A_1	e5a	172.17.26.10
e5b	172.17.27.10	node_A_2
e5a	172.17.26.11	e5b
172.17.27.11	node_B_1	e5a
172.17.26.13	e5b	172.17.27.13
node_B_2	e5a	172.17.26.12

下表总结了节点与每个节点的 MetroCluster IP 地址之间的关系。

节点	HA 配对节点	DR 配对节点	DR 辅助配对节点
node_A_1	node_A_2	node_B_1	node_B_2
• e5a : 172.17.26.10 • e5b : 172.17.27.10	• e5a : 172.17.26.11 • e5b : 172.17.27.11	• e5a : 172.17.26.13 • e5b : 172.17.27.13	• e5a : 172.17.26.12 • e5b : 172.17.27.12

node_A_2	node_A_1	node_B_2	node_B_1
• e5a : 172.17.26.11 • e5b : 172.17.27.11	• e5a : 172.17.26.10 • e5b : 172.17.27.10	• e5a : 172.17.26.12 • e5b : 172.17.27.12	• e5a : 172.17.26.13 • e5b : 172.17.27.13
node_B_1	node_B_2	node_A_1	node_A_2
• e5a : 172.17.26.13 • e5b : 172.17.27.13	• e5a : 172.17.26.12 • e5b : 172.17.27.12	• e5a : 172.17.26.10 • e5b : 172.17.27.10	• e5a : 172.17.26.11 • e5b : 172.17.27.11
node_B_2	node_B_1	node_A_2	node_A_1
• e5a : 172.17.26.12 • e5b : 172.17.27.12	• e5a : 172.17.26.13 • e5b : 172.17.27.13	• e5a : 172.17.26.11 • e5b : 172.17.27.11	• e5a : 172.17.26.10 • e5b : 172.17.27.10

- 您设置的MetroCluster启动程序值取决于新系统使用的是共享集群/HA端口还是共享MetroCluster或HA端口。使用以下信息确定系统的端口。

共享集群/HA端口

下表中列出的系统使用共享集群/HA端口：

AFF和ASA系统	FAS 系统
• AFF A20 • AFF A30 • AFF C30 • AFF A50 • AFF C60 • AFF C80 • AFF A70 • AFF A90 • AFF A1K	• FAS50 • FAS70 • FAS90

共享的MetroCluster或HA端口

下表中列出的系统使用共享的MetroCluster HA/HA端口：

AFF和ASA系统	FAS 系统
• AFF A150、ASA A150 • AFF A220 • AFF C250、ASA C250 • AFF A250、ASA A250 • AFF A300 • AFF A320 • AFF C400、ASA C400 • AFF A400、ASA A400 • AFF A700 • AFF C800、ASA C800 • AFF A800、ASA A800 • AFF A900、ASA A900	• FAS2750 • FAS500f • FAS8200 • FAS8300 • FAS8700 • FAS9000 • FAS9500

步骤

1. 从正常运行的站点收集灾难站点上 MetroCluster 接口的 IP 地址：

```
MetroCluster configuration-settings connection show
```

所需地址为 * 目标网络地址 * 列中显示的 DR 配对节点地址。

根据您的平台型号使用的是共享集群/HA端口还是共享MetroCluster端口、命令输出会有所不同。

使用共享集群/HA端口的系统

```
cluster_B::*> metrocluster configuration-settings connection show
DR                               Source           Destination
DR                               Source           Destination
Group Cluster Node      Network Address Network Address Partner Type
Config State
-----
1      cluster_B
      node_B_1
      Home Port: e5a
      172.17.26.13      172.17.26.10      DR Partner
completed
      Home Port: e5a
      172.17.26.13      172.17.26.11      DR Auxiliary
completed
      Home Port: e5b
      172.17.27.13      172.17.27.10      DR Partner
completed
      Home Port: e5b
      172.17.27.13      172.17.27.11      DR Auxiliary
completed
      node_B_2
      Home Port: e5a
      172.17.26.12      172.17.26.11      DR Partner
completed
      Home Port: e5a
      172.17.26.12      172.17.26.10      DR Auxiliary
completed
      Home Port: e5b
      172.17.27.12      172.17.27.11      DR Partner
completed
      Home Port: e5b
      172.17.27.12      172.17.27.10      DR Auxiliary
completed
12 entries were displayed.
```

使用共享的MetroCluster或HA端口的系统

以下输出显示了 AFF A700 和 FAS9000 系统配置的 IP 地址，这些系统的 MetroCluster IP 接口位于端口 e5a 和 e5b 上。接口可能因平台类型而异。

```
cluster_B::*> metrocluster configuration-settings connection show
DR                               Source           Destination
```

DR	Source	Destination
Group Cluster Node	Network Address	Network Address Partner Type
Config State		
-----	-----	-----
1	cluster_B	
	node_B_1	
	Home Port: e5a	
	172.17.26.13	172.17.26.12 HA Partner
completed		
	Home Port: e5a	
	172.17.26.13	172.17.26.10 DR Partner
completed		
	Home Port: e5a	
	172.17.26.13	172.17.26.11 DR Auxiliary
completed		
	Home Port: e5b	
	172.17.27.13	172.17.27.12 HA Partner
completed		
	Home Port: e5b	
	172.17.27.13	172.17.27.10 DR Partner
completed		
	Home Port: e5b	
	172.17.27.13	172.17.27.11 DR Auxiliary
completed		
	node_B_2	
	Home Port: e5a	
	172.17.26.12	172.17.26.13 HA Partner
completed		
	Home Port: e5a	
	172.17.26.12	172.17.26.11 DR Partner
completed		
	Home Port: e5a	
	172.17.26.12	172.17.26.10 DR Auxiliary
completed		
	Home Port: e5b	
	172.17.27.12	172.17.27.13 HA Partner
completed		
	Home Port: e5b	
	172.17.27.12	172.17.27.11 DR Partner
completed		
	Home Port: e5b	
	172.17.27.12	172.17.27.10 DR Auxiliary
completed		
12 entries were displayed.		

2. 如果需要确定接口的 VLAN ID 或网关地址，请从正常运行的站点确定 VLAN ID：

MetroCluster configuration-settings interface show

- 如果平台型号支持VLAN ID (请参见)，并且不使用默认VLAN ID，则需要确定VLAN [列表ID](#)。
- 如果使用，则需要网关地址 "[第 3 层广域网](#)"。

VLAN ID 包含在输出的 * 网络地址 * 列中。* 网关 * 列显示网关 IP 地址。

在此示例中，接口为 VLAN ID 为 120 的 e0a 和 VLAN ID 为 130 的 e0b：

```
Cluster-A::*> metrocluster configuration-settings interface show
DR
Config
Group Cluster Node      Network Address Netmask      Gateway
State
-----
1
    cluster_A
        node_A_1
            Home Port: e0a-120
                172.17.26.10  255.255.255.0  -
completed
            Home Port: e0b-130
                172.17.27.10  255.255.255.0  -
completed
```

3. 在每个灾难站点节点的提示符处 LOADER、根据您的平台型号是使用共享集群/HA端口还是共享MetroCluster/HA端口设置启动程序值：



- 如果接口使用的是默认VLAN，或者平台型号不使用VLAN ID (请参见 [列表](#))，则不需要_vla-id_。
- 如果配置未使用 "[第 3 层广域网](#)"， gateway-ip-address 的值为 * 0 *（零）。

使用共享集群/HA端口的系统

设置以下布塔格：

```
setenv bootarg.mcc.port_a_ip_config local-IP-address/local-IP-  
mask,0,0,DR-partner-IP-address,DR-aux-partnerIP-address,vlan-id
```

```
setenv bootarg.mcc.port_b_ip_config local-IP-address/local-IP-  
mask,0,0,DR-partner-IP-address,DR-aux-partnerIP-address,vlan-id
```

以下命令使用 VLAN 120 为第一个网络设置 node_A_1 的值，并使用 VLAN 130 为第二个网络设置 VLAN 130：

```
setenv bootarg.mcc.port_a_ip_config  
172.17.26.10/23,0,0,172.17.26.13,172.17.26.12,120
```

```
setenv bootarg.mcc.port_b_ip_config  
172.17.27.10/23,0,0,172.17.27.13,172.17.27.12,130
```

以下示例显示了不带 VLAN ID 的 node_A_1 的命令：

```
setenv bootarg.mcc.port_a_ip_config  
172.17.26.10/23,0,0,172.17.26.13,172.17.26.12
```

```
setenv bootarg.mcc.port_b_ip_config  
172.17.27.10/23,0,0,172.17.27.13,172.17.27.12
```

使用共享的MetroCluster或HA端口的系统

设置以下布塔格：

```
setenv bootarg.mcc.port_a_ip_config local-IP-address/local-IP-  
mask,0,HA-partner-IP-address,DR-partner-IP-address,DR-aux-partnerIP-  
address,vlan-id
```

```
setenv bootarg.mcc.port_b_ip_config local-IP-address/local-IP-  
mask,0,HA-partner-IP-address,DR-partner-IP-address,DR-aux-partnerIP-  
address,vlan-id
```

以下命令使用 VLAN 120 为第一个网络设置 node_A_1 的值，并使用 VLAN 130 为第二个网络设置 VLAN 130：

```
setenv bootarg.mcc.port_a_ip_config  
172.17.26.10/23,0,172.17.26.11,172.17.26.13,172.17.26.12,120  
  
setenv bootarg.mcc.port_b_ip_config  
172.17.27.10/23,0,172.17.27.11,172.17.27.13,172.17.27.12,130
```

以下示例显示了不带 VLAN ID 的 node_A_1 的命令：

```
setenv bootarg.mcc.port_a_ip_config  
172.17.26.10/23,0,172.17.26.11,172.17.26.13,172.17.26.12  
  
setenv bootarg.mcc.port_b_ip_config  
172.17.27.10/23,0,172.17.27.11,172.17.27.13,172.17.27.12
```

4. 从正常运行的站点收集灾难站点的 UUID：

```
MetroCluster node show -fields node-cluster-uuid , node-uuid
```

```

cluster_B::> metrocluster node show -fields node-cluster-uuid, node-uuid

(metrocluster node show)
dr-group-id cluster      node      node-uuid
node-cluster-uuid
-----
1            cluster_A   node_A_1 f03cb63c-9a7e-11e7-b68b-00a098908039
ee7db9d5-9a82-11e7-b68b-00a098
908039
1            cluster_A   node_A_2 aa9a7a7a-9a81-11e7-a4e9-00a098908c35
ee7db9d5-9a82-11e7-b68b-00a098
908039
1            cluster_B   node_B_1 f37b240b-9ac1-11e7-9b42-00a098c9e55d
07958819-9ac6-11e7-9b42-00a098
c9e55d
1            cluster_B   node_B_2 bf8e3f8f-9ac4-11e7-bd4e-00a098ca379f
07958819-9ac6-11e7-9b42-00a098
c9e55d
4 entries were displayed.
cluster_A::~*>

```

节点	UUID
集群 B	07958819-9ac6-11e7-9b42-00a098c9e55d
node_B_1	f37b240b-9ac1-11e7-9b42-00a098c9e55d
node_B_2	bf8e3f8f-9ac4-11e7-bd4e-00a098ca379f
cluster_A	ee7db9d5-9a82-11e7-b68b-00a098908039
node_A_1	f03cb63c-9a7e-11e7-b68b-00a098908039
node_A_2	aa9a7a7a-9a81-11e7-a4e9-00a098908c35

5. 在替代节点的 LOADER 提示符处，设置 UUID：

```
setenv bootarg.mgwd.partner_cluster_uuid partner-cluster-UUID

setenv bootarg.mgwd.cluster_uuid local-cluster-UUID

setenv bootarg.mcc.pri_partner_uuid DR-partner-node-UUID

setenv bootarg.mcc.aux_partner_uuid DR-aux-partner-node-UUID

setenv bootarg.mcc_iscsi.node_uuid local-node-UUID`
```

a. 设置 node_A_1 上的 UUID 。

以下示例显示了用于设置 node_A_1 上的 UUID 的命令：

```
setenv bootarg.mgwd.cluster_uuid ee7db9d5-9a82-11e7-b68b-00a098908039

setenv bootarg.mgwd.partner_cluster_uuid 07958819-9ac6-11e7-9b42-00a098c9e55d

setenv bootarg.mcc.pri_partner_uuid f37b240b-9ac1-11e7-9b42-00a098c9e55d

setenv bootarg.mcc.aux_partner_uuid bf8e3f8f-9ac4-11e7-bd4e-00a098ca379f

setenv bootarg.mcc_iscsi.node_uuid f03cb63c-9a7e-11e7-b68b-00a098908039
```

b. 设置 node_A_2 上的 UUID ：

以下示例显示了用于设置 node_A_2 上的 UUID 的命令：

```
setenv bootarg.mgwd.cluster_uuid ee7db9d5-9a82-11e7-b68b-00a098908039

setenv bootarg.mgwd.partner_cluster_uuid 07958819-9ac6-11e7-9b42-00a098c9e55d

setenv bootarg.mcc.pri_partner_uuid bf8e3f8f-9ac4-11e7-bd4e-00a098ca379f

setenv bootarg.mcc.aux_partner_uuid f37b240b-9ac1-11e7-9b42-00a098c9e55d

setenv bootarg.mcc_iscsi.node_uuid aa9a7a7a-9a81-11e7-a4e9-00a098908c35
```

6. 如果原始系统配置了 ADP ，请在每个替代节点的 LOADER 提示符处启用 ADP ：

```
setenv bootarg.mcc.ADP 启用 true
```

7. 如果运行的是 ONTAP 9.5 , 9.6 或 9.7 , 请在每个替代节点的 LOADER 提示符处启用以下变量:

```
setenv bootarg.mcc.lun_part true
```

- a. 设置 node_A_1 上的变量。

以下示例显示了在运行 ONTAP 9.6 时用于设置 node_A_1 上的值的命令:

```
setenv bootarg.mcc.lun_part true
```

- b. 设置 node_A_2 上的变量。

以下示例显示了在运行 ONTAP 9.6 时用于设置 node_A_2 上的值的命令:

```
setenv bootarg.mcc.lun_part true
```

8. 如果原始系统已配置端到端加密、请在每个替代节点的加载程序提示符处设置以下启动程序:

```
setenv bootarg.mccip.encryption_enabled 1
```

9. 如果原始系统配置了 ADP , 请在每个替代节点的 LOADER 提示符处设置原始系统 ID (* 不 * 替代控制器模块的系统 ID) 和节点的 DR 配对节点的系统 ID :

```
setenv bootarg.mcc.local_config_id original-sysID
```

```
setenv bootarg.mcc.dr_partner dr_partner-sysID
```

"确定旧控制器模块的系统ID"

- a. 设置 node_A_1 上的变量。

以下示例显示了用于设置 node_A_1 上的系统 ID 的命令:

- node_A_1 的旧系统 ID 为 4068741258 。
- node_B_1 的系统 ID 为 4068741254 。

```
setenv bootarg.mcc.local_config_id 4068741258  
setenv bootarg.mcc.dr_partner 4068741254
```

- b. 设置 node_A_2 上的变量。

以下示例显示了用于设置 node_A_2 上的系统 ID 的命令:

- node_A_1 的旧系统 ID 为 4068741260 。

- node_B_1 的系统 ID 为 4068741256。

```
setenv bootarg.mcc.local_config_id 4068741260
setenv bootarg.mcc.dr_partner 4068741256
```

启动灾难站点上的设备（**MetroCluster IP 配置**）

您必须打开灾难站点上的磁盘架和 MetroCluster IP 交换机组件的电源。灾难站点上的控制器模块将保留在 LOADER 提示符处。

关于此任务

此操作步骤中的示例假设如下：

- 站点 A 是灾难站点。
- 站点 B 是正常运行的站点。

步骤

1. 打开灾难站点上的磁盘架，并确保所有磁盘都在运行。
2. 如果尚未打开 MetroCluster IP 交换机，请将其打开。

配置 IP 交换机（**MetroCluster IP 配置**）

您必须配置已更换的任何 IP 交换机。

关于此任务

此任务仅限适用场景 MetroCluster IP 配置。

必须在两台交换机上执行此操作。在配置第一台交换机后，验证运行正常的站点上的存储访问是否不受影响。



如果正常运行的站点上的存储访问受到影响，则不能继续使用第二个交换机。

步骤

1. 请参见 ["MetroCluster IP 安装和配置：： ONTAP MetroCluster 配置之间的差异"](#) 有关为替代交换机布线和配置交换机的过程。

您可以使用以下各节中的过程：

- 为 IP 交换机布线
 - 配置 IP 交换机
2. 如果在正常运行的站点上禁用了 ISL，请启用这些 ISL 并验证这些 ISL 是否联机。
 - a. 在第一台交换机上启用 ISL 接口：

无关闭

以下示例显示了适用于 Broadcom IP 交换机或 Cisco IP 交换机的命令。

交换机供应商	命令
Broadcom	<pre>(IP_Switch_A_1)> enable (IP_switch_A_1)# configure (IP_switch_A_1) (Config) # interface 0/13-0/16 (IP_switch_A_1) (Interface 0/13-0/16)# no shutdown (IP_switch_A_1) (Interface 0/13-0/16)# exit (IP_switch_A_1) (Config) # exit</pre>
Cisco	<pre>IP_switch_A_1# conf t IP_switch_A_1(config)# int eth1/15-eth1/20 IP_switch_A_1(config)# no shutdown IP_switch_A_1(config)# copy running startup IP_switch_A_1(config)# show interface brief</pre>

- b. 在配对交换机上启用 ISL 接口：

无关闭

以下示例显示了适用于 Broadcom IP 交换机或 Cisco IP 交换机的命令。

交换机供应商	命令
Broadcom	<pre>(IP_Switch_A_2)> enable (IP_switch_A_2)# configure (IP_switch_A_2) (Config) # interface 0/13-0/16 (IP_switch_A_2) (Interface 0/13-0/16)# no shutdown (IP_switch_A_2) (Interface 0/13-0/16)# exit (IP_switch_A_2) (Config) # exit</pre>

Cisco

```
IP_switch_A_2# conf t
IP_switch_A_2(config)# int
eth1/15-eth1/20
IP_switch_A_2(config)# no
shutdown
IP_switch_A_2(config)# copy
running startup
IP_switch_A_2(config)# show
interface brief
```

c. 验证接口是否已启用：

s如何使用接口简介

以下示例显示了 Cisco 交换机的输出。

```
IP_switch_A_2(config)# show interface brief
```

```
-----
Port VRF Status IP Address Speed MTU
-----
```

```
mt0 -- up 10.10.99.10 100 1500
-----
```

```
Ethernet      VLAN Type Mode      Status Reason Speed   Port
Interface                                           Ch
#
```

```
-----
```

```
.
.
.
```

Eth1/15	10	eth	access	up	none	40G(D)	--
Eth1/16	10	eth	access	up	none	40G(D)	--
Eth1/17	10	eth	access	down	none	auto(D)	--
Eth1/18	10	eth	access	down	none	auto(D)	--
Eth1/19	10	eth	access	down	none	auto(D)	--
Eth1/20	10	eth	access	down	none	auto(D)	--

```
.
.
.
```

```
IP_switch_A_2#
```

验证与远程站点的存储连接（**MetroCluster IP** 配置）

您必须确认已更换的节点已连接到运行正常的站点上的磁盘架。

关于此任务

此任务将在灾难站点的替代节点上执行。

此任务在维护模式下执行。

步骤

1. 显示原始系统 ID 所拥有的磁盘。

```
disk show -s old-system-ID
```

远程磁盘可由 0m 设备识别。0m 表示磁盘已通过 MetroCluster iSCSI 连接进行连接。这些磁盘稍后必须在恢复操作步骤中重新分配。

```
*> disk show -s 4068741256
Local System ID: 1574774970

  DISK      OWNER          POOL  SERIAL NUMBER  HOME
DR HOME
-----
0m.i0.0L11 node_A_2 (4068741256) Pool1 S396NA0HA02128 node_A_2
(4068741256) node_A_2 (4068741256)
0m.i0.1L38 node_A_2 (4068741256) Pool1 S396NA0J148778 node_A_2
(4068741256) node_A_2 (4068741256)
0m.i0.0L52 node_A_2 (4068741256) Pool1 S396NA0J148777 node_A_2
(4068741256) node_A_2 (4068741256)
...
...
NOTE: Currently 49 disks are unowned. Use 'disk show -n' for additional
information.
*>
```

2. 对其他替代节点重复此步骤

为灾难站点上的池 1 磁盘重新分配磁盘所有权（**MetroCluster IP** 配置）

如果在灾难站点上更换了一个或两个控制器模块或 NVRAM 卡，则系统 ID 已更改，您必须将属于根聚合的磁盘重新分配给更换用的控制器模块。

关于此任务

由于节点处于切换模式，因此在此任务中仅会重新分配包含灾难站点的 pool1 根聚合的磁盘。此时，它们是唯一仍归旧系统 ID 所有的磁盘。

此任务将在灾难站点的替代节点上执行。

此任务在维护模式下执行。

这些示例假设以下条件：

- 站点 A 是灾难站点。
- node_A_1 已更换。
- node_A_2 已更换。
- 站点 B 是正常运行的站点。
- node_B_1 运行状况良好。
- node_B_2 运行状况良好。

旧的和新的系统ID在中进行了标识 ["更换硬件并启动新控制器"](#)。

此操作步骤中的示例使用具有以下系统 ID 的控制器：

节点	原始系统 ID	新系统 ID
node_A_1	4068741258	1574774970
node_A_2	4068741260	1574774991
node_B_1	4068741254	未更改
node_B_2	4068741256	未更改

步骤

1. 在替代节点处于维护模式的情况下，根据您的系统是否配置了 ADP 和 ONTAP 版本，使用正确的命令重新分配根聚合磁盘。

出现提示时，您可以继续重新分配。

如果系统正在使用 ADP ...	使用此命令重新分配磁盘 ...
是（ONTAP 9.8）	<code>dreassign -s old-system-ID -d new-system-ID -r dr-partner-system-ID</code>
是（ONTAP 9.7.x 及更早版本）	<code>dreassign -s old-system-ID -d new-system-ID -p old-partner-system-ID</code>
否	<code>dreassign -s old-system-ID -d new-system-ID</code>

以下示例显示了在非 ADP 系统上重新分配驱动器的情况：

```
*> disk reassign -s 4068741256 -d 1574774970
Partner node must not be in Takeover mode during disk reassignment from
maintenance mode.
Serious problems could result!!
Do not proceed with reassignment if the partner is in takeover mode.
Abort reassignment (y/n)? n

After the node becomes operational, you must perform a takeover and
giveback of the HA partner node to ensure disk reassignment is
successful.
Do you want to continue (y/n)? y
Disk ownership will be updated on all disks previously belonging to
Filer with sysid 537037643.
Do you want to continue (y/n)? y
disk reassign parameters: new_home_owner_id 537070473 ,
new_home_owner_name
Disk 0m.i0.3L14 will be reassigned.
Disk 0m.i0.1L6 will be reassigned.
Disk 0m.i0.1L8 will be reassigned.
Number of disks to be reassigned: 3
```

2. 销毁邮箱磁盘的内容：

m 邮箱销毁本地

出现提示时，您可以继续执行销毁操作。

以下示例显示了 mailbox destroy local 命令的输出：

```
*> mailbox destroy local
Destroying mailboxes forces a node to create new empty mailboxes,
which clears any takeover state, removes all knowledge
of out-of-date plexes of mirrored volumes, and will prevent
management services from going online in 2-node cluster
HA configurations.
Are you sure you want to destroy the local mailboxes? y
.....Mailboxes destroyed.
*>
```

3. 如果更换了磁盘，则会出现故障的本地丛，必须将其删除。

a. 显示聚合状态：

聚合状态

在以下示例中，从 node_A_1_aggr0/plex0 出现故障。

```
*> aggr status
Aug 18 15:00:07 [node_B_1:raid.vol.mirror.degraded:ALERT]: Aggregate
node_A_1_aggr0 is
    mirrored and one plex has failed. It is no longer protected by
    mirroring.
Aug 18 15:00:07 [node_B_1:raid.debug:info]: Mirrored aggregate
node_A_1_aggr0 has plex0
    clean(-1), online(0)
Aug 18 15:00:07 [node_B_1:raid.debug:info]: Mirrored aggregate
node_A_1_aggr0 has plex2
    clean(0), online(1)
Aug 18 15:00:07 [node_B_1:raid.mirror.vote.noRecord1Plex:error]:
WARNING: Only one plex
    in aggregate node_A_1_aggr0 is available. Aggregate might contain
    stale data.
Aug 18 15:00:07 [node_B_1:raid.debug:info]:
volobj_mark_sb_recovery_aggrs: tree:
    node_A_1_aggr0 vol_state:1 mcc_dr_opstate: unknown
Aug 18 15:00:07 [node_B_1:raid.fsm.commitStateTransit:debug]:
/node_A_1_aggr0 (VOL):
    raid state change UNINITD -> NORMAL
Aug 18 15:00:07 [node_B_1:raid.fsm.commitStateTransit:debug]:
/node_A_1_aggr0 (MIRROR):
    raid state change UNINITD -> DEGRADED
Aug 18 15:00:07 [node_B_1:raid.fsm.commitStateTransit:debug]:
/node_A_1_aggr0/plex0
    (PLEX): raid state change UNINITD -> FAILED
Aug 18 15:00:07 [node_B_1:raid.fsm.commitStateTransit:debug]:
/node_A_1_aggr0/plex2
    (PLEX): raid state change UNINITD -> NORMAL
Aug 18 15:00:07 [node_B_1:raid.fsm.commitStateTransit:debug]:
/node_A_1_aggr0/plex2/rg0
    (GROUP): raid state change UNINITD -> NORMAL
Aug 18 15:00:07 [node_B_1:raid.debug:info]: Topology updated for
aggregate node_A_1_aggr0
    to plex plex2
*>
```

b. 删除故障丛：

```
aggr destroy plex-id
```

```
*> aggr destroy node_A_1_aggr0/plex0
```

4. 暂停节点以显示 LOADER 提示符：

```
halt
```

5. 在灾难站点的另一个节点上重复上述步骤。

在 **MetroCluster IP** 配置中的替代控制器模块上启动到 **ONTAP**

您必须将灾难站点上的替代节点启动到 ONTAP 操作系统。

关于此任务

此任务从处于维护模式的灾难站点上的节点开始。

步骤

1. 在一个替代节点上，退出到 LOADER 提示符：halt
2. 显示启动菜单：boot_ontap menu
3. 从启动菜单中，选择选项 6 * 从备份配置更新闪存 *。

系统启动两次。系统提示您继续时，您应回答 yes。在第二次启动后，当系统 ID 不匹配时，您应响应 y。



如果未清除已用替代控制器模块的 NVRAM 内容，则可能会看到以下崩溃消息：panic : NVRAM 内容无效 ... 如果发生这种情况，请将系统重新启动到 ONTAP 提示符 (boot_ontap menu)。然后，您需要 [重置boot_recovery](#)和[rdb_Corrupt bootargs](#)

。确认继续提示：

```
Selection (1-9)? 6
```

```
This will replace all flash-based configuration with the last backup  
to  
disks. Are you sure you want to continue?: yes
```

。系统 ID 不匹配提示：

```
WARNING: System ID mismatch. This usually occurs when replacing a  
boot device or NVRAM cards!  
Override system ID? {y|n} y
```

4. 在正常运行的站点中，验证是否已将正确的配对系统 ID 应用于节点：

```
MetroCluster node show -fields node-systemID , ha-partner-systemID , dr-  
partner-systemID , dr-auxiliary-systemID
```

在此示例中，输出中应显示以下新的系统 ID：

- ° node_A_1： 1574774970
- ° node_A_2： 1574774991

"ha-partner-systemID" 列应显示新的系统 ID。

```
metrocluster node show -fields node-systemid,ha-partner-systemid,dr-
partner-systemid,dr-auxiliary-systemid

dr-group-id cluster      node      node-systemid ha-partner-systemid dr-
partner-systemid dr-auxiliary-systemid
-----
1          Cluster_A  Node_A_1  1574774970    1574774991
4068741254          4068741256
1          Cluster_A  Node_A_2  1574774991    1574774970
4068741256          4068741254
1          Cluster_B  Node_B_1  -             -
-
1          Cluster_B  Node_B_2  -             -
-
4 entries were displayed.
```

5. 如果未正确设置配对系统 ID，则必须手动设置正确的值：

- 暂停并显示节点上的 LOADER 提示符。
- 验证 partner-sysID bootarg 的当前值：

```
printenv
```

- 将此值设置为正确的配对系统 ID：

```
setenv partner-sysid partner-sysID
```

- 启动节点：

```
boot_ontap
```

- 如有必要，在另一节点上重复这些子步骤。

6. 确认灾难站点上的替代节点已做好切回准备：

```
MetroCluster node show
```

替代节点应处于等待切回恢复模式。如果它们处于正常模式，则可以重新启动替代节点。启动后，节点应处于等待切回恢复模式。

以下示例显示替代节点已做好切回准备：

```
cluster_B::> metrocluster node show
```

DR Group	Cluster	Node	Configuration State	DR Mirroring Mode
1	cluster_B	node_B_1	configured	enabled switchover
		node_B_2	configured	enabled switchover
	cluster_A	node_A_1	configured	enabled waiting for switchback recovery
		node_A_2	configured	enabled waiting for switchback recovery

4 entries were displayed.

```
cluster_B::>
```

7. 验证 MetroCluster 连接配置设置:

MetroCluster configuration-settings connection show

配置状态应指示已完成。

```
cluster_B::*> metrocluster configuration-settings connection show
```

DR Group	Cluster	Node	Source Network Address	Destination Network Address	Partner Type
1	cluster_B	node_B_2	Home Port: e5a 172.17.26.13	172.17.26.12	HA Partner
			Home Port: e5a 172.17.26.13	172.17.26.10	DR Partner
			Home Port: e5a 172.17.26.13	172.17.26.11	DR Auxiliary
			Home Port: e5b 172.17.27.13	172.17.27.12	HA Partner

	Home Port: e5b		
completed	172.17.27.13	172.17.27.10	DR Partner
	Home Port: e5b		
completed	172.17.27.13	172.17.27.11	DR Auxiliary
	node_B_1		
	Home Port: e5a		
completed	172.17.26.12	172.17.26.13	HA Partner
	Home Port: e5a		
completed	172.17.26.12	172.17.26.11	DR Partner
	Home Port: e5a		
completed	172.17.26.12	172.17.26.10	DR Auxiliary
	Home Port: e5b		
completed	172.17.27.12	172.17.27.13	HA Partner
	Home Port: e5b		
completed	172.17.27.12	172.17.27.11	DR Partner
	Home Port: e5b		
completed	172.17.27.12	172.17.27.10	DR Auxiliary
	cluster_A		
	node_A_2		
	Home Port: e5a		
completed	172.17.26.11	172.17.26.10	HA Partner
	Home Port: e5a		
completed	172.17.26.11	172.17.26.12	DR Partner
	Home Port: e5a		
completed	172.17.26.11	172.17.26.13	DR Auxiliary
	Home Port: e5b		
completed	172.17.27.11	172.17.27.10	HA Partner
	Home Port: e5b		
completed	172.17.27.11	172.17.27.12	DR Partner
	Home Port: e5b		
completed	172.17.27.11	172.17.27.13	DR Auxiliary
	node_A_1		

```

completed      Home Port: e5a
                 172.17.26.10      172.17.26.11      HA Partner
completed      Home Port: e5a
                 172.17.26.10      172.17.26.13      DR Partner
completed      Home Port: e5a
                 172.17.26.10      172.17.26.12      DR Auxiliary
completed      Home Port: e5b
                 172.17.27.10      172.17.27.11      HA Partner
completed      Home Port: e5b
                 172.17.27.10      172.17.27.13      DR Partner
completed      Home Port: e5b
                 172.17.27.10      172.17.27.12      DR Auxiliary
completed
24 entries were displayed.

cluster_B::*>

```

8. 在灾难站点的另一个节点上重复上述步骤。

【重置启动恢复】重置boot_recovery和rdb_Corrupt bootargs

如果需要、您可以重置boot_recovery和rdb_Corrupt_bootargs

步骤

1. 将节点暂停回LOADER提示符：

```
siteA::*> halt -node <node-name>
```

2. 检查是否已设置以下bootarg：

```
LOADER> printenv bootarg.init.boot_recovery
LOADER> printenv bootarg.rdb_corrupt
```

3. 如果已将任一bootarg设置为值、请取消设置并启动ONTAP：

```
LOADER> unsetenv bootarg.init.boot_recovery
LOADER> unsetenv bootarg.rdb_corrupt
LOADER> saveenv
LOADER> bye
```

将连接从正常运行的节点还原到灾难站点（**MetroCluster IP** 配置）

您必须从运行正常的节点还原 MetroCluster iSCSI 启动程序连接。

关于此任务

只有 MetroCluster IP 配置才需要此操作步骤。

步骤

1. 在任一正常运行的节点的提示符处，更改为高级权限级别：

```
set -privilege advanced
```

当系统提示您继续进入高级模式并显示高级模式提示符（*>）时，您需要使用 y 进行响应。

2. 在 DR 组中的两个运行正常的节点上连接 iSCSI 启动程序：

```
storage iscsi-initiator connect -node n幸存节点 -label *
```

以下示例显示了用于连接站点 B 上启动程序的命令：

```
site_B::*> storage iscsi-initiator connect -node node_B_1 -label *
site_B::*> storage iscsi-initiator connect -node node_B_2 -label *
```

3. 返回到管理权限级别：

```
set -privilege admin
```

验证自动分配或手动分配池 0 驱动器

在配置了 ADP 的系统上，您必须验证是否已自动分配池 0 驱动器。在未配置 ADP 的系统上，您必须手动分配池 0 驱动器。

验证灾难站点（**MetroCluster IP** 系统）上 **ADP** 系统上池 0 驱动器的驱动器分配

如果在灾难站点上更换了驱动器，并且系统配置了 ADP，则必须确认远程驱动器对节点可见且已正确分配。

步骤

1. 验证是否自动分配池 0 驱动器：

d展示

在以下示例中，如果 AFF A800 系统没有外部磁盘架，则会自动将四分之一（8 个驱动器）分配给 node_A_1，并将四分之一自动分配给 node_A_2。其余驱动器将是 node_B_1 和 node_B_2 的远程（pool1）驱动器。

```
cluster_A::*> disk show
```

Disk Owner	Usable Size	Disk Shelf	Bay	Container Type	Type	Container Name
node_A_1:0n.12	1.75TB	0	12	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.13	1.75TB	0	13	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.14	1.75TB	0	14	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.15	1.75TB	0	15	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.16	1.75TB	0	16	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.17	1.75TB	0	17	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.18	1.75TB	0	18	SSD-NVM	shared	aggr0
node_A_1						
node_A_1:0n.19	1.75TB	0	19	SSD-NVM	shared	-
node_A_1						
node_A_2:0n.0	1.75TB	0	0	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.1	1.75TB	0	1	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.2	1.75TB	0	2	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.3	1.75TB	0	3	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.4	1.75TB	0	4	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.5	1.75TB	0	5	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.6	1.75TB	0	6	SSD-NVM	shared	
aggr0_node_A_2_0 node_A_2						
node_A_2:0n.7	1.75TB	0	7	SSD-NVM	shared	-
node_A_2						
node_A_2:0n.24	-	0	24	SSD-NVM	unassigned	-
node_A_2:0n.25	-	0	25	SSD-NVM	unassigned	-
node_A_2:0n.26	-	0	26	SSD-NVM	unassigned	-
node_A_2:0n.27	-	0	27	SSD-NVM	unassigned	-

```
node_A_2:0n.28 - 0 28 SSD-NVM unassigned - -
node_A_2:0n.29 - 0 29 SSD-NVM unassigned - -
node_A_2:0n.30 - 0 30 SSD-NVM unassigned - -
node_A_2:0n.31 - 0 31 SSD-NVM unassigned - -
node_A_2:0n.36 - 0 36 SSD-NVM unassigned - -
node_A_2:0n.37 - 0 37 SSD-NVM unassigned - -
node_A_2:0n.38 - 0 38 SSD-NVM unassigned - -
node_A_2:0n.39 - 0 39 SSD-NVM unassigned - -
node_A_2:0n.40 - 0 40 SSD-NVM unassigned - -
node_A_2:0n.41 - 0 41 SSD-NVM unassigned - -
node_A_2:0n.42 - 0 42 SSD-NVM unassigned - -
node_A_2:0n.43 - 0 43 SSD-NVM unassigned - -
32 entries were displayed.
```

在灾难站点的非 **ADP** 系统上分配池 0 驱动器（**MetroCluster IP** 配置）

如果已在灾难站点上更换驱动器，并且系统未配置 ADP，则需要手动将新驱动器分配到池 0。

关于此任务

对于 ADP 系统，驱动器会自动分配。

步骤

1. 在灾难站点的一个替代节点上，重新分配节点的池 0 驱动器：

```
s存储磁盘 assign -n number-of-replacement disks -p 0
```

此命令将分配灾难站点上新添加的（无所有权的）驱动器。分配的驱动器数量和大小应与发生灾难之前节点所分配的驱动器数量和大小相同（或更大）。storage disk assign 手册页包含有关执行更精细的驱动器分配的详细信息。

2. 对灾难站点上的另一个替代节点重复此步骤。

在运行正常的站点上分配池 1 驱动器（**MetroCluster IP** 配置）

如果在灾难站点上更换了驱动器，并且系统未配置 ADP，则在正常运行的站点上，您需要手动将位于灾难站点上的远程驱动器分配给正常运行的节点的池 1。您必须确定要分配的驱动器数量。

关于此任务

对于 ADP 系统，驱动器会自动分配。

步骤

1. 在正常运行的站点上，分配第一个节点的池 1（远程）驱动器：storage disk assign -n number-of-replacement disks -p 1 0m*

此命令将在灾难站点上分配新添加的驱动器和未分配的驱动器。

以下命令分配 22 个驱动器：

```
cluster_B::> storage disk assign -n 22 -p 1 0m*
```

删除运行正常的站点拥有的故障丛（**MetroCluster IP** 配置）

更换硬件并分配磁盘后，您必须删除故障远程丛，这些丛属于运行正常的站点节点，但位于灾难站点上。

关于此任务

这些步骤将在运行正常的集群上执行。

步骤

1. 确定本地聚合：storage aggregate show -is-home true

```
cluster_B::> storage aggregate show -is-home true

cluster_B Aggregates:
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
node_B_1_aggr0 1.49TB   74.12GB 95% online    1 node_B_1
raid4,

mirror

degraded
node_B_2_aggr0 1.49TB   74.12GB 95% online    1 node_B_2
raid4,

mirror

degraded
node_B_1_aggr1 2.99TB   2.88TB  3% online   15 node_B_1
raid_dp,

mirror

degraded
node_B_1_aggr2 2.99TB   2.91TB  3% online   14 node_B_1
raid_tec,

mirror
```

```
degraded
node_B_2_aggr1 2.95TB  2.80TB    5% online      37 node_B_2
raid_dp,

mirror

degraded
node_B_2_aggr2 2.99TB  2.87TB    4% online      35 node_B_2
raid_tec,

mirror

degraded
6 entries were displayed.

cluster_B::>
```

2. 确定发生故障的远程丛：

```
storage aggregate plex show
```

以下示例将调用处于远程状态（而不是 plex0）且状态为 "Failed" 的丛：

```
cluster_B::> storage aggregate plex show -fields aggregate,status,is-
online,Plex,pool
aggregate      plex  status          is-online pool
-----
node_B_1_aggr0 plex0 normal,active true      0
node_B_1_aggr0 plex4 failed,inactive false - <<<<---Plex at remote site
node_B_2_aggr0 plex0 normal,active true      0
node_B_2_aggr0 plex4 failed,inactive false - <<<<---Plex at remote site
node_B_1_aggr1 plex0 normal,active true      0
node_B_1_aggr1 plex4 failed,inactive false - <<<<---Plex at remote site
node_B_1_aggr2 plex0 normal,active true      0
node_B_1_aggr2 plex1 failed,inactive false - <<<<---Plex at remote site
node_B_2_aggr1 plex0 normal,active true      0
node_B_2_aggr1 plex4 failed,inactive false - <<<<---Plex at remote site
node_B_2_aggr2 plex0 normal,active true      0
node_B_2_aggr2 plex1 failed,inactive false - <<<<---Plex at remote site
node_A_1_aggr1 plex0 failed,inactive false -
node_A_1_aggr1 plex4 normal,active true      1
node_A_1_aggr2 plex0 failed,inactive false -
node_A_1_aggr2 plex1 normal,active true      1
node_A_2_aggr1 plex0 failed,inactive false -
node_A_2_aggr1 plex4 normal,active true      1
node_A_2_aggr2 plex0 failed,inactive false -
node_A_2_aggr2 plex1 normal,active true      1
20 entries were displayed.

cluster_B::>
```

3. 使每个故障丛脱机，然后将其删除：

a. 使故障丛脱机：

```
storage aggregate plex offline -aggregate aggregate-name -plex plex-id
```

以下示例显示了正在脱机的聚合 "node_B_2_aggr1/plex1"：

```
cluster_B::> storage aggregate plex offline -aggregate node_B_1_aggr0
-plex plex4

Plex offline successful on plex: node_B_1_aggr0/plex4
```

b. 删除故障丛：

```
storage aggregate plex delete -aggregate aggregate-name -plex plex-id
```

出现提示时，您可以销毁丛。

以下示例显示了要删除的丛 node_B_2_aggr1/plex1。

```
cluster_B::> storage aggregate plex delete -aggregate node_B_1_aggr0
-plex plex4

Warning: Aggregate "node_B_1_aggr0" is being used for the local
management root
        volume or HA partner management root volume, or has been
marked as
        the aggregate to be used for the management root volume
after a
        reboot operation. Deleting plex "plex4" for this aggregate
could lead
        to unavailability of the root volume after a disaster
recovery
        procedure. Use the "storage aggregate show -fields
        has-mroot,has-partner-mroot,root" command to view such
aggregates.

Warning: Deleting plex "plex4" of mirrored aggregate "node_B_1_aggr0"
on node
        "node_B_1" in a MetroCluster configuration will disable its
synchronous disaster recovery protection. Are you sure you
want to
        destroy this plex? {y|n}: y
[Job 633] Job succeeded: DONE

cluster_B::>
```

您必须对每个故障丛重复这些步骤。

4. 确认已删除丛：

```
storage aggregate plex show -fields aggregate , status , is-online , plex ,
pool
```

```
cluster_B::> storage aggregate plex show -fields aggregate,status,is-
online,Plex,pool
aggregate      plex  status          is-online pool
-----
node_B_1_aggr0 plex0 normal,active true      0
node_B_2_aggr0 plex0 normal,active true      0
node_B_1_aggr1 plex0 normal,active true      0
node_B_1_aggr2 plex0 normal,active true      0
node_B_2_aggr1 plex0 normal,active true      0
node_B_2_aggr2 plex0 normal,active true      0
node_A_1_aggr1 plex0 failed,inactive false    -
node_A_1_aggr1 plex4 normal,active true      1
node_A_1_aggr2 plex0 failed,inactive false    -
node_A_1_aggr2 plex1 normal,active true      1
node_A_2_aggr1 plex0 failed,inactive false    -
node_A_2_aggr1 plex4 normal,active true      1
node_A_2_aggr2 plex0 failed,inactive false    -
node_A_2_aggr2 plex1 normal,active true      1
14 entries were displayed.

cluster_B::>
```

5. 确定已切换的聚合：

```
storage aggregate show -is-home false
```

您还可以使用 `storage aggregate plex show -fields aggregate , status , is-online , plex , pool` 命令来标识从 0 已切换聚合。它们的状态将为 "失败，非活动"。

以下命令显示了四个已切换的聚合：

- node_A_1_aggr1
- node_A_1_aggr2
- node_A_2_aggr1
- node_A_2_aggr2.

```

cluster_B::> storage aggregate show -is-home false

cluster_A Switched Over Aggregates:
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
node_A_1_aggr1 2.12TB  1.88TB   11% online    91 node_B_1
raid_dp,

mirror

degraded
node_A_1_aggr2 2.89TB  2.64TB    9% online    90 node_B_1
raid_tec,

mirror

degraded
node_A_2_aggr1 2.12TB  1.86TB   12% online    91 node_B_2
raid_dp,

mirror

degraded
node_A_2_aggr2 2.89TB  2.64TB    9% online    90 node_B_2
raid_tec,

mirror

degraded
4 entries were displayed.

cluster_B::>

```

6. 识别已切换的丛：

```
storage aggregate plex show -fields aggregate , status , is-online , Plex , pool
```

您希望确定状态为 " 失败，非活动 " 的丛。

以下命令显示了四个已切换的聚合：

```
cluster_B::> storage aggregate plex show -fields aggregate,status,is-
online,Plex,pool
aggregate      plex  status          is-online pool
-----
node_B_1_aggr0 plex0 normal,active true      0
node_B_2_aggr0 plex0 normal,active true      0
node_B_1_aggr1 plex0 normal,active true      0
node_B_1_aggr2 plex0 normal,active true      0
node_B_2_aggr1 plex0 normal,active true      0
node_B_2_aggr2 plex0 normal,active true      0
node_A_1_aggr1 plex0 failed,inactive false - <<<<-- Switched over
aggr/Plex0
node_A_1_aggr1 plex4 normal,active true      1
node_A_1_aggr2 plex0 failed,inactive false - <<<<-- Switched over
aggr/Plex0
node_A_1_aggr2 plex1 normal,active true      1
node_A_2_aggr1 plex0 failed,inactive false - <<<<-- Switched over
aggr/Plex0
node_A_2_aggr1 plex4 normal,active true      1
node_A_2_aggr2 plex0 failed,inactive false - <<<<-- Switched over
aggr/Plex0
node_A_2_aggr2 plex1 normal,active true      1
14 entries were displayed.

cluster_B::>
```

7. 删除故障丛：

```
storage aggregate plex delete -aggregate node_A_1_aggr1 -plex plex0
```

出现提示时，您可以销毁丛。

以下示例显示了要删除的丛 node_A_1_aggr1/plex0：

```

cluster_B::> storage aggregate plex delete -aggregate node_A_1_aggr1
-plex plex0

Warning: Aggregate "node_A_1_aggr1" hosts MetroCluster metadata volume
"MDV_CRS_e8457659b8a711e78b3b00a0988fe74b_A". Deleting plex
"plex0"
      for this aggregate can lead to the failure of configuration
      replication across the two DR sites. Use the "volume show
-vserver
      <admin-vserver> -volume MDV_CRS*" command to verify the
location of
      such volumes.

Warning: Deleting plex "plex0" of mirrored aggregate "node_A_1_aggr1" on
node
      "node_A_1" in a MetroCluster configuration will disable its
      synchronous disaster recovery protection. Are you sure you want
to
      destroy this plex? {y|n}: y
[Job 639] Job succeeded: DONE

cluster_B::>

```

您必须对每个故障聚合重复这些步骤。

8. 验证运行正常的站点上是否没有剩余出现故障的丛。

以下输出显示所有丛均为正常，活动和联机。

```
cluster_B::> storage aggregate plex show -fields aggregate,status,is-
online,Plex,pool
aggregate      plex  status          is-online pool
-----
node_B_1_aggr0 plex0 normal,active true      0
node_B_2_aggr0 plex0 normal,active true      0
node_B_1_aggr1 plex0 normal,active true      0
node_B_2_aggr2 plex0 normal,active true      0
node_B_1_aggr1 plex0 normal,active true      0
node_B_2_aggr2 plex0 normal,active true      0
node_A_1_aggr1 plex4 normal,active true      1
node_A_1_aggr2 plex1 normal,active true      1
node_A_2_aggr1 plex4 normal,active true      1
node_A_2_aggr2 plex1 normal,active true      1
10 entries were displayed.

cluster_B::>
```

执行聚合修复和还原镜像（MetroCluster IP 配置）

更换硬件并分配磁盘后，在运行 ONTAP 9.5 或更早版本的系统中，您可以执行 MetroCluster 修复操作。然后，在所有版本的 ONTAP 中，您必须确认聚合已镜像，并在必要时重新启动镜像。

关于此任务

从 ONTAP 9.6 开始，灾难站点节点启动时会自动执行修复操作。不需要使用修复命令。

这些步骤将在运行正常的集群上执行。

步骤

1. 如果您使用的是 ONTAP 9.6 或更高版本，则必须验证自动修复是否已成功完成：

- a. 确认 heal-aggr-auto 和 heal-root-aggr-auto 操作已完成：

MetroCluster 操作历史记录显示

以下输出显示 cluster_A 上的操作已成功完成

```
cluster_B::*> metrocluster operation history show
```

Operation Time	State	Start Time	End
-----	-----	-----	
heal-root-aggr-auto	successful	2/25/2019 06:45:58	
2/25/2019 06:46:02			
heal-aggr-auto	successful	2/25/2019 06:45:48	
2/25/2019 06:45:52			
.			
.			
.			

b. 确认灾难站点已做好切回准备:

MetroCluster node show

以下输出显示 cluster_A 上的操作已成功完成

```
cluster_B::*> metrocluster node show
```

DR	Configuration	DR
Group Cluster Node	State	Mirroring Mode
-----	-----	-----
1 cluster_A		
node_A_1	configured	enabled heal roots
completed		
node_A_2	configured	enabled heal roots
completed		
cluster_B		
node_B_1	configured	enabled waiting for
switchback recovery		
node_B_2	configured	enabled waiting for
switchback recovery		
4 entries were displayed.		

2. 如果您使用的是 ONTAP 9.5 或更早版本，则必须执行聚合修复:

a. 验证节点的状态:

MetroCluster node show

以下输出显示切换已完成，因此可以执行修复。

```
cluster_B::> metrocluster node show
```

DR Group	Cluster	Node	Configuration State	DR Mirroring Mode
1	cluster_B	node_B_1	configured	enabled switchover
		node_B_2	configured	enabled switchover
	cluster_A	node_A_1	configured	enabled waiting for
		node_A_2	configured	enabled waiting for

4 entries were displayed.

```
cluster_B::>
```

b. 执行聚合修复阶段:

```
MetroCluster heal -phase aggregates
```

以下输出显示了一个典型的聚合修复操作。

```
cluster_B::*> metrocluster heal -phase aggregates
[Job 647] Job succeeded: Heal Aggregates is successful.

cluster_B::*> metrocluster operation show
  Operation: heal-aggregates
    State: successful
  Start Time: 10/26/2017 12:01:15
  End Time: 10/26/2017 12:01:17
  Errors: -

cluster_B::*>
```

c. 确认聚合修复已完成, 并且灾难站点已做好切回准备:

```
MetroCluster node show
```

以下输出显示 cluster_A 上的 " 修复聚合 " 阶段已完成

```
cluster_B::> metrocluster node show
DR
Group Cluster Node          Configuration State      DR
Mirroring Mode
-----
1      cluster_A
      node_A_1      configured  enabled   heal
aggregates completed
      node_A_2      configured  enabled   heal
aggregates completed
      cluster_B
      node_B_1      configured  enabled   waiting for
switchback recovery
      node_B_2      configured  enabled   waiting for
switchback recovery
4 entries were displayed.

cluster_B::>
```

3. 如果已更换磁盘，则必须镜像本地聚合和切换后的聚合：

a. 显示聚合：

s存储聚合显示

```
cluster_B::> storage aggregate show
cluster_B Aggregates:
Aggregate      Size Available Used% State  #Vols  Nodes
RAID Status
-----
node_B_1_aggr0 1.49TB  74.12GB   95% online    1 node_B_1
raid4,
normal
node_B_2_aggr0 1.49TB  74.12GB   95% online    1 node_B_2
raid4,
normal
node_B_1_aggr1 3.14TB  3.04TB    3% online   15 node_B_1
raid_dp,
normal
node_B_1_aggr2 3.14TB  3.06TB    3% online   14 node_B_1
raid_tec,
```

```

normal
node_B_1_aggr1 3.14TB  2.99TB    5% online    37 node_B_2
raid_dp,

normal
node_B_1_aggr2 3.14TB  3.02TB    4% online    35 node_B_2
raid_tec,

normal

cluster_A Switched Over Aggregates:
Aggregate      Size Available Used% State   #Vols  Nodes
RAID Status
-----
node_A_1_aggr1 2.36TB  2.12TB   10% online    91 node_B_1
raid_dp,

normal
node_A_1_aggr2 3.14TB  2.90TB    8% online    90 node_B_1
raid_tec,

normal
node_A_2_aggr1 2.36TB  2.10TB   11% online    91 node_B_2
raid_dp,

normal
node_A_2_aggr2 3.14TB  2.89TB    8% online    90 node_B_2
raid_tec,

normal
12 entries were displayed.

cluster_B::>

```

b. 镜像聚合：

```
storage aggregate mirror -aggregate aggregate-name
```

以下输出显示了典型的镜像操作。

```
cluster_B::> storage aggregate mirror -aggregate node_B_1_aggr1
```

Info: Disks would be added to aggregate "node_B_1_aggr1" on node "node_B_1" in the following manner:

Second Plex

	RAID Group rg0, 6 disks (block checksum, raid_dp)		
Size	Position	Disk	Type
	-----	-----	-----
	dparity	5.20.6	SSD
-	parity	5.20.14	SSD
-	data	5.21.1	SSD
894.0GB	data	5.21.3	SSD
894.0GB	data	5.22.3	SSD
894.0GB	data	5.21.13	SSD
894.0GB			

Aggregate capacity available for volume use would be 2.99TB.

Do you want to continue? {y|n}: y

- c. 对正常运行的站点中的每个聚合重复上述步骤。
- d. 等待聚合重新同步；您可以使用 `storage aggregate show` 命令检查状态。

以下输出显示许多聚合正在重新同步。

```
cluster_B::> storage aggregate show
```

cluster_B Aggregates:

Aggregate	Size	Available	Used%	State	#Vols	Nodes
RAID Status						
-----	-----	-----	-----	-----	-----	-----
node_B_1_aggr0	1.49TB	74.12GB	95%	online	1	node_B_1
raid4,						

```

mirrored,

normal
node_B_2_aggr0 1.49TB  74.12GB  95% online  1 node_B_2
raid4,

mirrored,

normal
node_B_1_aggr1 2.86TB  2.76TB  4% online  15 node_B_1
raid_dp,

resyncing
node_B_1_aggr2 2.89TB  2.81TB  3% online  14 node_B_1
raid_tec,

resyncing
node_B_2_aggr1 2.73TB  2.58TB  6% online  37 node_B_2
raid_dp,

resyncing
node_B-2_aggr2 2.83TB  2.71TB  4% online  35 node_B_2
raid_tec,

resyncing

cluster_A Switched Over Aggregates:
Aggregate      Size Available Used% State  #Vols  Nodes
RAID Status
-----
node_A_1_aggr1 1.86TB  1.62TB  13% online  91 node_B_1
raid_dp,

resyncing
node_A_1_aggr2 2.58TB  2.33TB  10% online  90 node_B_1
raid_tec,

resyncing
node_A_2_aggr1 1.79TB  1.53TB  14% online  91 node_B_2
raid_dp,

resyncing
node_A_2_aggr2 2.64TB  2.39TB  9% online  90 node_B_2
raid_tec,

```

```
resyncing
12 entries were displayed.
```

e. 确认所有聚合均已联机并已重新同步：

```
storage aggregate plex show
```

以下输出显示所有聚合均已重新同步。

```
cluster_A::> storage aggregate plex show
()
```

Aggregate Plex	Is Online	Is Resyncing	Resyncing Percent	Status
node_B_1_aggr0 plex0	true	false		- normal,active
node_B_1_aggr0 plex8	true	false		- normal,active
node_B_2_aggr0 plex0	true	false		- normal,active
node_B_2_aggr0 plex8	true	false		- normal,active
node_B_1_aggr1 plex0	true	false		- normal,active
node_B_1_aggr1 plex9	true	false		- normal,active
node_B_1_aggr2 plex0	true	false		- normal,active
node_B_1_aggr2 plex5	true	false		- normal,active
node_B_2_aggr1 plex0	true	false		- normal,active
node_B_2_aggr1 plex9	true	false		- normal,active
node_B_2_aggr2 plex0	true	false		- normal,active
node_B_2_aggr2 plex5	true	false		- normal,active
node_A_1_aggr1 plex4	true	false		- normal,active
node_A_1_aggr1 plex8	true	false		- normal,active
node_A_1_aggr2 plex1	true	false		- normal,active
node_A_1_aggr2 plex5	true	false		- normal,active
node_A_2_aggr1 plex4	true	false		- normal,active
node_A_2_aggr1 plex8	true	false		- normal,active
node_A_2_aggr2 plex1	true	false		- normal,active
node_A_2_aggr2 plex5	true	false		- normal,active

20 entries were displayed.

4. 在运行 ONTAP 9.5 及更早版本的系统上，执行根聚合修复阶段：

```
MetroCluster heal -phase root-aggregates
```

```
cluster_B::> metrocluster heal -phase root-aggregates
[Job 651] Job is queued: MetroCluster Heal Root Aggregates Job.Oct 26
13:05:00
[Job 651] Job succeeded: Heal Root Aggregates is successful.
```

5. 确认 " 修复根 " 阶段已完成, 并且灾难站点已做好切回准备:

以下输出显示 cluster_A 上的 " 修复根 " 阶段已完成

```
cluster_B::> metrocluster node show
DR                               Configuration  DR
Group Cluster Node              State          Mirroring Mode
-----
1      cluster_A
      node_A_1      configured    enabled    heal roots
completed
      node_A_2      configured    enabled    heal roots
completed
      cluster_B
      node_B_1      configured    enabled    waiting for
switchback recovery
      node_B_2      configured    enabled    waiting for
switchback recovery
4 entries were displayed.

cluster_B::>
```

继续验证已更换节点上的许可证。

["验证已更换节点上的许可证"](#)

准备在 MetroCluster FC 配置中切回

验证端口配置 (仅限 **MetroCluster FC** 配置)

您必须在节点上设置环境变量, 然后将其关闭, 以便为 MetroCluster 配置做好准备。

关于此任务

在更换用的控制器模块处于维护模式时执行此操作步骤。

只有在启动程序模式下使用 FC 或 CNA 端口的系统上, 才需要执行检查端口配置的步骤。

步骤

1. 在维护模式下，还原 FC 端口配置：

```
ucadmin modify -m fc -t initiatoradapter_name
```

如果您只想在启动程序配置中使用端口对之一，请输入精确的适配器名称。

2. 根据您的配置，执行以下操作之一：

FC 端口配置	那么 ...
这两个端口相同	系统提示时出现问题解答 "y` "，因为修改端口对中的一个端口也会修改另一个端口。
不同	a. 系统提示时出现问题解答 "n` "。</div> <div>b. 还原 FC 端口配置：</div> <div>`ucadmin modify -m fc -t initiators

3. 退出维护模式：

```
halt
```

问题描述命令后，请等待，直到系统停留在 LOADER 提示符处。

4. 将节点重新启动至维护模式，以使配置更改生效：

```
boot_ontap maint
```

5. 验证变量的值：

```
ucadmin show
```

6. 退出维护模式并显示 LOADER 提示符：

```
halt
```

配置 **FC-SAS** 网桥（仅限 **MetroCluster FC** 配置）

如果您更换了 FC-SAS 网桥，则必须在还原 MetroCluster 配置时对其进行配置。操作步骤与 FC-SAS 网桥的初始配置相同。

步骤

1. 打开 FC-SAS 网桥的电源。
2. 使用 `set IPAddress port ipaddress` 命令设置以太网端口上的 IP 地址。
 - 端口 可以是 "MP1" 或 "MP2" 。
 - ipaddress 可以是 xxx.xxx.xxx.xxx 格式的 IP 地址。

在以下示例中，以太网端口 1 上的 IP 地址为 10.10.10.55：

```
Ready.  
set IPAddress MP1 10.10.10.55  
  
Ready. *
```

3. 使用 `set IPSubnetMask port mask` 命令设置以太网端口上的 IP 子网掩码。

- 端口 可以是 "MP1" 或 "MP2"。
- `mAsk` 可以是 xxx.xxx.xxx.xxx 格式的子网掩码。

在以下示例中，以太网端口 1 上的 IP 子网掩码为 255.255.255.0：

```
Ready.  
set IPSubnetMask MP1 255.255.255.0  
  
Ready. *
```

4. 使用 `set EthernetSpeed port speed` 命令设置以太网端口上的速度。

- 端口 可以是 "MP1" 或 "MP2"。
- `s` 对等 可以是 "100" 或 "1000"。

在以下示例中，以太网端口 1 上的以太网速度设置为 1000。

```
Ready.  
set EthernetSpeed MP1 1000  
  
Ready. *
```

5. 使用 `saveConfiguration` 命令保存配置，并在系统提示时重新启动网桥。

通过在配置以太网端口后保存配置，您可以使用 Telnet 继续配置网桥，并可以使用 FTP 访问网桥以执行固件更新。

以下示例显示了 `saveConfiguration` 命令以及重新启动网桥的提示。

```
Ready.  
SaveConfiguration  
  Restart is necessary....  
  Do you wish to restart (y/n) ?  
Confirm with 'y'. The bridge will save and restart with the new  
settings.
```

6. 在 FC-SAS 网桥重新启动后，重新登录。

7. 使用 `set fcdatarate port speed` 命令设置 FC 端口的速度。

- 端口 可以是 "1" 或 "2"。
- s对等 可以是 "2 GB"，"4 GB"，"8 GB" 或 "16 GB"，具体取决于您的网桥型号。

在以下示例中，端口 FC1 速度设置为 "8 GB"。

```
Ready.  
set fcdatarate 1 8Gb  
  
Ready. *
```

8. 使用 `set FCConnMode port mode` 命令设置 FC 端口上的拓扑。

- 端口 可以是 "1" 或 "2"。
- mode 可以是 "PTp"，"loop"，"PTP-loop" 或 "autode"。

在以下示例中，端口 FC1 拓扑设置为 "PPT"。

```
Ready.  
set FCConnMode 1 ptp  
  
Ready. *
```

9. 使用 `saveConfiguration` 命令保存配置，并在系统提示时重新启动网桥。

以下示例显示了 `saveConfiguration` 命令以及重新启动网桥的提示。

```
Ready.  
SaveConfiguration  
  Restart is necessary....  
  Do you wish to restart (y/n) ?  
Confirm with 'y'. The bridge will save and restart with the new  
settings.
```

10. 在 FC-SAS 网桥重新启动后，重新登录。
11. 如果 FC-SAS 网桥运行的是固件 1.60 或更高版本，请启用 SNMP。

```
Ready.  
set snmp enabled  
  
Ready. *  
saveconfiguration  
  
Restart is necessary....  
Do you wish to restart (y/n) ?  
  
Verify with 'y' to restart the FibreBridge.
```

12. 关闭 FC-SAS 网桥的电源。

配置 FC 交换机（仅限 **MetroCluster FC** 配置）

如果您更换了灾难站点中的 FC 交换机，则必须使用供应商专用的过程对其进行配置。您必须配置一个交换机，验证运行正常的站点上的存储访问不受影响，然后配置第二个交换机。

相关任务

["FC 交换机的端口分配"](#)

在站点灾难后配置 **Brocade FC** 交换机

您必须使用此 Brocade 专用的操作步骤配置替代交换机并启用 ISL 端口。

关于此任务

此操作步骤中的示例基于以下假设：

- 站点 A 是灾难站点。
- 已更换 FC_switch_A_1。
- 已更换 FC_switch_A_2。
- 站点 B 是正常运行的站点。
- FC_switch_B_1 运行状况良好。
- FC_switch_B_2 运行状况良好。

在为 FC 交换机布线时，您必须验证是否正在使用指定的端口分配：

- ["FC 交换机的端口分配"](#)

这些示例显示了两个 FC-SAS 网桥。如果有更多网桥，则必须禁用并随后启用其他端口。

步骤

1. 启动并预配置新交换机：

- a. 启动新交换机并让其启动。
- b. 检查交换机上的固件版本以确认其与其他 FC 交换机的版本匹配：

固件

- c. 按照以下主题中所述配置新交换机，跳过在交换机上配置分区的步骤。

["光纤连接的 MetroCluster 安装和配置"](#)

["延伸型 MetroCluster 安装和配置"](#)

- d. 持久禁用交换机：

```
sswitchcfgpersistentdisable
```

重新启动或快速启动后，交换机将保持禁用状态。如果此命令不可用，则应使用 `sswitch- disable` 命令。

以下示例显示了对 BrocadeSwitchA 执行的命令：

```
BrocadeSwitchA:admin> switchcfgpersistentdisable
```

以下示例显示了对 BrocadeSwitchB 执行的命令：

```
BrocadeSwitchA:admin> switchcfgpersistentdisable
```

2. 完成新交换机的配置：

- a. 在运行正常的站点上启用 ISL：

```
portcfgpersistentenable port-number
```

```
FC_switch_B_1:admin> portcfgpersistentenable 10  
FC_switch_B_1:admin> portcfgpersistentenable 11
```

- b. 在替代交换机上启用 ISL：

```
portcfgpersistentenable port-number
```

```
FC_switch_A_1:admin> portcfgpersistentenable 10  
FC_switch_A_1:admin> portcfgpersistentenable 11
```

c. 在替代交换机（在此示例中为 FC_switch_A_1）上，验证 ISL 是否联机：

```
sswitchshow
```

```
FC_switch_A_1:admin> switchshow
switchName: FC_switch_A_1
switchType: 71.2
switchState:Online
switchMode: Native
switchRole: Principal
switchDomain:      4
switchId:   fffc03
switchWwn:  10:00:00:05:33:8c:2e:9a
zoning:      OFF
switchBeacon: OFF

Index Port Address Media Speed State Proto
=====
...
10   10   030A00 id   16G   Online  FC E-Port 10:00:00:05:33:86:89:cb
"FC_switch_A_1"
11   11   030B00 id   16G   Online  FC E-Port 10:00:00:05:33:86:89:cb
"FC_switch_A_1" (downstream)
...
```

3. 持久启用交换机：

```
sswitchcfgpersistentenable
```

4. 验证端口是否联机：

```
sswitchshow
```

在站点灾难后配置 **Cisco FC** 交换机

您必须使用 Cisco 专用的操作步骤配置替代交换机并启用 ISL 端口。

关于此任务

此操作步骤中的示例基于以下假设：

- 站点 A 是灾难站点。
- 已更换 FC_switch_A_1。
- 已更换 FC_switch_A_2。
- 站点 B 是正常运行的站点。
- FC_switch_B_1 运行状况良好。

- FC_switch_B_2 运行状况良好。

步骤

1. 配置交换机：

- a. 请参见 ["光纤连接的 MetroCluster 安装和配置"](#)
- b. 按照中的步骤配置交换机 ["配置 Cisco FC 交换机"](#) 第节 ["在 Cisco FC 交换机上配置分区"](#) 部分的 `_except_`：

分区将在此操作步骤中稍后进行配置。

2. 在运行正常的交换机（在此示例中为 FC_switch_B_1）上，启用 ISL 端口。

以下示例显示了用于启用端口的命令：

```
FC_switch_B_1# conf t
FC_switch_B_1(config)# int fc1/14-15
FC_switch_B_1(config)# no shut
FC_switch_B_1(config)# end
FC_switch_B_1# copy running-config startup-config
FC_switch_B_1#
```

3. 使用 show interface brief 命令验证 ISL 端口是否已启动。

4. 从网络结构中检索分区信息。

以下示例显示了用于分发分区配置的命令：

```
FC_switch_B_1(config-zone)# zoneset distribute full vsan 10
FC_switch_B_1(config-zone)# zoneset distribute full vsan 20
FC_switch_B_1(config-zone)# end
```

FC_switch_B_1 将分发到 "vsan 10" 和 "vsan 20" 网络结构中的所有其他交换机，分区信息将从 FC_switch_A_1 中检索。

5. 在运行状况良好的交换机上，验证是否已从配对交换机正确检索分区信息：

s如何分区

```

FC_switch_B_1# show zone
zone name FC-VI_Zone_1_10 vsan 10
  interface fc1/1 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/2 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/1 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/2 swwn 20:00:54:7f:ee:b8:24:c0

zone name STOR_Zone_1_20_25A vsan 20
  interface fc1/5 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/8 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/9 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/10 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/11 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/8 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/9 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/10 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/11 swwn 20:00:54:7f:ee:b8:24:c0

zone name STOR_Zone_1_20_25B vsan 20
  interface fc1/8 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/9 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/10 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/11 swwn 20:00:54:7f:ee:e3:86:50
  interface fc1/5 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/8 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/9 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/10 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/11 swwn 20:00:54:7f:ee:b8:24:c0
FC_switch_B_1#

```

6. 确定交换机网络结构中交换机的全球通用名称（WWN）。

在此示例中，两个交换机 WWN 如下所示：

- FC_switch_A_1： 20： 00： 54： 7f： ee： B8： 24： c0
- FC_switch_B_1： 20： 00： 54： 7f： ee： c6： 80： 78

```

FC_switch_B_1# show wwn switch
Switch WWN is 20:00:54:7f:ee:c6:80:78
FC_switch_B_1#

FC_switch_A_1# show wwn switch
Switch WWN is 20:00:54:7f:ee:b8:24:c0
FC_switch_A_1#

```

7. 进入分区的配置模式，然后删除不属于这两个交换机的交换机 WWN 的分区成员：

无成员接口 `interface-ide swwn WWN`

在此示例中，以下成员不与网络结构中任一交换机的 WWN 关联，必须将其删除：

◦ 分区名称 `FC-VI_Zone_1_10 vsan 10`

- 接口 `fc1/1 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/2 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`



AFF A700 和 FAS9000 系统支持四个 FC-VI 端口。您必须从 FC-VI 区域中删除所有四个端口。

◦ 分区名称 `STOR_Zone_1_20_25 a vsan 20`

- 接口 `fc1/5 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/8 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/9 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/10 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/11 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`

◦ 分区名称 `STOR_Zone_1_20_25B vSAN 20`

- 接口 `fc1/8 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/9 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/10 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`
- 接口 `fc1/11 swwn 20 : 00 : 54 : 7f : ee : e3 : 86 : 50`

以下示例显示了如何删除这些接口：

```

FC_switch_B_1# conf t
FC_switch_B_1(config)# zone name FC-VI_Zone_1_10 vsan 10
FC_switch_B_1(config-zone)# no member interface fc1/1 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/2 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# zone name STOR_Zone_1_20_25A vsan 20
FC_switch_B_1(config-zone)# no member interface fc1/5 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/8 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/9 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/10 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/11 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# zone name STOR_Zone_1_20_25B vsan 20
FC_switch_B_1(config-zone)# no member interface fc1/8 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/9 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/10 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# no member interface fc1/11 swwn
20:00:54:7f:ee:e3:86:50
FC_switch_B_1(config-zone)# save running-config startup-config
FC_switch_B_1(config-zone)# zoneset distribute full 10
FC_switch_B_1(config-zone)# zoneset distribute full 20
FC_switch_B_1(config-zone)# end
FC_switch_B_1# copy running-config startup-config

```

8. 【第 8 步】将新交换机的端口添加到分区中。

以下示例假设替代交换机上的布线与旧交换机上的布线相同：

```

FC_switch_B_1# conf t
FC_switch_B_1(config)# zone name FC-VI_Zone_1_10 vsan 10
FC_switch_B_1(config-zone)# member interface fc1/1 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/2 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# zone name STOR_Zone_1_20_25A vsan 20
FC_switch_B_1(config-zone)# member interface fc1/5 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/8 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/9 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/10 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/11 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# zone name STOR_Zone_1_20_25B vsan 20
FC_switch_B_1(config-zone)# member interface fc1/8 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/9 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/10 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# member interface fc1/11 swwn
20:00:54:7f:ee:c6:80:78
FC_switch_B_1(config-zone)# save running-config startup-config
FC_switch_B_1(config-zone)# zoneset distribute full 10
FC_switch_B_1(config-zone)# zoneset distribute full 20
FC_switch_B_1(config-zone)# end
FC_switch_B_1# copy running-config startup-config

```

9. 验证是否已正确配置分区：show zone

以下示例输出显示了三个分区：

```

FC_switch_B_1# show zone
zone name FC-VI_Zone_1_10 vsan 10
  interface fc1/1 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/2 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/1 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/2 swwn 20:00:54:7f:ee:b8:24:c0

zone name STOR_Zone_1_20_25A vsan 20
  interface fc1/5 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/8 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/9 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/10 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/11 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/8 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/9 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/10 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/11 swwn 20:00:54:7f:ee:b8:24:c0

zone name STOR_Zone_1_20_25B vsan 20
  interface fc1/8 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/9 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/10 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/11 swwn 20:00:54:7f:ee:c6:80:78
  interface fc1/5 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/8 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/9 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/10 swwn 20:00:54:7f:ee:b8:24:c0
  interface fc1/11 swwn 20:00:54:7f:ee:b8:24:c0
FC_switch_B_1#

```

验证存储配置

您必须确认所有存储均可从正常运行的节点中查看。

步骤

1. 确认灾难站点上的所有存储组件在正常运行的站点上的数量和类型都相同。

正常运行的站点和灾难站点应具有相同数量的磁盘架堆栈，磁盘架和磁盘。在桥接或光纤连接的 MetroCluster 配置中，站点应具有相同数量的 FC-SAS 网桥。

2. 确认灾难站点上已更换的所有磁盘均为无主磁盘：

运行本地磁盘 `show-n`

磁盘应显示为未分配。

3. 如果未更换任何磁盘，请确认所有磁盘均存在：

d展示

启动灾难站点上的设备

准备切回时，您必须打开灾难站点上的 MetroCluster 组件的电源。此外，您还必须在直连 MetroCluster 配置中重新连接 SAS 存储连接，并在光纤连接 MetroCluster 配置中启用非交换机间链路端口。

开始之前

您必须已更换 MetroCluster 组件并完全按照旧组件的方式为其布线。

["光纤连接的 MetroCluster 安装和配置"](#)

["延伸型 MetroCluster 安装和配置"](#)

关于此任务

此操作步骤中的示例假设如下：

- 站点 A 是灾难站点。
 - 已更换 FC_switch_A_1。
 - 已更换 FC_switch_A_2。
- 站点 B 是正常运行的站点。
 - FC_switch_B_1 运行状况良好。
 - FC_switch_B_2 运行状况良好。

FC 交换机仅存在于光纤连接的 MetroCluster 配置中。

步骤

1. 在使用 SAS 布线的延伸型 MetroCluster 配置（不使用 FC 交换机网络结构或 FC-SAS 网桥）中，跨两个站点连接所有存储，包括远程存储。

灾难站点上的控制器必须保持关闭状态，或者在 LOADER 提示符处保持关闭状态。

2. 在运行正常的站点上，禁用磁盘自动分配：

s存储磁盘选项 `modify -autosign off *`

```
cluster_B::> storage disk option modify -autoassign off *  
2 entries were modified.
```

3. 在正常运行的站点上，确认磁盘自动分配已关闭：

s存储磁盘选项 `show`

```
cluster_B::> storage disk option show
Node      BKg. FW. Upd.  Auto Copy  Auto Assign  Auto Assign Policy
-----
node_B_1      on          on         off          default
node_B_2      on          on         off          default
2 entries were displayed.

cluster_B::>
```

4. 打开灾难站点上的磁盘架，并确保所有磁盘都在运行。
5. 在桥接或光纤连接的 MetroCluster 配置中，打开灾难站点上的所有 FC-SAS 网桥。
6. 如果更换了任何磁盘，请保持控制器关闭或处于 LOADER 提示符处。
7. 在光纤连接的 MetroCluster 配置中，启用 FC 交换机上的非 ISL 端口。

交换机供应商	然后，按照以下步骤启用端口 ...
--------	-------------------

- a. 持久启用连接到 FC-SAS 网桥的端口：
`portpersistentenable port-number`

在以下示例中，端口 6 和 7 已启用：

```
FC_switch_A_1:admin>  
portpersistentenable 6  
FC_switch_A_1:admin>  
portpersistentenable 7  
  
FC_switch_A_1:admin>
```

- b. 持久启用连接到 HBA 和 FC-VI 适配器的端口：
`portpersistentenable port-number`

在以下示例中，端口 6 和 7 已启用：

```
FC_switch_A_1:admin>  
portpersistentenable 1  
FC_switch_A_1:admin>  
portpersistentenable 2  
FC_switch_A_1:admin>  
portpersistentenable 4  
FC_switch_A_1:admin>  
portpersistentenable 5  
FC_switch_A_1:admin>
```



对于 AFF A700 和 FAS9000 系统，必须使用 `switchcfgpersistentenable` 命令持久启用所有四个 FC-VI 端口。

- c. 对运行正常的站点上的第二个 FC 交换机重复子步骤 a 和 b。

Cisco	a. 进入接口的配置模式，然后使用 no shut 命令启用端口。 在以下示例中，禁用了端口 fc1/36： <pre>FC_switch_A_1# conf t FC_switch_A_1(config)# interface fc1/36 FC_switch_A_1(config)# no shut FC_switch_A_1(config-if)# end FC_switch_A_1# copy running- config startup-config</pre> b. 验证交换机端口是否已启用： sHow interface brief c. 对连接到 FC-SAS 网桥， HBA 和 FC-VI 适配器的其他端口重复子步骤 a 和 b。 d. 对运行正常的站点上的第二个 FC 交换机重复子步骤 a， b 和 c。
-------	---

为更换的驱动器分配所有权

如果在灾难站点还原硬件时更换了驱动器，或者必须将驱动器置零或删除所有权，则必须为受影响的驱动器分配所有权。

开始之前

灾难站点的可用驱动器数量必须至少与灾难发生前相同。

驱动器架和驱动器布局必须满足中的要求 ["所需的 MetroCluster IP 组件和命名约定"](#) 部分 ["MetroCluster IP 安装和配置"](#)。

关于此任务

这些步骤在灾难站点的集群上执行。

此操作步骤显示了所有驱动器的重新分配以及在灾难站点上创建新丛的情况。新丛是运行正常的站点的远程丛和灾难站点的本地丛。

本节提供了双节点和四节点配置的示例。对于双节点配置，您可以忽略对每个站点上第二个节点的引用。对于八节点配置，您必须考虑第二个 DR 组上的其他节点。这些示例假设以下条件：

- 站点 A 是灾难站点。
 - node_A_1 已更换。
 - node_A_2 已更换。

仅存在于四节点 MetroCluster 配置中。

- 站点 B 是正常运行的站点。
 - node_B_1 运行状况良好。
 - node_B_2 运行状况良好。

仅存在于四节点 MetroCluster 配置中。

控制器模块具有以下原始系统 ID：

MetroCluster 配置中的节点数	节点	原始系统 ID
四个	node_A_1	4068741258
node_A_2	4068741260	node_B_1
4068741254	node_B_2	4068741256
两个	node_A_1	4068741258

在分配驱动器时，应记住以下几点：

- 对于灾难发生前存在的每个节点，old-count-of-disks 必须至少具有相同数量的磁盘。

如果指定的或存在的磁盘数量较少，则修复操作可能会由于空间不足而无法完成。

- 要创建的新丛是属于正常运行的站点（node_B_x pool1）的远程丛和属于灾难站点的本地丛（node_B_x pool0）。
- 所需驱动器总数不应包括根聚合磁盘。

如果将 n 个磁盘分配给运行正常的站点的 pool1，则应将 n-3 个磁盘分配给灾难站点，并假设根聚合使用三个磁盘。

- 不能将任何磁盘分配给与同一堆栈中所有其他磁盘分配到的池不同的池。
- 属于正常运行的站点的磁盘将分配给池 1，属于灾难站点的磁盘将分配给池 0。

步骤

1. 根据您使用的是四节点还是双节点 MetroCluster 配置，分配新的无主驱动器：

- 对于四节点 MetroCluster 配置，请在替代节点上使用以下一系列命令将新的无主磁盘分配给相应的磁盘池：
 - i. 系统地将每个节点的已更换磁盘分配给各自的磁盘池：

```
dassign -s sysid -n old-count-of-disks -p pool
```

在运行正常的站点中，您可以为每个节点问题描述一个 disk assign 命令：

```
cluster_B::> disk assign -s node_B_1-sysid -n old-count-of-disks  
-p 1 **\ (remote pool of surviving site\)**  
cluster_B::> disk assign -s node_B_2-sysid -n old-count-of-disks  
-p 1 **\ (remote pool of surviving site\)**  
cluster_B::> disk assign -s node_A_1-old-sysid -n old-count-of-  
disks -p 0 **\ (local pool of disaster site\)**  
cluster_B::> disk assign -s node_A_2-old-sysid -n old-count-of-  
disks -p 0 **\ (local pool of disaster site\)**
```

以下示例显示了具有系统 ID 的命令：

```
cluster_B::> disk assign -s 4068741254 -n 21 -p 1  
cluster_B::> disk assign -s 4068741256 -n 21 -p 1  
cluster_B::> disk assign -s 4068741258 -n 21 -p 0  
cluster_B::> disk assign -s 4068741260 -n 21 -p 0
```

i. 确认磁盘的所有权：

```
storage disk show -fields owner , pool
```

```
storage disk show -fields owner, pool
cluster_A:> storage disk show -fields owner, pool
disk      owner      pool
-----
0c.00.1   node_A_1      Pool0
0c.00.2   node_A_1      Pool0
.
.
.
0c.00.8   node_A_1      Pool1
0c.00.9   node_A_1      Pool1
.
.
.
0c.00.15  node_A_2      Pool0
0c.00.16  node_A_2      Pool0
.
.
.
0c.00.22  node_A_2      Pool1
0c.00.23  node_A_2      Pool1
.
.
.
```

- 对于双节点 MetroCluster 配置，请在替代节点上使用以下一系列命令将新的无主磁盘分配给相应的磁盘池：

- i. 显示本地磁盘架 ID：

运行 `local storage show shelf`

- ii. 将运行状况良好的节点的已更换磁盘分配到池 1：

运行本地磁盘分配 `-shelf shelf-id -n old-count-of-disks -p 1 -s node_B_1-sysid -f`

- iii. 将替代节点的已更换磁盘分配到池 0：

运行本地磁盘分配 `-shelf shelf-id -n old-count-of-disks -p 0 -s node_A_1-sysid -f`

2. 在运行正常的站点上，再次启用自动磁盘分配：

`storage disk option modify -autosassign on *`

```
cluster_B::> storage disk option modify -autoassign on *
2 entries were modified.
```

3. 在正常运行的站点上，确认自动磁盘分配已启用：

s存储磁盘选项 show

```
cluster_B::> storage disk option show
Node      BKg.  FW.  Upd.  Auto Copy  Auto Assign  Auto Assign Policy
-----
node_B_1      on           on          on          default
node_B_2      on           on          on          default
2 entries were displayed.

cluster_B::>
```

相关信息

["磁盘和聚合管理"](#)

["MetroCluster 配置如何使用 SyncMirror 提供数据冗余"](#)

执行聚合修复和还原镜像（**MetroCluster FC** 配置）

更换硬件并分配磁盘后，您可以执行 MetroCluster 修复操作。然后，您必须确认聚合已镜像，并在必要时重新启动镜像。

步骤

1. 在灾难站点上执行两个修复阶段（聚合修复和根修复）：

```
cluster_B::> metrocluster heal -phase aggregates

cluster_B::> metrocluster heal -phase root-aggregates
```

2. 监控修复过程，并验证聚合是否处于重新同步或镜像状态：

storage aggregate show -node local

如果聚合显示此状态 ...	那么 ...
正在重新同步	无需执行任何操作。让聚合完成重新同步。
镜像已降级	继续执行 如果一个或多个丛保持脱机，则需要执行其他步骤来重建镜像。

已镜像，正常	无需执行任何操作。
未知，脱机	如果更换了灾难站点上的所有磁盘，则根聚合将显示此状态。

```
cluster_B::> storage aggregate show -node local

Aggregate      Size Available Used% State  #Vols  Nodes      RAID
Status
-----
node_B_1_aggr1
      227.1GB    11.00GB   95% online      1 node_B_1  raid_dp,
resyncing

NodeA_1_aggr2
      430.3GB    28.02GB   93% online      2 node_B_1  raid_dp,
mirror
degraded

node_B_1_aggr3
      812.8GB    85.37GB   89% online      5 node_B_1  raid_dp,
mirrored,
normal

3 entries were displayed.

cluster_B::>
```

在以下示例中，三个聚合各自处于不同的状态：

节点	状态
node_B_1_aggr1	正在重新同步
node_B_1_aggr2	镜像已降级
node_B_1_aggr3.	已镜像，正常

3. 【第 3 步 _fc_aggr_healing】如果一个或多个丛保持脱机状态，则需要执行其他步骤来重建镜像。

在上表中，必须重建 node_B_1_aggr2 的镜像。

a. 查看聚合的详细信息以确定任何出现故障的丛：

```
storage aggregate show -r -aggregate node_B_1_aggr2
```

在以下示例中，丛 /node_B_1_aggr2/plex0 处于故障状态：

```

cluster_B::> storage aggregate show -r -aggregate node_B_1_aggr2

Owner Node: node_B_1
Aggregate: node_B_1_aggr2 (online, raid_dp, mirror degraded) (block
checksums)
Plex: /node_B_1_aggr2/plex0 (offline, failed, inactive, pool0)
RAID Group /node_B_1_aggr2/plex0/rg0 (partial)
Usable
Physical
Position Disk Pool Type RPM Size
Size Status
-----
-----

Plex: /node_B_1_aggr2/plex1 (online, normal, active, pool1)
RAID Group /node_B_1_aggr2/plex1/rg0 (normal, block checksums)
Usable
Physical
Position Disk Pool Type RPM Size
Size Status
-----
-----

dparity 1.44.8 1 SAS 15000 265.6GB
273.5GB (normal)
parity 1.41.11 1 SAS 15000 265.6GB
273.5GB (normal)
data 1.42.8 1 SAS 15000 265.6GB
273.5GB (normal)
data 1.43.11 1 SAS 15000 265.6GB
273.5GB (normal)
data 1.44.9 1 SAS 15000 265.6GB
273.5GB (normal)
data 1.43.18 1 SAS 15000 265.6GB
273.5GB (normal)
6 entries were displayed.

cluster_B::>

```

a. 删除故障丛:

```
storage aggregate plex delete -aggregate aggregate-name -plex plex
```

b. 重新建立镜像:

```
storage aggregate mirror -aggregate aggregate-name
```

c. 监控丛的重新同步和镜像状态，直到所有镜像均已重新建立且所有聚合均显示已镜像，正常状态：

s存储聚合显示

将根聚合的磁盘所有权重新分配给更换用的控制器模块（ **MetroCluster FC** 配置）

如果在灾难站点上更换了一个或两个控制器模块或 NVRAM 卡，则系统 ID 已更改，您必须将属于根聚合的磁盘重新分配给更换用的控制器模块。

关于此任务

由于节点处于切换模式且已完成修复，因此，本节仅会重新分配包含灾难站点的 pool1 根聚合的磁盘。此时，它们是唯一仍归旧系统 ID 所有的磁盘。

本节提供了双节点和四节点配置的示例。对于双节点配置，您可以忽略对每个站点上第二个节点的引用。对于八节点配置，您必须考虑第二个 DR 组上的其他节点。这些示例假设以下条件：

- 站点 A 是灾难站点。
 - node_A_1 已更换。
 - node_A_2 已更换。

仅存在于四节点 MetroCluster 配置中。

- 站点 B 是正常运行的站点。
 - node_B_1 运行状况良好。
 - node_B_2 运行状况良好。

仅存在于四节点 MetroCluster 配置中。

旧的和新的系统ID在中进行了标识 ["更换硬件并启动新控制器"](#)。

此操作步骤中的示例使用具有以下系统 ID 的控制器：

节点数	节点	原始系统 ID	新系统 ID
四个	node_A_1	4068741258	1574774970
	node_A_2	4068741260	1574774991
	node_B_1	4068741254	未更改
	node_B_2	4068741256	未更改
两个	node_A_1	4068741258	1574774970

步骤

1. 在替代节点处于维护模式的情况下，重新分配根聚合磁盘：

```
dreassign -s old-system-ID -d new-system-ID
```

```
*> disk reassign -s 4068741258 -d 1574774970
```

2. 查看磁盘以确认灾难站点的 pool1 根聚合磁盘的所有权已更改为替代节点：

d展示

输出可能会显示更多或更少的磁盘，具体取决于根聚合中的磁盘数量以及这些磁盘中是否有任何磁盘发生故障并已更换。如果更换了磁盘，则 Pool0 磁盘将不会显示在输出中。

现在，应将灾难站点的 pool1 根聚合磁盘分配给替代节点。

```
*> disk show
Local System ID: 1574774970

  DISK              OWNER              POOL  SERIAL NUMBER  HOME
DR HOME
-----
sw_A_1:6.126L19    node_A_1(1574774970) Pool0  serial-number
node_A_1(1574774970)
sw_A_1:6.126L3     node_A_1(1574774970) Pool0  serial-number
node_A_1(1574774970)
sw_A_1:6.126L7     node_A_1(1574774970) Pool0  serial-number
node_A_1(1574774970)
sw_B_1:6.126L8     node_A_1(1574774970) Pool1  serial-number
node_A_1(1574774970)
sw_B_1:6.126L24    node_A_1(1574774970) Pool1  serial-number
node_A_1(1574774970)
sw_B_1:6.126L2     node_A_1(1574774970) Pool1  serial-number
node_A_1(1574774970)

*> aggr status

      Aggr State      Status
node_A_1_root online  raid_dp, aggr
                      mirror degraded
                      64-bit

*>
```

3. 查看聚合状态：

聚合状态

输出可能会显示更多或更少的磁盘，具体取决于根聚合中的磁盘数量以及这些磁盘中是否有任何磁盘发生故障并已更换。如果更换了磁盘，则 Pool0 磁盘将不会显示在输出中。

```
*> aggr status
      Aggr State      Status
node_A_1_root online  raid_dp, aggr
                      mirror degraded
                      64-bit
*>
```

4. 删除邮箱磁盘的内容：

m邮箱销毁本地

5. 如果聚合未联机，请将其联机：

```
aggr online aggr_name
```

6. 暂停节点以显示 LOADER 提示符：

```
halt
```

启动新控制器模块（**MetroCluster FC** 配置）

数据聚合和根聚合的聚合修复完成后，您必须启动灾难站点上的一个或多个节点。

关于此任务

此任务从显示 LOADER 提示符的节点开始。

步骤

1. 显示启动菜单：

```
boot_ontap 菜单
```

2. 在启动菜单中，选择选项 6 * 从备份配置更新闪存 *。

3. 对以下提示回答 y：

此操作将使用上次备份到磁盘来替换所有基于闪存的配置。确实要继续？： y

系统将启动两次，第二次加载新配置。



如果未清除已用替代控制器的 NVRAM 内容，则可能会出现崩溃，并显示以下消息：panic : NVRAM contents are invalid... 如果发生这种情况，请重复 [从启动菜单中，选择选项 6 * 从备份配置更新闪存 *](#)。将系统启动至 ONTAP 提示符。然后、您需要 [重置启动恢复和rdb_Corrupt bootargs](#)

4. 在丛 0 上镜像根聚合：

- a. 将三个 pool0 磁盘分配给新控制器模块。
- b. 镜像根聚合 pool1 丛：

```
aggr mirror root-aggr-name
```

c. 将无主磁盘分配给本地节点上的 pool0

5. 如果您使用的是四节点配置，请在灾难站点的另一个节点上重复上述步骤。

6. 刷新 MetroCluster 配置：

a. 进入高级权限模式：

```
set -privilege advanced
```

b. 刷新配置：

```
MetroCluster configure -refresh true
```

c. 返回到管理权限模式：

```
set -privilege admin
```

7. 确认灾难站点上的替代节点已做好切回准备：

```
MetroCluster node show
```

替代节点应处于 "Waiting for switchback recovery" 模式。如果它们处于 "normal" 模式，则可以重新启动替代节点。启动后，节点应处于 "Waiting for switchback recovery" 模式。

以下示例显示替代节点已做好切回准备：

```
cluster_B::> metrocluster node show
DR
Grp Cluster Node      Configuration  DR
State          Mirroring Mode
-----
1   cluster_B
      node_B_1  configured    enabled      switchover completed
      node_B_2  configured    enabled      switchover completed
      cluster_A
      node_A_1  configured    enabled      waiting for switchback
recovery
      node_A_2  configured    enabled      waiting for switchback
recovery
4 entries were displayed.

cluster_B::>
```

下一步操作

继续执行 ["完成灾难恢复过程"](#)。

【重置启动恢复】重置boot_recovery和rdb_Corrupt bootargs

如果需要、您可以重置boot_recovery和rdb_Corrupt_bootargs

步骤

1. 将节点暂停回LOADER提示符：

```
siteA::*> halt -node <node-name>
```

2. 检查是否已设置以下bootarg：

```
LOADER> printenv bootarg.init.boot_recovery  
LOADER> printenv bootarg.rdb_corrupt
```

3. 如果已将任一bootarg设置为值、请取消设置并启动ONTAP：

```
LOADER> unsetenv bootarg.init.boot_recovery  
LOADER> unsetenv bootarg.rdb_corrupt  
LOADER> saveenv  
LOADER> bye
```

在混合配置中准备切回（过渡期间恢复）

您必须执行某些任务，以便为切回操作准备混合 MetroCluster IP 和 FC 配置。此操作步骤仅适用于在 MetroCluster FC 到 IP 过渡过程中遇到故障的适用场景配置。

关于此任务

只有在发生故障时处于过渡中的系统上执行恢复时，才应使用此操作步骤。

在这种情况下， MetroCluster 是一种混合配置：

- 一个 DR 组由光纤连接的 MetroCluster FC 节点组成。

您必须对这些节点执行 MetroCluster FC 恢复步骤。

- 一个 DR 组由 MetroCluster IP 节点组成。

您必须对这些节点执行 MetroCluster IP 恢复步骤。

步骤

按以下顺序执行步骤。

1. 通过按顺序执行以下任务，准备要切回的 FC 节点：
 - a. ["验证端口配置（仅限 MetroCluster FC 配置）"](#)

- b. "配置 FC-SAS 网桥 (仅限 MetroCluster FC 配置) "
- c. "配置 FC 交换机 (仅限 MetroCluster FC 配置) "
- d. "验证存储配置" (仅对 MetroCluster FC 节点上更换的驱动器执行这些步骤)
- e. "启动灾难站点上的设备" (仅对 MetroCluster FC 节点上更换的驱动器执行这些步骤)
- f. "为更换的驱动器分配所有权" (仅对 MetroCluster FC 节点上更换的驱动器执行这些步骤)
- g. 执行中的步骤 "将根聚合的磁盘所有权重新分配给更换用的控制器模块 (MetroCluster FC 配置) ", 直到并包括对 mailbox destroy 命令执行问题描述的步骤。
- h. 销毁根聚合的本地丛 (丛 0) :

```
aggr destroy plex-id
```

- i. 如果根聚合未联机, 请将其联机。

2. 启动 MetroCluster FC 节点。

您必须在两个 MetroCluster FC 节点上执行这些步骤。

- a. 显示启动菜单:

```
boot_ontap 菜单
```

- b. 从启动菜单中, 选择选项 6 * 从备份配置更新闪存 *。
- c. 对以下提示回答 y :

此操作将使用上次备份到磁盘来替换所有基于闪存的配置。确实要继续? : y

系统将启动两次, 第二次加载新配置。



如果未清除已用替代控制器的 NVRAM 内容, 则可能会出现崩溃, 并显示以下消息: panic : NVRAM contents are invalid... 如果发生这种情况, 请重复这些子步骤, 将系统启动到 ONTAP 提示符。然后、您需要 [重置启动恢复和rdb_Corrupt bootargs](#)

3. 在丛 0 上镜像根聚合:

您必须在两个 MetroCluster FC 节点上执行这些步骤。

- a. 将三个 pool0 磁盘分配给新控制器模块。
- b. 镜像根聚合 pool1 丛:

```
aggr mirror root-aggr-name
```

- c. 将无主磁盘分配给本地节点上的 pool0

4. 返回维护模式。

您必须在两个 MetroCluster FC 节点上执行这些步骤。

- a. 暂停节点:

```
halt
```

- b. 将节点启动到维护模式:

```
boot_ontap maint
```

5. 删除邮箱磁盘的内容:

m邮箱销毁本地

您必须在两个 MetroCluster FC 节点上执行这些步骤。

6. halt the nodes : + halt

7. 节点启动后, 验证节点的状态:

```
MetroCluster node show
```

```
siteA::*> metrocluster node show
DR
Group Cluster Node
-----
1      siteA
      wmc66-a1      configured      enabled      waiting for
switchback recovery
      wmc66-a2      configured      enabled      waiting for
switchback recovery
      siteB
      wmc66-b1      configured      enabled      switchover
completed
      wmc66-b2      configured      enabled      switchover
completed
2      siteA
      wmc55-a1      -
      wmc55-a2      unreachable      -
      siteB
      wmc55-b1      configured      enabled      switchover
completed
      wmc55-b2      configured
```

8. 通过执行中的任务, 准备要切回的 MetroCluster IP 节点 "在 MetroCluster IP 配置中准备切回" 最高及包括 "删除运行正常的站点拥有的故障丛 (MetroCluster IP 配置) "。
9. 在 MetroCluster FC 节点上, 执行中的步骤 "执行聚合修复和还原镜像 (MetroCluster FC 配置) "。
10. 在 MetroCluster IP 节点上, 执行中的步骤 "执行聚合修复和还原镜像 (MetroCluster IP 配置) "。
11. 从开始, 继续执行恢复过程的其余任务 "为 FabricPool 配置重新建立对象存储"。

【重置启动恢复】重置boot_recovery和rdb_Corrupt bootargs

如果需要、您可以重置boot_recovery和rdb_Corrupt_bootargs

步骤

1. 将节点暂停回LOADER提示符：

```
siteA::*> halt -node <node-name>
```

2. 检查是否已设置以下bootarg：

```
LOADER> printenv bootarg.init.boot_recovery  
LOADER> printenv bootarg.rdb_corrupt
```

3. 如果已将任一bootarg设置为值、请取消设置并启动ONTAP：

```
LOADER> unsetenv bootarg.init.boot_recovery  
LOADER> unsetenv bootarg.rdb_corrupt  
LOADER> saveenv  
LOADER> bye
```

正在完成恢复

执行所需任务、完成从多控制器或存储故障中恢复的过程。

为 **FabricPool** 配置重新建立对象存储

如果 FabricPool 镜像中的某个对象存储与 MetroCluster 灾难站点位于同一位置并被销毁，则必须重新建立对象存储和 FabricPool 镜像。

关于此任务

- 如果对象存储是远程的，并且 MetroCluster 站点已被销毁，则无需重建对象存储，原始对象存储配置以及冷数据内容会保留下来。
- 有关 FabricPool 配置的详细信息，请参见 ["磁盘和聚合管理"](#)。

步骤

1. 按照中的操作步骤 " 更换 MetroCluster 配置上的 FabricPool 镜像 " 进行操作 ["磁盘和聚合管理"](#)。

验证已更换节点上的许可证

如果受损节点正在使用需要标准（节点锁定）许可证的 ONTAP 功能，则必须为更换节点安装新许可证。对于具有标准许可证的功能，集群中的每个节点都应具有自己的功能密钥。

关于此任务

在安装许可证密钥之前，需要标准许可证的功能仍可供替代节点使用。但是，如果受损节点是集群中唯一具有此功能许可证的节点，则不允许更改此功能的配置。此外，在节点上使用未经许可的功能可能会使您不符合您的许可协议，因此您应尽快在替代节点上安装替代许可证密钥。

许可证密钥必须采用 28 个字符的格式。

您有 90 天的宽限期来安装许可证密钥。宽限期过后，所有旧许可证将失效。安装有效的许可证密钥后，您可以在 24 小时内安装所有密钥，直到宽限期结束。



如果已更换站点上的所有节点（对于双节点 MetroCluster 配置，为单个节点），则必须在切回之前在替代节点上安装许可证密钥。

步骤

- 1. 确定节点上的许可证：

许可证显示

以下示例显示了有关系统中许可证的信息：

```
cluster_B::> license show
               (system license show)

Serial Number: 1-80-00050
Owner: site1-01
Package      Type      Description      Expiration
-----
Base         license   Cluster Base License   -
NFS          site      NFS License        -
CIFS         site      CIFS License        -
iSCSI        site      iSCSI License        -
FCP          site      FCP License         -
FlexClone    site      FlexClone License     -

6 entries were displayed.
```

- 2. 在切回后验证这些许可证是否适用于节点：

MetroCluster check license show

以下示例显示了适用于此节点的许可证：

```
cluster_B::> metrocluster check license show
```

Cluster	Check	Result
-----	-----	-----
Cluster_B	negotiated-switchover-ready	not-applicable
NFS	switchback-ready	not-applicable
CIFS	job-schedules	ok
iSCSI	licenses	ok
FCP	periodic-check-enabled	ok

- 如果您需要新的许可证密钥，请在 NetApp 支持站点的软件许可证下的我的支持部分中获取替代许可证密钥。



系统会自动生成所需的新许可证密钥，并将其发送到文件中的电子邮件地址。如果您未能在30天内收到包含许可证密钥的电子邮件，请参阅知识库文章中的["如果我的许可证出现问题、应联系谁？"部分](#) ["主板更换后流程、用于更新AFF/FAS系统上的许可。"](#)

- 安装每个许可证密钥：

```
ssystem license add -license-code license-key , license-key...+
```

- 如果需要，删除旧许可证：

- 检查未使用的许可证：

```
license clean-up -unused -simulate
```

- 如果列表显示正确，请删除未使用的许可证：

```
license clean-up -unused
```

正在还原密钥管理

如果数据卷已加密，则必须还原密钥管理。如果根卷已加密，则必须恢复密钥管理。

步骤

- 如果数据卷已加密，请使用适用于您的密钥管理配置的正确命令还原密钥。

如果您使用的是 ...	使用此命令 ...
• 板载密钥管理 *	<pre>sSecurity key-manager 板载同步</pre> 有关详细信息，请参见 "还原板载密钥管理加密密钥"。
• 外部密钥管理 *	<pre>sSecurity key-manager key query -node node-name</pre> 有关详细信息，请参见 "还原外部密钥管理加密密钥"。

2. 如果根卷已加密，请使用中的操作步骤 ["如果根卷已加密，则恢复密钥管理"](#)。

执行切回

修复 MetroCluster 配置后，您可以执行 MetroCluster 切回操作。MetroCluster 切回操作会将配置恢复到其正常运行状态，灾难站点上的 sync-source Storage Virtual Machine （SVM）处于活动状态，并从本地磁盘池提供数据。

开始之前

- 灾难集群必须已成功切换到正常运行的集群。
- 必须已对数据和根聚合执行修复。
- 正常运行的集群节点不能处于 HA 故障转移状态（对于每个 HA 对，所有节点都必须已启动且正在运行）。
- 灾难站点控制器模块必须完全启动，而不是处于 HA 接管模式。
- 必须镜像根聚合。
- 交换机间链路（ISL）必须处于联机状态。
- 必须在系统上安装所有必需的许可证。

步骤

1. 确认所有节点均处于已启用状态：

MetroCluster node show

以下示例显示了处于已启用状态的节点：

```
cluster_B::> metrocluster node show
```

DR		Configuration	DR
Group	Cluster Node	State	Mirroring Mode
1	cluster_A		
	node_A_1	configured	enabled heal roots completed
	node_A_2	configured	enabled heal roots completed
	cluster_B		
	node_B_1	configured	enabled waiting for
switchback recovery			
	node_B_2	configured	enabled waiting for
switchback recovery			

4 entries were displayed.

2. 确认所有 SVM 上的重新同步均已完成：

MetroCluster SVM show

3. 验证修复操作正在执行的任何自动 LIF 迁移是否已成功完成：

MetroCluster check lif show

4. 从运行正常的集群中的任何节点运行 `MetroCluster switchback` 命令，以执行切回。
5. 检查切回操作的进度：

MetroCluster show

当输出显示 "waiting-for-switchback" 时，切回操作仍在进行中：

```
cluster_B::> metrocluster show
Cluster                               Entry Name                               State
-----
Local: cluster_B                      Configuration state configured
                                      Mode                                     switchover
                                      AUSO Failure Domain -
Remote: cluster_A                     Configuration state configured
                                      Mode                                     waiting-for-switchback
                                      AUSO Failure Domain -
```

当输出显示 "Normal" 时，切回操作完成：

```
cluster_B::> metrocluster show
Cluster                               Entry Name                               State
-----
Local: cluster_B                      Configuration state configured
                                      Mode                                     normal
                                      AUSO Failure Domain -
Remote: cluster_A                     Configuration state configured
                                      Mode                                     normal
                                      AUSO Failure Domain -
```

如果切回需要很长时间才能完成，您可以在高级权限级别使用以下命令来检查正在进行的基线的状态：

MetroCluster config-replication resync-status show

6. 重新建立任何 SnapMirror 或 SnapVault 配置。

在 ONTAP 8.3 中，您需要在执行 MetroCluster 切回操作后手动重新建立丢失的 SnapMirror 配置。在 ONTAP 9.0 及更高版本中，系统会自动重新建立此关系。

验证切回是否成功

执行切回后，您需要确认所有聚合和 Storage Virtual Machine （SVM）均已切回并联机。

步骤

1. 验证切换后的数据聚合是否已切回：

s存储聚合显示

在以下示例中，节点 B2 上的 aggr_b2 已切回：

```
node_B_1::> storage aggregate show
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
...
aggr_b2        227.1GB    227.1GB    0% online      0 node_B_2  raid_dp,
mirrored,
normal

node_A_1::> aggr show
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
...
aggr_b2        -          -          - unknown      - node_A_1
```

如果灾难站点包含未镜像聚合，而未镜像聚合不再存在，则该聚合可能会在 storage aggregate show 命令的输出中显示为 "unknown" 状态。要删除未镜像聚合的过期条目、请联系技术支持、请参阅知识库文章 [如何在存储丢失的灾难发生后删除MetroCluster 中陈旧的未镜像聚合条目。](#)

2. 验证运行正常的集群上的所有同步目标SVM是否均处于休眠状态(显示运行状态为's已')：

vserver show -subtype sync-destination

```
node_B_1::> vserver show -subtype sync-destination
Vserver      Type      Subtype      Admin      Operational  Root
Aggregate
-----
...
cluster_A-vs1a-mc data sync-destination
running      stopped      vs1a_vol    aggr_b2
```

MetroCluster 配置中的 sync-destination 聚合会在其名称中自动附加后缀 "-mc"，以帮助标识它们。

3. 验证灾难集群上的同步源SVM是否已启动且正在运行：

```
vserver show -subtype sync-source
```

```
node_A_1::> vserver show -subtype sync-source
Admin      Operational  Root
Vserver    Type      Subtype    State      State      Volume
Aggregate
-----
.....
...
vs1a       data      sync-source
                                         running   running   vs1a_vol  aggr_b2
```

4. 使用 MetroCluster operation show 命令确认切回操作成功。

如果命令输出显示 ...	那么 ...
切回操作状态为成功。	切回过程已完成，您可以继续操作系统。
切回操作或切回 - 继续 - 代理操作已部分成功。	执行 MetroCluster operation show 命令输出中建议的修复操作。

完成后

您必须重复前面的部分，以反向执行切回。如果 site_A 已切换 site_B ，请让 site_B 切换 site_A

镜像替代节点的根聚合

如果更换了磁盘，则必须镜像灾难站点上新节点的根聚合。

步骤

- 1. 在灾难站点上，确定未镜像的聚合：

s存储聚合显示

```
cluster_A::> storage aggregate show
```

Aggregate Status	Size	Available	Used%	State	#Vols	Nodes	RAID
node_A_1_aggr0	1.49TB	74.12GB	95%	online	1	node_A_1	
raid4,							
normal							
node_A_2_aggr0	1.49TB	74.12GB	95%	online	1	node_A_2	
raid4,							
normal							
node_A_1_aggr1	1.49TB	74.12GB	95%	online	1	node_A_1	raid
4, normal							
mirrored							
node_A_2_aggr1	1.49TB	74.12GB	95%	online	1	node_A_2	raid
4, normal							
mirrored							

4 entries were displayed.

```
cluster_A::>
```

2. 镜像其中一个根聚合:

```
storage aggregate mirror -aggregate root-aggregate
```

以下示例显示了在镜像聚合时命令如何选择磁盘并提示确认。

```

cluster_A::> storage aggregate mirror -aggregate node_A_2_aggr0

Info: Disks would be added to aggregate "node_A_2_aggr0" on node
"node_A_2" in
    the following manner:

    Second Plex

        RAID Group rg0, 3 disks (block checksum, raid4)
        Position    Disk                                Type
Size
-----
-----
-      parity      2.10.0                                SSD
      data         1.11.19                               SSD
894.0GB      data   2.10.2                                SSD
894.0GB

    Aggregate capacity available for volume use would be 1.49TB.

Do you want to continue? {y|n}: y

cluster_A::>

```

3. 验证根聚合的镜像是否已完成：

s存储聚合显示

以下示例显示根聚合已镜像。

```
cluster_A::> storage aggregate show
```

Aggregate Status	Size	Available	Used%	State	#Vols	Nodes	RAID
node_A_1_aggr0	1.49TB	74.12GB	95%	online	1	node_A_1	raid4, mirrored, normal
node_A_2_aggr0	2.24TB	838.5GB	63%	online	1	node_A_2	raid4, mirrored, normal
node_A_1_aggr1	1.49TB	74.12GB	95%	online	1	node_A_1	raid4, mirrored, normal
node_A_2_aggr1	1.49TB	74.12GB	95%	online	1	node_A_2	raid4 mirrored, normal

4 entries were displayed.

```
cluster_A::>
```

4. 对其他根聚合重复上述步骤。

任何状态不为已镜像的根聚合都必须进行镜像。

重新配置 ONTAP 调解器（MetroCluster IP 配置）

如果您具有使用 ONTAP 调解器配置的 MetroCluster IP 配置，则必须删除并重新配置与 ONTAP 调解器的关联。

开始之前

- 您必须拥有 ONTAP Mediator 的 IP 地址、用户名和密码。
- ONTAP Mediator 必须在 Linux 主机上配置并运行。

步骤

1. 删除现有 ONTAP 调解器配置：

```
MetroCluster configuration-settings mediator remove
```

2. 重新配置 ONTAP 调解器配置：

```
MetroCluster configuration-settings mediator add -mediate-address mediate-ip-address
```

验证 **MetroCluster** 配置的运行状况

您应检查 MetroCluster 配置的运行状况以验证是否正常运行。

步骤

1. 检查每个集群上是否已配置 MetroCluster 并处于正常模式：

```
MetroCluster show
```

```
cluster_A::> metrocluster show
Cluster                               Entry Name                               State
-----
Local: cluster_A                      Configuration state configured
                                         Mode normal
                                         AUSO Failure Domain auso-on-cluster-disaster
Remote: cluster_B                     Configuration state configured
                                         Mode normal
                                         AUSO Failure Domain auso-on-cluster-disaster
```

2. 检查是否已在每个节点上启用镜像：

```
MetroCluster node show
```

```
cluster_A::> metrocluster node show
DR                               Configuration  DR
Group Cluster Node              State          Mirroring Mode
-----
1      cluster_A
        node_A_1      configured    enabled    normal
        cluster_B
        node_B_1      configured    enabled    normal
2 entries were displayed.
```

3. 检查 MetroCluster 组件是否运行正常：

```
MetroCluster check run
```

```
cluster_A::> metrocluster check run
```

Last Checked On: 10/1/2014 16:03:37

Component	Result
nodes	ok
lifs	ok
config-replication	ok
aggregates	ok

4 entries were displayed.

Command completed. Use the `metrocluster check show -instance` command or sub-commands in `metrocluster check` directory for detailed results. To check if the nodes are ready to do a switchover or switchback operation, run `metrocluster switchover -simulate` or `metrocluster switchback -simulate`, respectively.

4. 检查是否没有运行状况警报:

s系统运行状况警报显示

5. 模拟切换操作:

- 在任何节点的提示符处, 更改为高级权限级别:

```
set -privilege advanced
```

当系统提示您继续进入高级模式并显示高级模式提示符 (*>) 时, 您需要使用 y 进行响应。

- 使用 `-simulate` 参数执行切换操作:

```
MetroCluster switchover -simulate
```

- 返回到管理权限级别:

```
set -privilege admin
```

6. 对于使用 ONTAP 调解器的 MetroCluster IP 配置, 请确认 ONTAP 调解器已启动并正在运行。

- 检查调解器磁盘是否对系统可见:

```
storage failover mailbox-disk show
```

以下示例显示已识别邮箱磁盘。

```
node_A_1::*> storage failover mailbox-disk show
```

Mailbox				
Node	Owner	Disk	Name	Disk UUID
-----	-----	-----	-----	-----
still13-vsim-ucs626g				
.				
.				
local	0m.i2.3L26			
7BBA77C9:AD702D14:831B3E7E:0B0730EE:00000000:00000000:00000000:00000000				
local	0m.i2.3L27			
928F79AE:631EA9F9:4DCB5DE6:3402AC48:00000000:00000000:00000000:00000000				
local	0m.i1.0L60			
B7BCDB3C:297A4459:318C2748:181565A3:00000000:00000000:00000000:00000000				
.				
.				
.				
partner	0m.i1.0L14			
EA71F260:D4DD5F22:E3422387:61D475B2:00000000:00000000:00000000:00000000				
partner	0m.i2.3L64			
4460F436:AAE5AB9E:D1ED414E:ABF811F7:00000000:00000000:00000000:00000000				
28 entries were displayed.				

b. 更改为高级权限级别:

```
set -privilege advanced
```

c. 检查邮箱 LUN 是否对系统可见:

```
storage iscsi-initiator show
```

输出将显示存在邮箱 LUN :

```

Node      Type      Label      Target Portal      Target Name
Admin/Op
-----
.
.
.
.node_A_1
      mailbox
      mediator 172.16.254.1      iqn.2012-
05.local:mailbox.target.db5f02d6-e3d3      up/up
.
.
.
17 entries were displayed.

```

a. 返回到管理权限级别：

```
set -privilege admin
```

从非控制器故障中恢复

在灾难站点上的设备进行了任何必要的维护或更换，但未更换任何控制器后，您可以开始将 MetroCluster 配置恢复为完全冗余状态的过程。其中包括修复配置（首先修复数据聚合，然后修复根聚合）以及执行切回操作。

开始之前

- 灾难集群中的所有 MetroCluster 硬件都必须正常运行。
- 整个 MetroCluster 配置必须处于切换状态。
- 在光纤连接的 MetroCluster 配置中，ISL 必须在 MetroCluster 站点之间正常运行。

启用控制台日志记录

NetApp强烈建议您在使用的设备上启用控制台日志记录、并在执行此过程时执行以下操作：

- 在维护期间保持AutoSupport处于启用状态。
- 在维护前后触发维护AutoSupport消息、以便在维护活动期间禁用案例创建。

请参阅知识库文章 ["如何在计划的维护时段禁止自动创建案例"](#)。

- 为任何命令行界面会话启用会话日志记录。有关如何启用会话日志记录的说明，请查看知识库文章中的“日志记录会话输出”部分 ["如何配置PuTTY以优化与ONTAP系统的连接"](#)。

修复MetroCluster配置中的配置

在MetroCluster FC配置中、您可以按特定顺序执行修复操作、以便在切换后还原MetroCluster功能。

在MetroCluster IP配置中、修复操作应在切换后自动启动。否则、您可以手动执行修复操作。

开始之前

- 必须已执行切换，并且正常运行的站点必须正在提供数据。
- 灾难站点上的节点必须暂停或保持关闭状态。

在修复过程中，不能完全启动它们。

- 灾难站点上的存储必须可访问（磁盘架已启动，正常运行且可访问）。
- 在光纤连接的 MetroCluster 配置中，交换机间链路（ISL）必须已启动且正在运行。
- 在四节点 MetroCluster 配置中，运行正常的站点中的节点不能处于 HA 故障转移状态（对于每个 HA 对，所有节点都必须已启动且正在运行）。

关于此任务

必须先对数据聚合执行修复操作，然后再对根聚合执行此操作。

修复数据聚合

修复并更换灾难站点上的任何硬件后，您必须修复数据聚合。此过程会重新同步数据聚合，并使（现已修复）灾难站点做好正常运行的准备。在修复根聚合之前，您必须修复数据聚合。

关于此任务

以下示例显示了强制切换，在此可以使已切换的聚合联机。远程集群中的所有配置更新均已成功复制到本地集群。您在此操作步骤中启动灾难站点上的存储，但不会也不能启动灾难站点上的控制器模块。

步骤

1. 验证切换是否已完成：

MetroCluster 操作显示

```
controller_A_1:> metrocluster operation show
Operation: switchover
State: successful
Start Time: 7/25/2014 20:01:48
End Time: 7/25/2014 20:02:14
Errors: -
```

2. 在运行正常的集群中运行以下命令，以重新同步数据聚合：

```
MetroCluster heal -phase aggregates
```

```
controller_A_1::> metrocluster heal -phase aggregates
[Job 130] Job succeeded: Heal Aggregates is successful.
```

如果修复被否决，您可以使用`-override-vetoes`参数重新发出 MetroCluster heal 命令。如果使用此可选参数，则系统将覆盖任何阻止修复操作的软否决。

3. 验证操作是否已完成：

MetroCluster 操作显示

```
controller_A_1::> metrocluster operation show
Operation: heal-aggregates
State: successful
Start Time: 7/25/2014 18:45:55
End Time: 7/25/2014 18:45:56
Errors: -
```

4. 检查聚合的状态：

storage aggregate show 命令。

```
controller_A_1::> storage aggregate show
Aggregate Size      Available Used% State   #Vols  Nodes      RAID
Status
-----
...
aggr_b2    227.1GB  227.1GB   0%   online  0      mcc1-a2    raid_dp,
mirrored, normal...
```

5. 如果已在灾难站点上更换存储，您可能需要重新镜像聚合。

发生灾难后修复根聚合

修复数据聚合之后，您必须修复根聚合，以便为切回操作做好准备。

开始之前

MetroCluster 修复过程的数据聚合阶段必须已成功完成。

步骤

1. 切回镜像聚合：

```
MetroCluster heal -phase root-aggregates
```

```
mcclA::> metrocluster heal -phase root-aggregates
[Job 137] Job succeeded: Heal Root Aggregates is successful
```

如果修复被否决，您可以使用 `-override-vetoes` 参数重新发出 MetroCluster heal 命令。如果使用此可选参数，则系统将覆盖任何阻止修复操作的软否决。

2. 在目标集群上运行以下命令，以确保修复操作已完成：

MetroCluster 操作显示

```
mcclA::> metrocluster operation show
Operation: heal-root-aggregates
State: successful
Start Time: 7/29/2014 20:54:41
End Time: 7/29/2014 20:54:42
Errors: -
```

验证您的系统是否已做好切回准备

如果您的系统已处于切换状态，您可以使用 `-simulate` 选项预览切回操作的结果。

步骤

1. 启动灾难站点上的每个控制器模块。

如果节点已关闭：

启动节点。

如果节点位于加载程序提示符处：

运行命令： `boot_ontap`

2. 节点启动完成后、验证根聚合是否已镜像。

如果 **plex** 发生故障：

- a. 摧毁失败的 plex：

```
storage aggregate plex delete -aggregate <aggregate_name> -plex  
<plex_name>
```

- b. 通过重新创建镜像来重新建立镜像关系：

```
storage aggregate mirror -aggregate <aggregate-name>
```

如果丛处于脱机状态：

在线丛：

```
storage aggregate plex online -aggregate <aggregate_name> -plex <plex_name>
```

如果两个丛都存在：

重新同步将自动启动。

3. 模拟切回操作：

- a. 在任一正常运行的节点的提示符处，更改为高级权限级别：

```
set -privilege advanced
```

当系统提示您继续进入高级模式并显示高级模式提示符（*>）时，您需要使用 y 进行响应。

- a. 使用 `-simulate` 参数执行切回操作：

```
MetroCluster switchback -simulate
```

- b. 返回到管理权限级别：

```
set -privilege admin
```

4. 查看返回的输出。

输出将显示切回操作是否会出错。

验证结果示例

以下示例显示了对切回操作的成功验证：

```

cluster4::*> metrocluster switchback -simulate
(metrocluster switchback)
[Job 130] Setting up the nodes and cluster components for the switchback
operation...DBG:backup_api.c:327:backup_nso_sb_vetochek : MetroCluster
Switch Back
[Job 130] Job succeeded: Switchback simulation is successful.

cluster4::*> metrocluster op show
(metrocluster operation show)
Operation: switchback-simulate
State: successful
Start Time: 5/15/2014 16:14:34
End Time: 5/15/2014 16:15:04
Errors: -

cluster4::*> job show -name Me*

```

Job ID	Name	Owning Vserver	Node	State
130	MetroCluster Switchback	cluster4	cluster4-01	Success

```

Description: MetroCluster Switchback Job - Simulation

```

执行切回

修复 MetroCluster 配置后，您可以执行 MetroCluster 切回操作。MetroCluster 切回操作会将配置恢复到其正常运行状态，灾难站点上的 sync-source Storage Virtual Machine （SVM）处于活动状态，并从本地磁盘池提供数据。

开始之前

- 灾难集群必须已成功切换到正常运行的集群。
- 必须已对数据和根聚合执行修复。
- 正常运行的集群节点不能处于 HA 故障转移状态（对于每个 HA 对，所有节点都必须已启动且正在运行）。
- 灾难站点控制器模块必须完全启动，而不是处于 HA 接管模式。
- 必须镜像根聚合。
- 交换机间链路（ISL）必须处于联机状态。
- 必须在系统上安装所有必需的许可证。

步骤

1. 确认所有节点均处于已启用状态：

MetroCluster node show

以下示例显示了处于 "enabled" 状态的节点：

```
cluster_B::> metrocluster node show
```

DR	Configuration	DR	
Group	Cluster Node	State	Mirroring Mode
1	cluster_A		
	node_A_1	configured	enabled heal roots completed
	node_A_2	configured	enabled heal roots completed
	cluster_B		
	node_B_1	configured	enabled waiting for
	switchback recovery		
	node_B_2	configured	enabled waiting for
	switchback recovery		

4 entries were displayed.

2. 确认所有 SVM 上的重新同步均已完成：

MetroCluster SVM show

3. 验证修复操作正在执行的任何自动 LIF 迁移是否已成功完成：

MetroCluster check lif show

4. 从运行正常的集群中的任何节点运行以下命令，以执行切回。

MetroCluster 切回

5. 检查切回操作的进度：

MetroCluster show

当输出显示 "waiting-for-switchback" 时，切回操作仍在进行中：

```
cluster_B::> metrocluster show
```

Cluster	Entry Name	State
Local: cluster_B	Configuration state	configured
	Mode	switchover
	AUSO Failure Domain	-
Remote: cluster_A	Configuration state	configured
	Mode	waiting-for-switchback
	AUSO Failure Domain	-

当输出显示 "Normal" 时，切回操作完成：

```
cluster_B::> metrocluster show
Cluster                               Entry Name                               State
-----
Local: cluster_B                      Configuration state configured
                                         Mode normal
                                         AUSO Failure Domain -
Remote: cluster_A                     Configuration state configured
                                         Mode normal
                                         AUSO Failure Domain -
```

如果切回需要很长时间才能完成，您可以在高级权限级别使用以下命令来检查正在进行的基线的状态。

```
MetroCluster config-replication resync-status show
```

6. 重新建立任何 SnapMirror 或 SnapVault 配置。

在 ONTAP 8.3 中，您需要在执行 MetroCluster 切回操作后手动重新建立丢失的 SnapMirror 配置。在 ONTAP 9.0 及更高版本中，系统会自动重新建立此关系。

验证切回是否成功

执行切回后，您需要确认所有聚合和 Storage Virtual Machine （ SVM ）均已切回并联机。

步骤

1. 验证切换后的数据聚合是否已切回：

s 存储聚合显示

在以下示例中，节点 B2 上的 aggr_b2 已切回：

```
node_B_1::> storage aggregate show
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
...
aggr_b2        227.1GB    227.1GB    0% online      0 node_B_2  raid_dp,
mirrored,
normal

node_A_1::> aggr show
Aggregate      Size Available Used% State   #Vols  Nodes      RAID
Status
-----
...
aggr_b2        -          -          - unknown      - node_A_1
```

如果灾难站点包含未镜像聚合且未镜像聚合不再存在，则聚合可能会在 `storage aggregate show` 命令的输出中显示为 "unknown" 状态。请联系技术支持以删除未镜像聚合的过期条目、并参考知识库文章 ["如何在存储丢失的灾难发生后删除MetroCluster 中陈旧的未镜像聚合条目。"](#)

2. 验证运行正常的集群上的所有同步目标SVM是否均处于休眠状态(显示运行状态为's已'):

```
vserver show -subtype sync-destination
```

```
node_B_1::> vserver show -subtype sync-destination
Admin      Operational  Root
Vserver    Type      Subtype    State      State      Volume
Aggregate
-----
...
cluster_A-vs1a-mc data sync-destination
running      stopped    vs1a_vol    aggr_b2
```

MetroCluster 配置中的 sync-destination 聚合会在其名称中自动附加后缀 " -mc "，以帮助标识它们。

3. 验证灾难集群上的同步源SVM是否已启动且正在运行:

```
vserver show -subtype sync-source
```

```
node_A_1::> vsriver show -subtype sync-source

Vserver      Type      Subtype      Admin      Operational      Root
Aggregate
-----
...
vs1a          data      sync-source
running      running      vs1a_vol     aggr_b2
```

4. 确认切回操作成功:

MetroCluster 操作显示

如果命令输出显示 ...	那么 ...
切回操作状态为成功。	切回过程已完成，您可以继续操作系统。
切回操作或 <code>sswitchback-continuation-agent</code> 操作已部分成功。	执行 <code>MetroCluster operation show</code> 命令输出中建议的修复操作。

完成后

您必须重复前面的部分，以反向执行切回。如果 `site_A` 已切换 `site_B`，请让 `site_B` 切换 `site_A`

在切回后删除陈旧的聚合列表

在某些情况下，切回后，您可能会注意到存在 *stal* 聚合。陈旧的聚合是指已从 ONTAP 中删除但其信息仍记录在磁盘上的聚合。陈旧的聚合会使用 `nodeshell aggr status -r` 命令显示，但不会使用 `storage aggregate show` 命令显示。您可以删除这些记录，使其不再显示。

关于此任务

如果在 MetroCluster 配置处于切换状态时重新定位了聚合，则可能会发生陈旧的聚合。例如：

- 1. 站点 A 切换到站点 B
- 2. 删除聚合的镜像并将聚合从 `node_B_1` 重新定位到 `node_B_2` 以实现负载平衡。
- 3. 您可以执行聚合修复。

此时，`node_B_1` 上会显示一个陈旧的聚合，即使已从该节点中删除实际聚合也是如此。此聚合显示在 `nodeshell aggr status -r` 命令的输出中。它不会显示在 `storage aggregate show` 命令的输出中。

- 1. 比较以下命令的输出：

s存储聚合显示

```
run local aggr status -r
```

陈旧的聚合显示在 `run local aggr status -r` 输出中，但不显示在 `storage aggregate show` 输出中。例如，以下聚合可能显示在 `run local aggr status -r` 输出中：

```
Aggregate aggr05 (failed, raid_dp, partial) (block checksums)
Plex /aggr05/plex0 (offline, failed, inactive)
  RAID group /myaggr/plex0/rg0 (partial, block checksums)

  RAID Disk Device  HA  SHELF BAY CHAN Pool Type  RPM  Used (MB/blks)
Phys (MB/blks)
  -----
  -----
  dparity          FAILED          N/A          82/ -
  parity           0b.5      0b    -    -    SA:A    0 VMDISK  N/A 82/169472
88/182040
  data             FAILED          N/A          82/ -
  data             FAILED          N/A          82/ -
  data             FAILED          N/A          82/ -
  data             FAILED          N/A          82/ -
  data             FAILED          N/A          82/ -
  data             FAILED          N/A          82/ -
  Raid group is missing 7 disks.
```

2. 删除陈旧的聚合：

- a. 在任一节点的提示符处，更改为高级权限级别：

```
set -privilege advanced
```

当系统提示您继续进入高级模式并显示高级模式提示符（*>）时，您需要使用 `y` 进行响应。

- a. 删除陈旧的聚合：

```
aggregate remove-stale-record -aggregate aggregate_name
```

- b. 返回到管理权限级别：

```
set -privilege admin
```

3. 确认已删除陈旧的聚合记录：

```
run local aggr status -r
```

版权信息

版权所有 © 2026 NetApp, Inc.。保留所有权利。中国印刷。未经版权所有者事先书面许可，本文档中受版权保护的任何部分不得以任何形式或通过任何手段（图片、电子或机械方式，包括影印、录音、录像或存储在电子检索系统中）进行复制。

从受版权保护的 NetApp 资料派生的软件受以下许可和免责声明的约束：

本软件由 NetApp 按“原样”提供，不含任何明示或暗示担保，包括但不限于适销性以及针对特定用途的适用性的隐含担保，特此声明不承担任何责任。在任何情况下，对于因使用本软件而以任何方式造成的任何直接性、间接性、偶然性、特殊性、惩罚性或后果性损失（包括但不限于购买替代商品或服务；使用、数据或利润方面的损失；或者业务中断），无论原因如何以及基于何种责任理论，无论出于合同、严格责任或侵权行为（包括疏忽或其他行为），NetApp 均不承担责任，即使已被告知存在上述损失的可能性。

NetApp 保留在不另行通知的情况下随时对本文档所述的任何产品进行更改的权利。除非 NetApp 以书面形式明确同意，否则 NetApp 不承担因使用本文档所述产品而产生的任何责任或义务。使用或购买本产品不表示获得 NetApp 的任何专利权、商标权或任何其他知识产权许可。

本手册中描述的产品可能受一项或多项美国专利、外国专利或正在申请的专利的保护。

有限权利说明：政府使用、复制或公开本文档受 DFARS 252.227-7013（2014 年 2 月）和 FAR 52.227-19（2007 年 12 月）中“技术数据权利 — 非商用”条款第 (b)(3) 条规定的限制条件的约束。

本文档中所含数据与商业产品和/或商业服务（定义见 FAR 2.101）相关，属于 NetApp, Inc. 的专有信息。根据本协议提供的所有 NetApp 技术数据和计算机软件具有商业性质，并完全由私人出资开发。美国政府对这些数据的使用权具有非排他性、全球性、受限且不可撤销的许可，该许可既不可转让，也不可再许可，但仅限在与交付数据所依据的美国政府合同有关且受合同支持的情况下使用。除本文档规定的情形外，未经 NetApp, Inc. 事先书面批准，不得使用、披露、复制、修改、操作或显示这些数据。美国政府对国防部的授权仅限于 DFARS 的第 252.227-7015(b)（2014 年 2 月）条款中明确的权利。

商标信息

NetApp、NetApp 标识和 <http://www.netapp.com/TM> 上所列的商标是 NetApp, Inc. 的商标。其他公司和产品名称可能是其各自所有者的商标。