

Oracle® Solaris Cluster 4.3 系统管理指南



文件号码 E62256
2016 年 7 月

文件号码 E62256

版权所有 © 2000, 2016, Oracle 和/或其附属公司。保留所有权利。

本软件和相关文档是根据许可证协议提供的，该许可证协议中规定了关于使用和公开本软件和相关文档的各种限制，并受知识产权法的保护。除非在许可证协议中明确许可或适用法律明确授权，否则不得以任何形式、任何方式使用、拷贝、复制、翻译、广播、修改、授权、传播、分发、展示、执行、发布或显示本软件和相关文档的任何部分。除非法律要求实现互操作，否则严禁对本软件进行逆向工程设计、反汇编或反编译。

此文档所含信息可能随时被修改，恕不另行通知，我们不保证该信息没有错误。如果贵方发现任何问题，请书面通知我们。

如果将本软件或相关文档交付给美国政府，或者交付给以美国政府名义获得许可证的任何机构，则适用以下注意事项：

U.S. GOVERNMENT END USERS: Oracle programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, delivered to U.S. Government end users are "commercial computer software" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, use, duplication, disclosure, modification, and adaptation of the programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, shall be subject to license terms and license restrictions applicable to the programs. No other rights are granted to the U.S. Government.

本软件或硬件是为了在各种信息管理应用领域内的一般使用而开发的。它不应被应用于任何存在危险或潜在危险的应用领域，也不是为此而开发的，其中包括可能会产生人身伤害的应用领域。如果在危险应用领域内使用本软件或硬件，贵方应负责采取所有适当的防范措施，包括备份、冗余和其它确保安全使用本软件或硬件的措施。对于因在危险应用领域内使用本软件或硬件所造成的一切损失或损害，Oracle Corporation 及其附属公司概不负责。

Oracle 和 Java 是 Oracle 和/或其附属公司的注册商标。其他名称可能是各自所有者的商标。

Intel 和 Intel Xeon 是 Intel Corporation 的商标或注册商标。所有 SPARC 商标均是 SPARC International, Inc 的商标或注册商标，并应按照许可证的规定使用。AMD、Opteron、AMD 徽标以及 AMD Opteron 徽标是 Advanced Micro Devices 的商标或注册商标。UNIX 是 The Open Group 的注册商标。

本软件或硬件以及文档可能提供了访问第三方内容、产品和服务的方式或有关这些内容、产品和服务的信息。除非您与 Oracle 签订的相应协议另行规定，否则对于第三方内容、产品和服务，Oracle Corporation 及其附属公司明确表示不承担任何种类的保证，亦不对其承担任何责任。除非您和 Oracle 签订的相应协议另行规定，否则对于因访问或使用第三方内容、产品或服务所造成的任何损失、成本或损害，Oracle Corporation 及其附属公司概不负责。

文档可访问性

有关 Oracle 对可访问性的承诺，请访问 Oracle Accessibility Program 网站 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=docacc>。

获得 Oracle 支持

购买了支持服务的 Oracle 客户可通过 My Oracle Support 获得电子支持。有关信息，请访问 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=info>；如果您听力受损，请访问 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=trs>。

目录

使用本文档	21
1 Oracle Solaris Cluster 管理介绍	23
Oracle Solaris Cluster 管理概述	23
使用区域群集	24
Oracle Solaris OS 功能限制	25
管理工具	25
Oracle Solaris Cluster Manager 浏览器界面	26
命令行界面	26
群集管理准备	28
记录 Oracle Solaris Cluster 的硬件配置	28
使用管理控制台	28
备份群集	29
管理群集	29
远程登录到群集	30
如何安全地连接到群集控制台	31
▼ 如何访问群集配置实用程序	31
▼ 如何显示 Oracle Solaris Cluster 发行版本信息和版本信息	32
▼ 如何显示已配置的资源类型、资源组和资源	34
▼ 如何检查群集组件的状态	36
▼ 如何检查公共网络的状态	39
▼ 如何查看群集配置	39
▼ 如何验证基本群集配置	49
▼ 如何检查全局挂载点	54
▼ 如何查看 Oracle Solaris Cluster 命令日志的内容	55
2 Oracle Solaris Cluster 和 RBAC	59
通过 Oracle Solaris Cluster 设置和使用 RBAC	59
Oracle Solaris Cluster RBAC 权限配置文件	59
使用 Oracle Solaris Cluster 管理权限配置文件创建和分配 RBAC 角色	61

▼ 如何从命令行创建角色	61
修改用户的 RBAC 属性	62
▼ 如何使用用户帐户工具修改用户的 RBAC 属性	62
▼ 如何从命令行修改用户的 RBAC 属性	63
3 关闭和引导群集	65
关闭和引导群集概述	65
▼ 如何关闭群集	66
▼ 如何引导群集	68
▼ 如何重新引导群集	72
关闭和引导群集中的单个节点	79
▼ 如何关闭节点	80
▼ 如何引导节点	82
▼ 如何重新引导节点	86
▼ 如何以非群集模式引导节点	89
修复已满的 /var 文件系统	91
▼ 如何修复已满的 /var 文件系统	91
4 数据复制方法	93
了解数据复制	93
支持的数据复制方法	94
在校园群集内使用基于存储的数据复制	95
在校园群集内使用基于存储的数据复制的要求和限制	97
在校园群集内使用基于存储的数据复制时的手动恢复注意事项	98
使用基于存储的数据复制时的最佳做法	98
5 管理全局设备、磁盘路径监视和群集文件系统	99
管理全局设备和全局名称空间概述	99
Solaris Volume Manager 的全局设备许可	100
动态重新配置全局设备	100
配置和管理基于存储的复制设备	100
管理 EMC Symmetrix Remote Data Facility 复制设备	101
管理群集文件系统概述	112
群集文件系统限制	112
管理设备组	112
▼ 如何更新全局设备名称空间	114
▼ 如何更改用于全局设备名称空间的 lofi 设备的大小	114
迁移全局设备名称空间	115

▼ 如何将全局设备名称空间从专用分区迁移到 lofi 设备	116
▼ 如何将全局设备名称空间从 lofi 设备迁移到专用分区	117
添加并注册设备组	118
▼ 如何添加和注册设备组 (Solaris Volume Manager)	118
▼ 如何添加并注册设备组 (原始磁盘)	120
▼ 如何添加并注册复制设备组 (ZFS)	121
▼ 如何配置无 HAStoragePlus 的本地 ZFS 存储池	122
维护设备组	124
如何删除和取消注册设备组 (Solaris Volume Manager)	124
▼ 如何将节点从所有设备组中删除	124
▼ 如何将节点从设备组中删除 (Solaris Volume Manager)	125
▼ 如何从原始磁盘设备组删除节点	127
▼ 如何更改设备组属性	128
▼ 如何设置设备组所需的辅助节点数	130
▼ 如何列出设备组配置	132
▼ 如何切换设备组的主节点	133
▼ 如何将设备组置于维护状态	134
管理存储设备的 SCSI 协议设置	135
▼ 如何显示所有存储设备的默认全局 SCSI 协议设置	136
▼ 如何显示单个存储设备的 SCSI 协议	136
▼ 如何更改所有存储设备的默认全局隔离协议设置	137
▼ 如何更改单个存储设备的隔离协议	138
管理群集文件系统	140
▼ 如何添加群集文件系统	140
▼ 如何删除群集文件系统	142
▼ 如何检查群集中的全局挂载点	144
管理磁盘路径监视	145
▼ 如何监视磁盘路径	146
▼ 如何取消监视磁盘路径	147
▼ 如何打印故障磁盘路径	148
▼ 如何解决磁盘路径状态错误	148
▼ 如何从文件监视磁盘路径	149
▼ 如何启用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能	150
▼ 如何禁用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能	151
6 管理法定	153
管理法定设备	153

动态重新配置法定设备	154
添加法定设备	155
删除或替换法定设备	160
维护法定设备	163
更改法定设备的默认超时时间	169
管理 Oracle Solaris Cluster 法定服务器	169
启动和停止法定服务器软件	170
▼ 如何启动法定服务器	170
▼ 如何停止法定服务器	170
显示有关法定服务器的信息	171
清除过时的法定服务器群集信息	172
7 管理群集互连和公共网络	175
管理群集互连	175
动态重新配置群集互连	176
▼ 如何检查群集互连的状态	177
▼ 如何添加群集传输电缆、传输适配器或传输交换机	178
▼ 如何删除传输电缆、传输适配器或传输交换机	180
▼ 如何启用群集传输电缆	182
▼ 如何禁用群集传输电缆	183
▼ 如何确定传输适配器的实例编号	184
▼ 如何更改现有群集的专用网络地址或地址范围	185
群集互连故障排除	187
管理公共网络	188
如何在群集中管理 IP 网络多路径组	188
动态重新配置公共网络接口	189
8 管理群集节点	191
向群集或区域群集添加节点	191
▼ 如何向现有的群集或区域群集添加节点	192
恢复群集节点	194
▼ 如何从统一归档文件恢复节点	194
从群集中删除节点	198
▼ 如何从区域群集中删除节点	198
▼ 如何从群集软件配置中删除节点	199
▼ 如何在节点连接多于两个的群集中删除阵列与单个节点之间的连接 ..	202
▼ 如何纠正错误消息	204
9 管理群集	205

管理群集概述	205
▼ 如何更改群集名称	206
▼ 如何将节点 ID 映射到节点名称	207
▼ 如何使用对新群集节点的验证	208
▼ 如何在群集中重置时间	209
▼ SPARC: 如何在节点上显示 OpenBoot PROM (OBP)	211
▼ 如何更改节点专用主机名	212
▼ 如何重命名节点	214
▼ 如何更改现有 Oracle Solaris Cluster 逻辑主机名资源使用的逻辑主机名	216
▼ 如何使节点进入维护状态	216
▼ 如何使节点脱离维护状态	218
▼ 如何从群集节点卸载 Oracle Solaris Cluster 软件	220
对节点卸载进行故障排除	222
创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB	223
配置负载限制	229
执行区域群集管理任务	231
▼ 如何从统一归档文件配置区域群集	232
▼ 如何从统一归档文件安装区域群集	233
▼ 如何向区域群集添加网络地址	234
▼ 如何删除区域群集	235
▼ 如何从区域群集中删除文件系统	236
▼ 如何从区域群集中删除存储设备	239
用于测试的故障排除过程	241
在全局群集外部运行应用程序	241
恢复损坏的磁盘集	243
10 对 CPU 使用控制的配置	247
CPU 控制介绍	247
选择方案	247
公平份额调度器	248
配置 CPU 控制	248
▼ 如何在全局群集节点中控制 CPU 使用情况	248
11 更新您的软件	251
Oracle Solaris Cluster 软件更新概述	251
更新 Oracle Solaris Cluster 软件	252
更新技巧	253
更新特定软件包	253

更新或修补区域群集	254
更新法定服务器或 AI 安装服务器	256
卸载软件包	257
▼ 如何卸载软件包	257
▼ 如何卸载法定服务器或 AI 安装服务器软件包	257
12 备份和恢复群集	259
备份群集	259
▼ 如何为镜像执行联机备份 (Solaris Volume Manager)	259
▼ 如何备份群集配置	261
恢复群集文件	261
▼ 如何恢复 ZFS 根 (/) 文件系统 (Solaris Volume Manager)	262
13 使用 Oracle Solaris Cluster Manager 浏览器界面	265
Oracle Solaris Cluster Manager 概述	265
访问 Oracle Solaris Cluster Manager 软件	266
▼ 如何访问 Oracle Solaris Cluster Manager	266
排除 Oracle Solaris Cluster Manager 的故障	267
为 Oracle Solaris Cluster Manager 配置辅助功能支持	269
使用拓扑监视群集	269
▼ 如何使用拓扑监视和更新群集	269
A 部署示例：使用 Availability Suite 软件在群集间配置基于主机的数据复制	271
理解 Availability Suite 软件在群集中的应用	271
Availability Suite 软件使用的数据复制方法	272
远程镜像复制	272
时间点快照	273
复制示例配置	273
在群集间配置基于主机的数据复制的准则	274
配置复制资源组	275
配置应用程序资源组	275
管理接管的准则	278
任务列表：数据复制配置示例	280
连接和安装群集	280
如何配置设备组和资源组的示例	282
▼ 如何在主群集上配置设备组	283
▼ 如何在辅助群集上配置设备组	284
▼ 如何在主群集上为 NFS 应用程序配置文件系统	285
▼ 如何在辅助群集上为 NFS 应用程序配置文件系统	286

▼ 如何在主群集上创建复制资源组	287
▼ 如何在辅助群集上创建复制资源组	289
▼ 如何在主群集上创建 NFS 应用程序资源组	291
▼ 如何在辅助群集上创建 NFS 应用程序资源组	293
如何启用数据复制的示例	295
▼ 如何在主群集上启用复制	296
▼ 如何在辅助群集上启用复制	297
如何执行数据复制的示例	298
▼ 如何执行远程镜像复制	299
▼ 如何执行时间点快照	300
▼ 如何检验复制是否已正确配置	301
用以管理接管的方法示例	304
▼ 如何更新 DNS 条目	304
索引	307



图 1	使用基于存储的数据复制的双工作间配置	96
图 2	远程镜像复制	272
图 3	时间点快照	273
图 4	复制示例配置	274
图 5	配置故障转移应用程序中的资源组	277
图 6	配置可伸缩应用程序中的资源组	278
图 7	客户机到群集的 DNS 映射	279
图 8	群集配置示例	281

表

表 1	Oracle Solaris Cluster 服务	25
表 2	Oracle Solaris Cluster 管理工具	29
表 3	任务列表：关闭和引导群集	66
表 4	任务列表：关闭并引导节点	79
表 5	任务列表：动态重新配置磁盘和磁带设备	100
表 6	任务列表：管理 EMC SRDF 基于存储的复制设备	101
表 7	任务列表：管理设备组	113
表 8	任务列表：管理群集文件系统	140
表 9	任务列表：管理磁盘路径监视	145
表 10	任务列表：法定管理	154
表 11	任务列表：动态重新配置法定设备	155
表 12	任务列表：管理群集互连	176
表 13	任务列表：动态重新配置公共网络接口	177
表 14	任务列表：动态重新配置公共网络接口	190
表 15	任务列表：向现有的全局或区域群集添加节点	192
表 16	任务列表：删除节点	198
表 17	任务列表：管理群集	205
表 18	任务列表：创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB	224
表 19	其他区域群集任务	232
表 20	更新 Oracle Solaris Cluster 软件	252
表 21	任务列表：备份群集文件	259
表 22	任务列表：恢复群集文件	262
表 23	任务列表：数据复制配置示例	280
表 24	需要的硬件和软件	281
表 25	配置示例中的组和资源的摘要	282

示例

例 1	显示 Oracle Solaris Cluster 发行版本信息和版本信息	32
例 2	显示已配置的资源类型、资源组和资源	35
例 3	检查群集组件的状态	37
例 4	检查公共网络状态	39
例 5	查看全局群集配置	40
例 6	查看区域群集配置	47
例 7	检查全局群集配置并且所有基本检查均通过	51
例 8	列出交互式验证检查	51
例 9	运行功能验证检查	52
例 10	检查全局群集配置并且某项检查未通过	53
例 11	检查全局挂载点	54
例 12	查看 Oracle Solaris Cluster 命令日志的内容	57
例 13	使用 smrole 命令创建定制操作员角色	61
例 14	关闭区域群集	67
例 15	SPARC: 关闭全局群集	67
例 16	x86: 关闭全局群集	68
例 17	SPARC: 引导全局群集	69
例 18	x86: 引导群集	70
例 19	重新引导区域群集	74
例 20	SPARC: 重新引导全局群集	75
例 21	x86: 重新引导群集	76
例 22	SPARC: 关闭全局群集节点	81
例 23	x86: 关闭全局群集节点	81
例 24	关闭区域群集节点	82
例 25	SPARC: 引导全局群集节点	84
例 26	x86: 引导群集节点	84
例 27	SPARC: 重新引导全局群集节点	87
例 28	重新引导区域群集节点	88
例 29	SPARC: 在非群集模式下引导全局群集节点	90
例 30	创建副本对	106

例 31	检验数据复制设置	106
例 32	显示与所用磁盘相对应的 DID	108
例 33	组合 DID 实例	109
例 34	显示组合的 DID	109
例 35	主站点故障转移后手动恢复 EMC SRDF 数据	111
例 36	更新全局设备名称空间	114
例 37	添加 Solaris Volume Manager 设备组	120
例 38	从设备组中删除一个节点 (Solaris Volume Manager)	126
例 39	从原始设备组中删除节点	127
例 40	更改设备组属性	129
例 41	更改所需的辅助节点数 (Solaris Volume Manager)	131
例 42	将所需的辅助节点数设置为默认值	131
例 43	列出所有设备组的状态	132
例 44	列出特定设备组的配置	133
例 45	切换设备组的主节点	134
例 46	将设备组置于维护状态	135
例 47	显示所有存储设备的默认全局 SCSI 协议设置	136
例 48	显示单个设备的 SCSI 协议	137
例 49	为所有存储设备设置默认全局隔离协议设置	138
例 50	设置单个设备的隔离协议	140
例 51	删除群集文件系统	144
例 52	监视单个节点上的磁盘路径	146
例 53	监视所有节点上的磁盘路径	146
例 54	从 CCR 重新读取磁盘配置	147
例 55	取消监视磁盘路径	147
例 56	打印故障磁盘路径	148
例 57	从文件监视磁盘路径	150
例 58	删除最后一个法定设备	162
例 59	将法定设备置于维护状态	165
例 60	重新设置法定选票计数 (法定设备)	167
例 61	列出法定配置	168
例 62	启动所有已配置的法定服务器	170
例 63	启动特定法定服务器	170
例 64	停止所有已配置的法定服务器	171
例 65	停止特定法定服务器	171
例 66	显示一个法定服务器的配置信息	172
例 67	显示多个法定服务器的配置信息	172
例 68	显示所有正在运行的法定服务器的配置信息	172
例 69	从法定服务器配置中清除过时的群集信息	174

例 70	检查群集互连的状态	177
例 71	检验添加群集传输电缆、传输适配器或传输交换机	179
例 72	检验删除传输电缆、传输适配器或传输交换机	181
例 73	从群集软件配置中删除节点	201
例 74	更改群集的名称	207
例 75	将节点 ID 映射到节点名称	208
例 76	防止将新计算机添加到全局群集中	209
例 77	允许将所有新计算机添加到全局群集中	209
例 78	指定要添加到全局群集中的新计算机	209
例 79	将验证设置为标准 UNIX	209
例 80	将验证设置为 DES	209
例 81	更改专用主机名	214
例 82	将全局群集节点置于维护状态	218
例 83	使群集节点脱离维护状态并重置法定选票计数	219
例 84	从全局群集中删除区域群集	236
例 85	删除区域群集中的高可用性本地文件系统	238
例 86	删除区域群集中的高可用性 ZFS 文件系统	239
例 87	从区域群集中删除 Solaris Volume Manager 磁盘集	240
例 88	从区域群集中删除 DID 设备	240
例 89	恢复 ZFS 根 (/) 文件系统 (Solaris Volume Manager)	263

使用本文档

《Oracle Solaris Cluster 系统管理指南》提供了有关在基于 SPARC 和基于 x86 的系统上管理 Oracle Solaris Cluster 配置的过程。

- 概述 – 介绍如何配置 Oracle Solaris Cluster 配置
- 目标读者 – 技术人员、系统管理员和授权服务提供商
- 必备知识 – 对故障排除和硬件更换具有丰富经验

产品文档库

有关该产品及相关产品的文档和资源，可从以下网址获得：<http://www.oracle.com/pls/topic/lookup?ctx=E62283>。

反馈

可以在 <http://www.oracle.com/goto/docfeedback> 上提供有关本文档的反馈。

Oracle Solaris Cluster 管理介绍

本章提供以下有关管理全局群集和区域群集的信息，并且包括使用 Oracle Solaris Cluster 管理工具的过程：

- “Oracle Solaris Cluster 管理概述” [23]
- “使用区域群集” [24]
- “Oracle Solaris OS 功能限制” [25]
- “管理工具” [25]
- “群集管理准备” [28]
- “管理群集” [29]

本指南中的所有过程均适用于 Oracle Solaris 11 操作系统。

全局群集由一个或多个全局群集节点组成。全局群集还可以包含 solaris 或 solaris10 标记的非全局区域，这些区域不是节点，但是配置有 HA for Zones 数据服务。

区域群集包含通过 cluster 属性设置的一个或多个 solaris、solaris10 或 labeled 标记的非全局区域。在区域群集中，不允许使用任何其他标记类型。labeled 标记的区域群集仅供 Oracle Solaris 软件的 Trusted Extensions 功能使用。可以通过使用 clzonecluster 命令、clsetup 实用程序或 Oracle Solaris Cluster Manager 浏览器界面创建区域群集。

您可以使用 Oracle Solaris Zones 提供的隔离功能，在类似于全局群集的区域群集上运行受支持的服务。区域群集依赖于全局群集，因此需要全局群集。全局群集不包含区域群集。区域群集在一台计算机上至多具有一个区域群集节点。只有当同一计算机上的全局群集节点能够继续正常工作时，区域群集节点才能继续正常工作。如果计算机上的全局群集节点发生故障，则该计算机上的所有区域群集节点也将发生故障。有关区域群集的常规信息，请参见《Oracle Solaris Cluster 4.3 Concepts Guide》。

Oracle Solaris Cluster 管理概述

Oracle Solaris Cluster 高可用性环境能够确保关键应用程序对最终用户可用。系统管理员的职责就是保证 Oracle Solaris Cluster 配置的稳定性和可操作性。

请先熟悉以下手册中的规划信息，然后再开始管理任务。

- [《Oracle Solaris Cluster 4.3 软件安装指南》](#) 中的 第 1 章, “规划 Oracle Solaris Cluster 配置”
- [《Oracle Solaris Cluster 4.3 Concepts Guide》](#)

Oracle Solaris Cluster 管理组织成以下手册中的任务。

注 - 其中的一些任务可以使用 Oracle Solaris Cluster Manager 浏览器界面执行。这在单独的任务过程中进行了说明。有关 Oracle Solaris Cluster Manager 登录说明, 请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

- 标准任务, 用于定期 (甚至每天) 管理和维护全局群集或区域群集。本指南中讲述了这些任务。
- 数据服务任务, 比如安装、配置和更改特性。这些任务在 [《Oracle Solaris Cluster 4.3 数据服务规划和管理指南》](#) 中进行介绍。
- 服务任务, 比如添加或检修存储或网络硬件。这些任务在 [《Oracle Solaris Cluster Hardware Administration Manual》](#) 中进行介绍。

一般情况下, 可以在群集运行期间执行 Oracle Solaris Cluster 管理任务。如果需从群集中去掉某个节点乃至关闭该节点, 您可以在其余节点继续执行群集操作的同时执行该操作。除非另有说明, 否则 Oracle Solaris Cluster 管理任务应在全局群集节点中执行。对于那些要求关闭整个群集的操作过程, 应通过将停机时间安排在正常工作时间之外来尽量减小对系统的影响。如果打算关闭群集或某个群集节点, 请提前通知用户。

使用区域群集

在区域群集中还可以运行两个 Oracle Solaris Cluster 管理命令 (`cluster` 和 `clnode`)。但是, 这些命令的作用范围仅限于发出命令的区域群集。例如, 在全局群集节点中使用 `cluster` 命令时, 将检索有关全局群集和所有区域群集的所有信息。在区域群集中使用 `cluster` 命令时, 将检索有关该特定区域群集的信息。

当您在全局群集节点中使用 `clzonecluster` 命令时, 该命令将影响全局群集中的所有区域群集。区域群集命令还影响区域群集中的所有节点, 即使这些节点在命令发出时处于关闭状态也是如此。

区域群集支持对处于资源组管理器 (Resource Group Manager, RGM) 控制下的资源进行委托管理。因此, 区域群集管理员可以查看 (但不能更改) 跨区域群集边界的区域群集依赖性。只有全局群集节点中的管理员可以创建、修改或删除跨区域群集边界的依赖性。

以下列表包含对区域群集执行的主要管理任务。

- 启动和重新引导区域群集 – 请参见 [第 3 章 关闭和引导群集](#)。还可以使用 Oracle Solaris Cluster Manager 浏览器界面引导和重新引导区域群集。有关 Oracle Solaris Cluster Manager 登录说明, 请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

- 向区域群集添加节点 – 请参见[第 8 章 管理群集节点](#)。
- 从区域群集删除节点 – 请参见[如何从区域群集中删除节点 \[198\]](#)。还可以使用 Oracle Solaris Cluster Manager 浏览器界面从区域群集节点卸载软件。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。
- 查看区域群集的配置 – 请参见[如何查看群集配置 \[39\]](#)。还可以使用 Oracle Solaris Cluster Manager 浏览器界面查看区域群集的配置。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。
- 验证区域群集的配置 – 请参见[如何验证基本群集配置 \[49\]](#)。
- 停止区域群集 – 请参见[第 3 章 关闭和引导群集](#)。还可以使用 Oracle Solaris Cluster Manager 浏览器界面关闭区域群集。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

Oracle Solaris OS 功能限制

请勿使用服务管理工具 (Service Management Facility, SMF) 管理界面来启用或禁用以下 Oracle Solaris Cluster 服务。

表 1 Oracle Solaris Cluster 服务

Oracle Solaris Cluster 服务	FMRI
pnm	svc:/system/cluster/pnm:default
cl_event	svc:/system/cluster/cl_event:default
cl_eventlog	svc:/system/cluster/cl_eventlog:default
rpc_pmf	svc:/system/cluster/rpc_pmf:default
rpc_fed	svc:/system/cluster/rpc_fed:default
rgm	svc:/system/cluster/rgm:default
scdpm	svc:/system/cluster/scdpm:default
cl_ccra	svc:/system/cluster/cl_ccra:default
scsymon_srv	svc:/system/cluster/scsymon_srv:default
spm	svc:/system/cluster/spm:default
cl_svc_cluster_milestone	svc:/system/cluster/cl_svc_cluster_milestone:default
cl_svc_enable	svc:/system/cluster/cl_svc_enable:default
network-multipathing	svc:/system/cluster/network-multipathing

管理工具

您可以使用命令行或 Oracle Solaris Cluster Manager 浏览器界面执行 Oracle Solaris Cluster 配置的管理任务。下节概述了 Oracle Solaris Cluster Manager 和命令行工具。

Oracle Solaris Cluster Manager 浏览器界面

Oracle Solaris Cluster 软件支持可用于在群集上执行各种管理任务的浏览器界面 Oracle Solaris Cluster Manager。有关更多信息，请参见[第 13 章 使用 Oracle Solaris Cluster Manager 浏览器界面](#)。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

以下是可以在 Oracle Solaris Cluster Manager 中执行的一些任务：

- 创建和更新区域群集
- 创建资源和资源组
- 将文件系统、逻辑主机或共享存储添加到区域群集
- 创建 Oracle Database 数据服务
- 管理全局群集或区域群集中的节点
- 添加并管理法定设备和服务器
- 添加并管理 NAS 存储设备，以及管理磁盘和设备组
- 管理 Geographic Edition 伙伴关系

命令行界面

您可以通过 `clsetup` 实用程序以交互方式执行大多数 Oracle Solaris Cluster 管理任务。本指南中的管理过程尽可能使用了 `clsetup` 实用程序。

您可以通过 `clsetup` 实用程序管理下列主菜单项。

- Quorum (法定)
- Resource groups (资源组)
- Data services (数据服务)
- Cluster interconnect (群集互连)
- Device groups and volumes (设备组和卷)
- Private hostnames (专用主机名)
- New nodes (新节点)
- Zone cluster (区域群集)
- Other cluster tasks (其他群集任务)

以下列表中提供了用于管理 Oracle Solaris Cluster 配置的其他命令。有关更多详细信息，请参阅手册页。

`if_mpadm(1M)`

将 IP 地址从 IP 网络多路径组 (IP Network Multipathing, IPMP) 中的一个适配器切换到另一个适配器。

`claccess(1CL)`

管理 Oracle Solaris Cluster 访问策略以添加节点。

`cldevice(1CL)`

管理 Oracle Solaris Cluster 设备。

`cldevicegroup(1CL)`

管理 Oracle Solaris Cluster 设备组。

`clinterconnect(1CL)`

管理 Oracle Solaris Cluster 互连。

`clnasdevice(1CL)`

管理对 Oracle Solaris Cluster 配置的 NAS 设备的访问。

`clnode(1CL)`

管理 Oracle Solaris Cluster 节点。

`clquorum(1CL)`

管理 Oracle Solaris Cluster 法定。

`clreslogicalhostname(1CL)`

管理 Oracle Solaris Cluster 的逻辑主机名资源。

`clresource(1CL)`

管理 Oracle Solaris Cluster 数据服务资源。

`clresourcegroup(1CL)`

管理 Oracle Solaris Cluster 数据服务资源。

`clresourcetype(1CL)`

管理 Oracle Solaris Cluster 数据服务资源。

`clressharedaddress(1CL)`

管理 Oracle Solaris Cluster 的共享地址资源。

`clsetup(1CL)`

创建区域群集并以交互方式执行 Oracle Solaris Cluster 配置。

`clsnmphost(1CL)`

管理 Oracle Solaris Cluster SNMP 主机。

`clsnmpmib(1CL)`

管理 Oracle Solaris Cluster SNMP MIB。

`clsnmpuser(1CL)`

管理 Oracle Solaris Cluster SNMP 用户。

`cltelemetryattribute(1CL)`

配置系统资源监视。

`cluster(1CL)`

管理 Oracle Solaris Cluster 配置的全局配置和全局状态。

`clzonecluster(1CL)`

创建和修改区域群集。

此外，您还可以使用命令来管理 Oracle Solaris Cluster 配置的卷管理器部分。这些命令依赖于您的群集使用的特定卷管理器。

群集管理准备

本节介绍如何做好管理群集前的准备工作。

记录 Oracle Solaris Cluster 的硬件配置

在改变 Oracle Solaris Cluster 配置时，记录针对您的站点的硬件配置。为了减轻管理工作量，请在更改或升级群集时参阅硬件文档。标注各个群集组件之间的电缆和连接也可以使管理变得更加容易。

记录原始群集配置和后来进行的更改，以便帮助第三方服务供应商在为您的群集提供服务时节省所需的时间。

使用管理控制台

可以将专用工作站或通过管理网络连接的工作站用作管理控制台，来管理活动群集。

可以使用 `pconsole` 实用程序运行每个群集节点的终端窗口以及可以向所有节点同时发出您在那里键入的命令的主窗口。有关在管理控制台上安装 `pconsole` 软件的信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[如何在管理控制台上安装 pconsole 软件](#)”。

可以使用 Oracle Solaris Cluster Manager 浏览器界面配置、监视和管理群集及群集组件。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

管理控制台并不是一个群集节点。管理控制台用来远程访问群集节点，或者通过公共网络，或者通过基于网络的终端集中器。

Oracle Solaris Cluster 不需要专用管理控制台，但使用控制台可带来以下好处：

- 通过在同一计算机上给控制台和管理工具分组来启用集中化的群集管理
- 通过 Enterprise Services 或服务提供商来提供可能更快的故障解决方案

备份群集

应定期备份您的群集。尽管 Oracle Solaris Cluster 软件可提供高可用环境（在若干个存储设备上保存着数据的镜像副本），但 Oracle Solaris Cluster 软件并不能代替定期备份。Oracle Solaris Cluster 配置可以承受多种故障，但并不能防止用户错误或程序错误或者灾难性故障。因此，您必须采用适当的备份过程，以防数据丢失。

备份应包含以下信息：

- 所有文件系统分区
- 所有的数据库数据（如果正在运行 DBMS 数据服务）
- 所有群集磁盘的磁盘分区信息

管理群集

表 2 “Oracle Solaris Cluster 管理工具”提供了管理群集的起点。

表 2 Oracle Solaris Cluster 管理工具

任务	工具	说明
远程登录群集	从命令行使用 Oracle Solaris pconsole 实用程序远程登录到群集。	“远程登录到群集” [30] “如何安全地连接到群集控制台” [31]
以交互方式配置群集	使用 clzonecluster 命令或 clsetup 实用程序。	如何访问群集配置实用程序 [31]
显示 Oracle Solaris Cluster 发行版本号和版本信息	组合使用 clnode 命令与 showrev -v -node 子命令和选项。	如何显示 Oracle Solaris Cluster 发行版本信息和版本信息 [32]
显示已安装的资源、资源组和资源类型	使用以下命令显示资源信息： ■ clresource ■ clresourcegroup ■ clresourcetype	如何显示已配置的资源类型、资源组和资源 [34]
以图形方式监视群集组件	使用 Oracle Solaris Cluster Manager。	请参见联机帮助。
以图形方式管理某些群集组件	使用 Oracle Solaris Cluster Manager。	请参见联机帮助。

任务	工具	说明
检查群集组件状态	组合使用 <code>cluster</code> 命令与 <code>status</code> 子命令。	如何检查群集组件的状态 [36]
检查公共网络上 IPMP 组的状态	对于全局群集，请组合使用 <code>clnode status</code> 命令与 <code>-m</code> 选项。 对于区域群集，请组合使用 <code>clzonecluster</code> 命令与 <code>show</code> 子命令。	如何检查公共网络的状态 [39]
查看群集配置	对于区域群集，请组合使用 <code>cluster</code> 命令与 <code>show</code> 子命令。 对于区域群集，请组合使用 <code>clzonecluster</code> 命令与 <code>show</code> 子命令。	如何查看群集配置 [39]
查看并显示配置的 NAS 设备	对于全局群集或区域群集，请组合使用 <code>clzonecluster</code> 命令与 <code>show</code> 子命令。	clnasdevice(1CL)
检查全局挂载点或检验群集配置	对于全局群集，请组合使用 <code>cluster</code> 命令与 <code>check</code> 子命令。 对于区域群集，请使用 <code>clzonecluster verify</code> 命令。	如何验证基本群集配置 [49]
查看 Oracle Solaris Cluster 命令日志的内容	检查 <code>/var/cluster/logs/commandlog</code> 文件。	如何查看 Oracle Solaris Cluster 命令日志的内容 [55]
查看 Oracle Solaris Cluster 系统消息	检查 <code>/var/adm/messages</code> 文件。	《在 Oracle Solaris 11.3 中排除系统管理问题》中的“系统消息格式” 还可以在 Oracle Solaris Cluster Manager 浏览器界面中查看节点的系统消息。有关 Oracle Solaris Cluster Manager 登录说明，请参见 如何访问 Oracle Solaris Cluster Manager [266] 。
监视 Solaris Volume Manager 的状态	使用 <code>metastat</code> 命令。	《Solaris Volume Manager 管理指南》

远程登录到群集

可以从命令行使用并行控制台访问 (pconsole) 实用程序远程登录到群集。pconsole 实用程序是 Oracle Solaris `terminal/pconsole` 软件包的一部分。可通过执行 `pkg install terminal/pconsole` 安装该软件包。pconsole 实用程序会为您在命令行上指定的每个远程主机创建一个主机终端窗口。该实用程序还会打开一个中央（或主）控制台窗口，该窗口将您在其中输入的内容传播到您打开的每个连接。

pconsole 实用程序可以从 X Windows 内运行，也可以在控制台模式下运行。在要用作群集的管理控制台的计算机上安装 pconsole。如果您的终端服务器允许连接到服务器 IP

地址上的特定端口号，那么除了指定主机名和 IP 地址外，您还可以指定端口号，如下所示：*terminal-server:portnumber*。

有关更多信息，请参见 *pconsole(1)* 手册页。

如何安全地连接到群集控制台

如果您的终端集中器或系统控制器支持 ssh，则可使用 *pconsole* 实用程序连接到这些系统的控制台。*pconsole* 实用程序是 Oracle Solaris terminal/pconsole 软件包的一部分，是在安装该软件包时安装的。*pconsole* 实用程序会为您在命令行上指定的每个远程主机创建一个主机终端窗口。该实用程序还会打开一个中央（或主）控制台窗口，该窗口将您在其中输入的内容传播到您打开的每个连接。有关更多信息，请参见 *pconsole(1)* 手册页。

▼ 如何访问群集配置实用程序

使用 *clsetup* 实用程序，可以交互方式创建区域群集，并可以为全局群集配置法定设备、资源组、群集传输、专用主机名、设备组以及新节点选项。*clzonecluster* 实用程序可对区域群集执行类似的配置任务。有关更多信息，请参见 *clsetup(1CL)* 和 *clzonecluster(1CL)* 手册页。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面完成此过程。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在全局群集的某个活动成员节点上承担 **root** 角色。
从全局群集的节点执行此过程中的所有步骤。

2. 启动配置实用程序。

```
phys-schost# clsetup
```

- 对于全局群集，使用 **clsetup** 命令启动该实用程序。

```
phys-schost# clsetup
```

此时将显示 q。

- 对于区域群集，使用 `clzonecluster` 命令启动该实用程序。本示例中的区域群集是 `SCZONE`。

```
phys-schost# clzonecluster configure SCZONE
```

您可以使用以下选项查看该实用程序中的可用操作：

```
clzc:sczone> ?
```

您还可以使用交互式的 `clsetup` 实用程序或 Oracle Solaris Cluster Manager 浏览器界面在群集范围中创建区域群集或者添加文件系统或存储设备。所有其他区域群集配置任务可通过 `clzonecluster configure` 命令来执行。有关使用 `clsetup` 实用程序或 Oracle Solaris Cluster Manager 配置区域群集的说明，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》。

3. 从菜单中选择配置。

按照屏幕上的说明完成任务。有关更多详细信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[创建和配置区域群集](#)”中的说明。

另请参见 有关更多信息，请参见 `clsetup` 或 `clzonecluster` 联机帮助手册页。

▼ 如何显示 Oracle Solaris Cluster 发行版本信息和版本信息

执行此过程不需要以 `root` 角色登录。从全局群集的节点执行此过程中的所有步骤。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- 显示 Oracle Solaris Cluster 发行版本信息和版本信息：

```
phys-schost# clnode show-rev -v -node
```

该命令显示了 Oracle Solaris Cluster 的发行版本号和所有 Oracle Solaris Cluster 软件包的版本字符串。

例 1 显示 Oracle Solaris Cluster 发行版本信息和版本信息

以下示例显示了群集的发行版本信息以及随 Oracle Solaris Cluster 4.2 提供的软件包的版本信息。

```
phys-schost# clnode show-rev
4.2
```



```
phys-schost#% clnode show-rev -v
```

```
Oracle Solaris Cluster 4.2 for Oracle Solaris 11 sparc
ha-cluster/data-service/apache           :4.2-0.30
ha-cluster/data-service/dhcp             :4.2-0.30
ha-cluster/data-service/dns              :4.2-0.30
ha-cluster/data-service/goldengate        :4.2-0.30
ha-cluster/data-service/glassfish-message-queue :4.2-0.30
ha-cluster/data-service/ha-ldom           :4.2-0.30
ha-cluster/data-service/ha-zones          :4.2-0.30
ha-cluster/data-service/iplanet-web-server :4.2-0.30
ha-cluster/data-service/jd-edwards-enterpriseone :4.2-0.30
ha-cluster/data-service/mysql            :4.2-0.30
ha-cluster/data-service/nfs              :4.2-0.30
ha-cluster/data-service/obiee            :4.2-0.30
ha-cluster/data-service/oracle-database   :4.2-0.30
ha-cluster/data-service/oracle-ebs        :4.2-0.30
ha-cluster/data-service/oracle-external-proxy :4.2-0.30
ha-cluster/data-service/oracle-http-server :4.2-0.30
ha-cluster/data-service/oracle-pmn-server :4.2-0.30
ha-cluster/data-service/oracle-traffic-director :4.2-0.30
ha-cluster/data-service/peoplesoft        :4.2-0.30
ha-cluster/data-service/postgresql        :4.2-0.30
ha-cluster/data-service/samba             :4.2-0.30
ha-cluster/data-service/sap-livecache     :4.2-0.30
ha-cluster/data-service/sapdb             :4.2-0.30
ha-cluster/data-service/sapnetweaver      :4.2-0.30
ha-cluster/data-service/siebel            :4.2-0.30
ha-cluster/data-service/sybase            :4.2-0.30
ha-cluster/data-service/timesten          :4.2-0.30
ha-cluster/data-service/tomcat            :4.2-0.30
ha-cluster/data-service/weblogic          :4.2-0.30
ha-cluster/developer/agent-builder        :4.2-0.30
ha-cluster/developer/api                 :4.2-0.30
ha-cluster/geo/geo-framework              :4.2-0.30
ha-cluster/geo/manual                     :4.2-0.30
ha-cluster/geo/replication/availability-suite :4.2-0.30
ha-cluster/geo/replication/data-guard     :4.2-0.30
ha-cluster/geo/replication/sbp            :4.2-0.30
ha-cluster/geo/replication/srdf           :4.2-0.30
ha-cluster/geo/replication/zfs-sa         :4.2-0.30
ha-cluster/group-package/ha-cluster-data-services-full :4.2-0.30
ha-cluster/group-package/ha-cluster-framework-full :4.2-0.30
ha-cluster/group-package/ha-cluster-framework-l10n :4.2-0.30
ha-cluster/group-package/ha-cluster-framework-minimal :4.2-0.30
ha-cluster/group-package/ha-cluster-framework-scm :4.2-0.30
ha-cluster/group-package/ha-cluster-framework-slm :4.2-0.30
ha-cluster/group-package/ha-cluster-full :4.2-0.30
ha-cluster/group-package/ha-cluster-geo-full :4.2-0.30
ha-cluster/group-package/ha-cluster-geo-incorporation :4.2-0.30
ha-cluster/group-package/ha-cluster-incorporation :4.2-0.30
ha-cluster/group-package/ha-cluster-minimal :4.2-0.30
ha-cluster/group-package/ha-cluster-quorum-server-full :4.2-0.30
```

```
ha-cluster/group-package/ha-cluster-quorum-server-l10n :4.2-0.30
ha-cluster/ha-service/derby :4.2-0.30
ha-cluster/ha-service/gds :4.2-0.30
ha-cluster/ha-service/gds2 :4.2-0.30
ha-cluster/ha-service/logical-hostname :4.2-0.30
ha-cluster/ha-service/smf-proxy :4.2-0.30
ha-cluster/ha-service/telemetry :4.2-0.30
ha-cluster/library/cacao :4.2-0.30
ha-cluster/library/ucmm :4.2-0.30
ha-cluster/locale :4.2-0.30
ha-cluster/release/name :4.2-0.30
ha-cluster/service/management :4.2-0.30
ha-cluster/service/management/slm :4.2-0.30
ha-cluster/service/quorum-server :4.2-0.30
ha-cluster/service/quorum-server/locale :4.2-0.30
ha-cluster/service/quorum-server/manual/locale :4.2-0.30
ha-cluster/storage/svm-mediator :4.2-0.30
ha-cluster/system/cfgchk :4.2-0.30
ha-cluster/system/core :4.2-0.30
ha-cluster/system/dsconfig-wizard :4.2-0.30
ha-cluster/system/install :4.2-0.30
ha-cluster/system/manual :4.2-0.30
ha-cluster/system/manual/data-services :4.2-0.30
ha-cluster/system/manual/locale :4.2-0.30
ha-cluster/system/manual/manager :4.2-0.30
ha-cluster/system/manual/manager-glassfish3:4.2-0.30
```

▼ 如何显示已配置的资源类型、资源组和资源

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以通过 Oracle Solaris Cluster Manager 浏览器界面查看资源和资源组。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

开始之前 非 root 角色的用户需要具有 solaris.cluster.read RBAC 授权才能使用该子命令。

- 显示群集的已配置资源类型、资源组和资源。

```
phys-schost# cluster show -t resource,resourcetype,resourcegroup
```

从全局群集的节点执行此过程中的所有步骤。要获取各个资源、资源组和资源类型的信息，请将 show 子命令与下列命令之一配合使用：

- resource
- resource group

■ resourcetype

例 2 显示已配置的资源类型、资源组和资源

以下示例显示了为群集 `schost` 配置的资源类型 (RT Name)、资源组 (RG Name) 和资源 (RS Name)。

```
phys-schost# cluster show -t resource,resourcetype,resourcegroup

=== Registered Resource Types ===

Resource Type:                SUNW.sctelemetry
RT_description:                sctelemetry service for Oracle Solaris Cluster
RT_version:                    1
API_version:                    7
RT_basedir:                    /usr/cluster/lib/rgm/rt/sctelemetry
Single_instance:              True
Proxy:                          False
Init_nodes:                    All potential masters
Installed_nodes:               <All>
Failover:                      False
Pkglist:                       <NULL>
RT_system:                     True
Global_zone:                   True

=== Resource Groups and Resources ===

Resource Group:                tel-rg
RG_description:                <NULL>
RG_mode:                       Failover
RG_state:                      Managed
Failback:                      False
Nodelist:                      phys-schost-2 phys-schost-1

--- Resources for Group tel-rg ---

Resource:                      tel-res
Type:                          SUNW.sctelemetry
Type_version:                  4.0
Group:                         tel-rg
R_description:
Resource_project_name:         default
Enabled{phys-schost-2}:        True
Enabled{phys-schost-1}:        True
Monitored{phys-schost-2}:      True
Monitored{phys-schost-1}:      True

Resource Type:                SUNW.qfs
RT_description:                SAM-QFS Agent on Oracle Solaris Cluster
RT_version:                    3.1
API_version:                    3
RT_basedir:                    /opt/SUNWsamfs/sc/bin
```

```
Single_instance:           False
Proxy:                     False
Init_nodes:                All potential masters
Installed_nodes:           <All>
Failover:                  True
Pkglist:                   <NULL>
RT_system:                 False
Global_zone:               True

=== Resource Groups and Resources ===
Resource Group:             qfs-rg
RG_description:             <NULL>
RG_mode:                    Failover
RG_state:                   Managed
Failback:                   False
Nodelist:                   phys-schost-2 phys-schost-1

--- Resources for Group qfs-rg ---
Resource:                   qfs-res
Type:                       SUNW.qfs
Type_version:               3.1
Group:                       qfs-rg
R_description:
Resource_project_name:      default
Enabled{phys-schost-2}:     True
Enabled{phys-schost-1}:     True
Monitored{phys-schost-2}:   True
Monitored{phys-schost-1}:   True
```

▼ 如何检查群集组件的状态

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面检查群集组件的状态。有关更多信息，请参见 Oracle Solaris Cluster Manager 联机帮助。

Oracle Solaris Cluster Manager 和 cluster status 命令还可以显示区域群集的状态。

开始之前 非 root 角色的用户需要具有 solaris.cluster.read RBAC 授权才能使用 status 子命令。

● 检查群集组件的状态。

```
phys-schost# cluster status
```

从全局群集的节点执行此过程中的所有步骤。

例 3 检查群集组件的状态

以下示例提供了 `cluster status` 命令返回的群集组件状态信息的样例。

```
phys-schost# cluster status
=== Cluster Nodes ===

--- Node Status ---

Node Name                               Status
-----
phys-schost-1                           Online
phys-schost-2                           Online

=== Cluster Transport Paths ===

Endpoint1           Endpoint2           Status
-----
phys-schost-1:nge1  phys-schost-4:nge1  Path online
phys-schost-1:e1000g1 phys-schost-4:e1000g1 Path online

=== Cluster Quorum ===

--- Quorum Votes Summary ---

Needed   Present   Possible
-----
3         3         4

--- Quorum Votes by Node ---

Node Name      Present   Possible   Status
-----
phys-schost-1  1         1         Online
phys-schost-2  1         1         Online

--- Quorum Votes by Device ---

Device Name      Present   Possible   Status
-----
/dev/did/rdisk/d2s2  1         1         Online
/dev/did/rdisk/d8s2  0         1         Offline

=== Cluster Device Groups ===

--- Device Group Status ---

Device Group Name  Primary      Secondary   Status
-----
```

```
schost-2          phys-schost-2      -          Degraded
```

```
--- Spare, Inactive, and In Transition Nodes ---
```

Device Group Name	Spare Nodes	Inactive Nodes	In Transition Nodes
schost-2	-	-	-

```
=== Cluster Resource Groups ===
```

Group Name	Node Name	Suspended	Status
test-rg	phys-schost-1	No	Offline
	phys-schost-2	No	Online
test-rg	phys-schost-1	No	Offline
	phys-schost-2	No	Error--stop failed
test-rg	phys-schost-1	No	Online
	phys-schost-2	No	Online

```
=== Cluster Resources ===
```

Resource Name	Node Name	Status	Message
test_1	phys-schost-1	Offline	Offline
	phys-schost-2	Online	Online
test_1	phys-schost-1	Offline	Offline
	phys-schost-2	Stop failed	Faulted
test_1	phys-schost-1	Online	Online
	phys-schost-2	Online	Online

Device Instance	Node	Status
/dev/did/rdsk/d2	phys-schost-1	Ok
/dev/did/rdsk/d3	phys-schost-1	Ok
	phys-schost-2	Ok
/dev/did/rdsk/d4	phys-schost-1	Ok
	phys-schost-2	Ok
/dev/did/rdsk/d6	phys-schost-2	Ok

```
=== Zone Clusters ===
```

```
--- Zone Cluster Status ---
```

Name	Node Name	Zone HostName	Status	Zone Status
sczone	schost-1	sczone-1	Online	Running
schost-2	sczone-2	OnlineRunning		

▼ 如何检查公共网络的状态

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

要检查 IP 网络多路径组的状态，请将该命令与 `clnode status` 命令一起使用。

开始之前 非 root 角色的用户需要具有 `solaris.cluster.read` RBAC 授权才能使用该子命令。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面检查节点的状态。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

● 检查群集组件的状态。

```
phys-schost# clnode status -m
```

从全局群集的节点执行此过程中的所有步骤。

例 4 检查公共网络状态

以下示例提供了 `clnode status` 命令返回的群集组件状态信息的样例。

```
% clnode status -m
--- Node IPMP Group Status ---

Node Name      Group Name    Status  Adapter  Status
-----
phys-schost-1  test-rg      Online  nge2     Online
phys-schost-2 test-rg      Online  nge3     Online
```

▼ 如何查看群集配置

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以通过 Oracle Solaris Cluster Manager 浏览器界面查看群集的配置。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

开始之前 非 root 角色的用户需要具有 `solaris.cluster.read` RBAC 授权才能使用 `status` 子命令。

- 查看全局群集或区域群集的配置。

```
% cluster show
```

从全局群集的节点执行此过程中的所有步骤。

从全局群集节点运行 `cluster show` 命令可显示有关群集的详细配置信息以及区域群集的信息（如果已配置这些群集）。

您还可以使用 `clzonecluster show` 命令仅查看区域群集的配置信息。区域群集的属性包括区域群集名称、IP 类型、自动引导和区域路径。`show` 子命令在区域群集内部运行，并且仅适用于该特定区域群集。从区域群集节点运行 `clzonecluster show` 命令时，将仅检索有关对该特定区域群集可见的对象的状态。

要显示有关 `cluster` 命令的更多信息，请使用详细 (verbose) 选项。有关详细信息，请参见 [cluster\(1CL\)](#) 手册页。有关 `clzonecluster` 的更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

例 5 查看全局群集配置

以下示例列出了有关全局群集的配置信息。如果您配置了区域群集，则也会列出相应的信息。

```
phys-schost# cluster show

=== Cluster ===

Cluster Name:                cluster-1
clusterid:                   0x4DA2C888
installmode:                  disabled
heartbeat_timeout:           10000
heartbeat_quantum:           1000
private_netaddr:              172.11.0.0
private_netmask:              255.255.248.0
max_nodes:                    64
max_privatenets:              10
num_zoneclusters:             12
udp_session_timeout:          480
concentrate_load:             False
global_fencing:               prefer3
```



```

Node List:                phys-schost-1
Node Zones:               phys_schost-2:za

=== Host Access Control ===

Cluster name:             clustser-1
Allowed hosts:            phys-schost-1, phys-schost-2:za
Authentication Protocol:  sys

=== Cluster Nodes ===

Node Name:                phys-schost-1
Node ID:                  1
Enabled:                  yes
privatehostname:          clusternode1-priv
reboot_on_path_failure:   disabled
globalzoneshares:        3
defaultpsetmin:           1
quorum_vote:              1
quorum_defaultvote:       1
quorum_resv_key:          0x43CB1E1800000001
Transport Adapter List:   net1, net3

--- Transport Adapters for phys-schost-1 ---

Transport Adapter:        net1
Adapter State:            Enabled
Adapter Transport Type:   dlpi
Adapter Property(device_name): net
Adapter Property(device_instance): 1
Adapter Property(lazy_free): 1
Adapter Property(dlpi_heartbeat_timeout): 10000
Adapter Property(dlpi_heartbeat_quantum): 1000
Adapter Property(nw_bandwidth): 80
Adapter Property(bandwidth): 10
Adapter Property(ip_address): 172.16.1.1
Adapter Property(netmask): 255.255.255.128
Adapter Port Names:       0
Adapter Port State(0):    Enabled

Transport Adapter:        net3
Adapter State:            Enabled
Adapter Transport Type:   dlpi
Adapter Property(device_name): net
Adapter Property(device_instance): 3
Adapter Property(lazy_free): 0
Adapter Property(dlpi_heartbeat_timeout): 10000
Adapter Property(dlpi_heartbeat_quantum): 1000
Adapter Property(nw_bandwidth): 80
Adapter Property(bandwidth): 10
Adapter Property(ip_address): 172.16.0.129
Adapter Property(netmask): 255.255.255.128
Adapter Port Names:       0
Adapter Port State(0):    Enabled

```

```

--- SNMP MIB Configuration on phys-schost-1 ---

SNMP MIB Name:          Event
State:                  Disabled
Protocol:                SNMPv2

--- SNMP Host Configuration on phys-schost-1 ---

--- SNMP User Configuration on phys-schost-1 ---

SNMP User Name:          foo
Authentication Protocol: MD5
Default User:            No

Node Name:               phys-schost-2:za
Node ID:                  2
Type:                     cluster
Enabled:                  yes
privatehostname:          clusternode2-priv
reboot_on_path_failure:   disabled
globalzoneshares:         1
defaultpsetmin:           2
quorum_vote:              1
quorum_defaultvote:       1
quorum_resv_key:          0x43CB1E1800000002
Transport Adapter List:    e1000g1, nge1

--- Transport Adapters for phys-schost-2 ---

Transport Adapter:        e1000g1
Adapter State:            Enabled
Adapter Transport Type:   dlpi
Adapter Property(device_name): e1000g
Adapter Property(device_instance): 2
Adapter Property(lazy_free): 0
Adapter Property(dlpi_heartbeat_timeout): 10000
Adapter Property(dlpi_heartbeat_quantum): 1000
Adapter Property(nw_bandwidth): 80
Adapter Property(bandwidth): 10
Adapter Property(ip_address): 172.16.0.130
Adapter Property(netmask): 255.255.255.128
Adapter Port Names:       0
Adapter Port State(0):    Enabled

Transport Adapter:        nge1
Adapter State:            Enabled
Adapter Transport Type:   dlpi
Adapter Property(device_name): nge
Adapter Property(device_instance): 3
Adapter Property(lazy_free): 1
Adapter Property(dlpi_heartbeat_timeout): 10000
Adapter Property(dlpi_heartbeat_quantum): 1000
Adapter Property(nw_bandwidth): 80

```

```

Adapter Property(bandwidth):          10
Adapter Property(ip_address):         172.16.1.2
Adapter Property(netmask):            255.255.255.128
Adapter Port Names:                   0
Adapter Port State(0):                Enabled

--- SNMP MIB Configuration on phys-schost-2 ---

SNMP MIB Name:                        Event
State:                                Disabled
Protocol:                             SNMPv2

--- SNMP Host Configuration on phys-schost-2 ---

--- SNMP User Configuration on phys-schost-2 ---

=== Transport Cables ===

Transport Cable:                      phys-schost-1:e1000g1,switch2@1
Cable Endpoint1:                     phys-schost-1:e1000g1
Cable Endpoint2:                     switch2@1
Cable State:                         Enabled

Transport Cable:                      phys-schost-1:nge1,switch1@1
Cable Endpoint1:                     phys-schost-1:nge1
Cable Endpoint2:                     switch1@1
Cable State:                         Enabled

Transport Cable:                      phys-schost-2:nge1,switch1@2
Cable Endpoint1:                     phys-schost-2:nge1
Cable Endpoint2:                     switch1@2
Cable State:                         Enabled

Transport Cable:                      phys-schost-2:e1000g1,switch2@2
Cable Endpoint1:                     phys-schost-2:e1000g1
Cable Endpoint2:                     switch2@2
Cable State:                         Enabled

=== Transport Switches ===

Transport Switch:                     switch2
Switch State:                         Enabled
Switch Type:                         switch
Switch Port Names:                   1 2
Switch Port State(1):                Enabled
Switch Port State(2):                Enabled

Transport Switch:                     switch1
Switch State:                         Enabled
Switch Type:                         switch
Switch Port Names:                   1 2
Switch Port State(1):                Enabled
Switch Port State(2):                Enabled

```

=== Quorum Devices ===

Quorum Device Name:	d3
Enabled:	yes
Votes:	1
Global Name:	/dev/did/rdisk/d3s2
Type:	shared_disk
Access Mode:	scsi3
Hosts (enabled):	phys-schost-1, phys-schost-2

Quorum Device Name:	qs1
Enabled:	yes
Votes:	1
Global Name:	qs1
Type:	quorum_server
Hosts (enabled):	phys-schost-1, phys-schost-2
Quorum Server Host:	10.11.114.83
Port:	9000

=== Device Groups ===

Device Group Name:	testdg3
Type:	SVM
failback:	no
Node List:	phys-schost-1, phys-schost-2
preferenced:	yes
numsecondaries:	1
diskset name:	testdg3

=== Registered Resource Types ===

Resource Type:	SUNW.LogicalHostname:2
RT_description:	Logical Hostname Resource Type
RT_version:	4
API_version:	2
RT_basedir:	/usr/cluster/lib/rgm/rt/hafoip
Single_instance:	False
Proxy:	False
Init_nodes:	All potential masters
Installed_nodes:	<All>
Failover:	True
Pkglist:	<NULL>
RT_system:	True
Global_zone:	True

Resource Type:	SUNW.SharedAddress:2
RT_description:	HA Shared Address Resource Type
RT_version:	2
API_version:	2
RT_basedir:	/usr/cluster/lib/rgm/rt/hascip
Single_instance:	False
Proxy:	False

```

Init_nodes:                <Unknown>
Installed_nodes:           <All>
Failover:                  True
Pkglist:                   <NULL>
RT_system:                True
Global_zone:              True
Resource Type:             SUNW.HAStoragePlus:4
RT_description:            HA Storage Plus
RT_version:                4
API_version:               2
RT_basedir:                /usr/cluster/lib/rgm/rt/hastorageplus
Single_instance:          False
Proxy:                    False
Init_nodes:                All potential masters
Installed_nodes:           <All>
Failover:                  False
Pkglist:                   <NULL>
RT_system:                True
Global_zone:              True
Resource Type:             SUNW.haderby
RT_description:            haderby server for Oracle Solaris Cluster
RT_version:                1
API_version:               7
RT_basedir:                /usr/cluster/lib/rgm/rt/haderby
Single_instance:          False
Proxy:                    False
Init_nodes:                All potential masters
Installed_nodes:           <All>
Failover:                  False
Pkglist:                   <NULL>
RT_system:                True
Global_zone:              True
Resource Type:             SUNW.sctelemetry
RT_description:            sctelemetry service for Oracle Solaris Cluster
RT_version:                1
API_version:               7
RT_basedir:                /usr/cluster/lib/rgm/rt/sctelemetry
Single_instance:          True
Proxy:                    False
Init_nodes:                All potential masters
Installed_nodes:           <All>
Failover:                  False
Pkglist:                   <NULL>
RT_system:                True
Global_zone:              True
=== Resource Groups and Resources ===

Resource Group:            HA_RG
RG_description:            <Null>
RG_mode:                   Failover
RG_state:                  Managed
Failback:                  False
Nodelist:                  phys-schost-1 phys-schost-2

```

```

--- Resources for Group HA_RG ---

Resource:                                HA_R
Type:                                    SUNW.HAStoragePlus:4
Type_version:                            4
Group:                                    HA_RG
R_description:
Resource_project_name:                    SCSLM_HA_RG
Enabled{phys-schost-1}:                   True
Enabled{phys-schost-2}:                   True
Monitored{phys-schost-1}:                 True
Monitored{phys-schost-2}:                 True

Resource Group:                           cl-db-rg
RG_description:                           <Null>
RG_mode:                                   Failover
RG_state:                                   Managed
Failback:                                   False
Nodelist:                                   phys-schost-1 phys-schost-2

--- Resources for Group cl-db-rg ---

Resource:                                cl-db-rs
Type:                                    SUNW.haderby
Type_version:                            1
Group:                                    cl-db-rg
R_description:
Resource_project_name:                    default
Enabled{phys-schost-1}:                   True
Enabled{phys-schost-2}:                   True
Monitored{phys-schost-1}:                 True
Monitored{phys-schost-2}:                 True

Resource Group:                           cl-tlmtry-rg
RG_description:                           <Null>
RG_mode:                                   Scalable
RG_state:                                   Managed
Failback:                                   False
Nodelist:                                   phys-schost-1 phys-schost-2

--- Resources for Group cl-tlmtry-rg ---

Resource:                                cl-tlmtry-rs
Type:                                    SUNW.sctelemetry
Type_version:                            1
Group:                                    cl-tlmtry-rg
R_description:
Resource_project_name:                    default
Enabled{phys-schost-1}:                   True
Enabled{phys-schost-2}:                   True
Monitored{phys-schost-1}:                 True
Monitored{phys-schost-2}:                 True

=== DID Device Instances ===

```

```

DID Device Name:                /dev/did/rdisk/d1
Full Device Path:                phys-schost-1:/dev/rdisk/c0t2d0
Replication:                    none
default_fencing:                global

DID Device Name:                /dev/did/rdisk/d2
Full Device Path:                phys-schost-1:/dev/rdisk/c1t0d0
Replication:                    none
default_fencing:                global

DID Device Name:                /dev/did/rdisk/d3
Full Device Path:                phys-schost-2:/dev/rdisk/c2t1d0
Full Device Path:                phys-schost-1:/dev/rdisk/c2t1d0
Replication:                    none
default_fencing:                global

DID Device Name:                /dev/did/rdisk/d4
Full Device Path:                phys-schost-2:/dev/rdisk/c2t2d0
Full Device Path:                phys-schost-1:/dev/rdisk/c2t2d0
Replication:                    none
default_fencing:                global

DID Device Name:                /dev/did/rdisk/d5
Full Device Path:                phys-schost-2:/dev/rdisk/c0t2d0
Replication:                    none
default_fencing:                global

DID Device Name:                /dev/did/rdisk/d6
Full Device Path:                phys-schost-2:/dev/rdisk/c1t0d0
Replication:                    none
default_fencing:                global

=== NAS Devices ===

Nas Device:                     nas_filer1
Type:                           sun_uss
nodeIPs{phys-schost-2}:         10.134.112.112
nodeIPs{phys-schost-1}:         10.134.112.113
User ID: root

```

例 6 查看区域群集配置

以下示例列出了包含 RAC 的区域群集配置的属性。

```

% clzonecluster show
=== Zone Clusters ===

Zone Cluster Name:              sczone
zonename:                       sczone
zonepath:                       /zones/sczone
autoboot:                       TRUE
ip-type:                        shared
enable_priv_net:                TRUE

```

```

--- Solaris Resources for sczone ---

Resource Name:          net
address:                172.16.0.1
physical:               auto

Resource Name:          net
address:                172.16.0.2
physical:               auto

Resource Name:          fs
dir:                    /local/ufs-1
special:                /dev/md/ds1/dsk/d0
raw:                    /dev/md/ds1/rdisk/d0
type:                   ufs
options:                [logging]

Resource Name:          fs
dir:                    /gz/db_qfs/CrsHome
special:                CrsHome
raw:
type:                   samfs
options:                []

Resource Name:          fs
dir:                    /gz/db_qfs/CrsData
special:                CrsData
raw:
type:                   samfs
options:                []

Resource Name:          fs
dir:                    /gz/db_qfs/OraHome
special:                OraHome
raw:
type:                   samfs
options:                []

Resource Name:          fs
dir:                    /gz/db_qfs/OraData
special:                OraData
raw:
type:                   samfs
options:                []

--- Zone Cluster Nodes for sczone ---

Node Name:              sczone-1
physical-host:          sczone-1
hostname:               lzzone-1

Node Name:              sczone-2

```



```
physical-host:                               sczone-2
hostname:lzzone-2
```

您还可以使用 `clnasdevice show` 子命令或 Oracle Solaris Cluster Manager 查看为全局群集或区域群集配置的 NAS 设备。有关更多信息，请参见 [clnasdevice\(1CL\)](#) 手册页。

▼ 如何验证基本群集配置

`cluster` 命令使用 `check` 子命令验证全局群集正常工作所需的基本配置。如果所有检查均未失败，`cluster check` 将返回到 shell 提示符。如果某项检查失败，`cluster check` 将在指定的输出目录或默认输出目录中生成报告。如果对多个节点运行 `cluster check`，则 `cluster check` 将为每个节点生成一个报告并为多节点检查生成一个报告。也可使用 `cluster list-checks` 命令显示所有可用群集检查的列表。

除不需要用户交互即可运行的基本检查之外，该命令还可以运行交互式检查和功能检查。未指定 `-k keyword` 选项时将运行基本检查。

- 交互式检查需要从用户获取检查无法确定的信息。该检查会提示用户提供所需的信息，例如，固件版本号。使用 `-k interactive` 关键字可指定一个或多个交互式检查。
- 功能检查执行群集的特定功能或行为。该检查会提示用户输入信息（例如，要故障转移到的节点）以及确认是否开始或继续检查。使用 `-k functional check-id` 关键字指定功能检查。一次只能执行一项功能检查。

注 - 因为某些功能检查涉及到中断群集服务，所以在您阅读完检查的详细描述并确定是否需要先使群集脱离生产环境之前，请勿启动任何功能检查。要显示此信息，请使用以下命令：

```
% cluster list-checks -v -C checkID
```

可在详细模式下运行带有 `-v` 标志的 `cluster check` 命令，以显示进度信息。

注 - 在执行可能导致设备、卷管理组件或 Oracle Solaris Cluster 的配置发生更改的管理过程之后，应运行 `cluster check`。

在全局群集节点上运行 [clzonecluster\(1CL\)](#) 命令时，会运行一组检查，以验证区域群集正常工作所需的配置。如果所有检查都通过，`clzonecluster verify` 将返回到 shell 提示符，您可以放心地安装该区域群集。如果某项检查失败，则 `clzonecluster verify` 将报告检验失败的全局群集节点。如果对多个节点运行 `clzonecluster verify`，将针对每个节点和多个节点的检查分别生成一个报告。不允许在区域群集内部运行 `verify` 子命令。

1. 在全局群集的某个活动成员节点上承担 **root** 角色。

```
phys-schost# su
```

从全局群集的节点执行此过程中的所有步骤。

2. 确保您具有最新的检查。
 - a. 转到 [My Oracle Support](#) 的 "Patches & Updates" (修补程序和更新) 选项卡。
 - b. 在 "Advanced Search" (高级搜索) 中, 选择 **Solaris Cluster** 作为 "Product" (产品), 并在 "Description" (描述) 字段中键入 **check**。
该搜索将找到包含 **check** 的 Oracle Solaris Cluster 软件更新。
 - c. 应用任何尚未安装在群集上的软件更新。
3. 运行基本验证检查。

```
phys-schost# cluster check -v -o outputdir
```

-v 详细模式。

-o outputdir 将输出重定向到 *outputdir* 子目录。

该命令会运行所有可用的基本检查。不会影响任何群集功能。

4. 运行交互式验证检查。

```
phys-schost# cluster check -v -k interactive -o outputdir
```

-k interactive 指定运行交互式验证检查。

该命令会运行所有可用的交互式检查并提示您提供所需的群集相关信息。不会影响任何群集功能。

5. 运行功能验证检查。

- a. 以非详细模式列出所有可用的功能检查。

```
phys-schost# cluster list-checks -k functional
```

- b. 确定哪些功能检查执行的操作会干扰生产环境中的群集可用性或服务。
例如, 功能检查可能会引起节点出现紧急情况或故障转移到其他节点。

```
phys-schost# cluster list-checks -v -c check-ID
```

-c check-ID 指定特定检查。

c. 如果要执行的功能检查可能会中断群集的正常工作的，请确保群集不在生产环境中。

d. 启动功能检查。

```
phys-schost# cluster check -v -k functional -C check-ID -o outputdir
```

-k functional 指定运行功能验证检查。

响应来自检查的提示，确认应运行该检查以及必须执行的任何信息或操作。

e. 对于要运行的其余每个功能检查，重复执行步骤 c 和步骤 d。

注 - 为了进行记录，请为所运行的每个检查指定唯一 *outputdir* 子目录名称。如果重用 *outputdir* 名称，则新检查的输出将覆盖重用的 *outputdir* 子目录的现有内容。

6. 如果配置了区域群集，请检验区域群集的配置以了解是否可以安装区域群集。

```
phys-schost# clzonecluster verify zone-cluster-name
```

7. 记录群集配置以供将来诊断使用。

请参见《Oracle Solaris Cluster 4.3 软件安装指南》中的“如何记录群集配置的诊断数据”。

例 7 检查全局群集配置并且所有基本检查均通过

以下示例显示了针对节点 *phys-schost-1* 和 *phys-schost-2* 在详细模式下运行的 *cluster check*，其中节点通过了所有检查。

```
phys-schost# cluster check -v -n phys-schost-1,phys-schost-2
```

```
cluster check: Requesting explorer data and node report from phys-schost-1.
cluster check: Requesting explorer data and node report from phys-schost-2.
cluster check: phys-schost-1: Explorer finished.
cluster check: phys-schost-1: Starting single-node checks.
cluster check: phys-schost-1: Single-node checks finished.
cluster check: phys-schost-2: Explorer finished.
cluster check: phys-schost-2: Starting single-node checks.
cluster check: phys-schost-2: Single-node checks finished.
cluster check: Starting multi-node checks.
cluster check: Multi-node checks finished
```

例 8 列出交互式验证检查

以下示例列出了可用于在群集上运行的所有交互式检查。示例输出显示了可能的检查样例；实际的可用检查因各配置而异。

```
# cluster list-checks -k interactive
```

```
Some checks might take a few moments to run (use -v to see progress)...
I6994574:(Moderate)Fix for GLDV3 interfaces on cluster transport vulnerability applied?
```

例 9 运行功能验证检查

以下示例首先显示了功能检查的详细列表。随后列出了检查 F6968101 的详细描述，指出该检查会中断群集服务。群集将脱离生产环境。然后将运行功能检查，且详细输出会记录到 funct.test.F6968101.12Jan2011 子目录中。示例输出显示了可能的检查样例；实际的可用检查因各配置而异。

```
# cluster list-checks -k functional
F6968101: (Critical) Perform resource group switchover
F6984120: (Critical) Induce cluster transport network failure - single adapter.
F6984121: (Critical) Perform cluster shutdown
F6984140: (Critical) Induce node panic
# cluster list-checks -v -C F6968101
F6968101: (Critical) Perform resource group switchover
Keywords: SolarisCluster3.x, functional
Applicability: Applicable if multi-node cluster running live.
Check Logic: Select a resource group and destination node. Perform
'clresourcegroup switch' on specified resource group
either to specified node or to all nodes in succession.
Version: 1.2
Revision Date: 12/10/10

    使群集脱离生产环境

# cluster list-checks -k functional -C F6968101 -o funct.test.F6968101.12Jan2011
F6968101
initializing...
initializing xml output...
loading auxiliary data...
starting check run...
  pschost1, pschost2, pschost3, pschost4: F6968101.... starting:
Perform resource group switchover
```

```
=====
```

```
>>> Functional Check
```

```
'Functional' checks exercise cluster behavior. It is recommended that you
do not run this check on a cluster in production mode.' It is recommended
that you have access to the system console for each cluster node and
observe any output on the consoles while the check is executed.
```

```
If the node running this check is brought down during execution the check
must be rerun from this same node after it is rebooted into the cluster in
order for the check to be completed.
```

```
Select 'continue' for more details on this check.
```

```
1) continue
```

```
2) exit
```

```
choice: 1
```

```
=====
```

```
>>> Check Description <<<
```

请按照屏幕上的说明进行操作

例 10 检查全局群集配置并且某项检查未通过

以下示例显示，群集 suncluster 中的节点 phys-schost-2 缺少挂载点 /global/phys-schost-1。将在输出目录 /var/cluster/logs/cluster_check/<timestamp> 中创建报告。

```
phys-schost# cluster check -v -n phys-schost-1,phys-schost-2 -o/var/cluster/logs/
cluster_check/Dec5/

cluster check: Requesting explorer data and node report from phys-schost-1.
cluster check: Requesting explorer data and node report from phys-schost-2.
cluster check: phys-schost-1: Explorer finished.
cluster check: phys-schost-1: Starting single-node checks.
cluster check: phys-schost-1: Single-node checks finished.
cluster check: phys-schost-2: Explorer finished.
cluster check: phys-schost-2: Starting single-node checks.
cluster check: phys-schost-2: Single-node checks finished.
cluster check: Starting multi-node checks.
cluster check: Multi-node checks finished.
cluster check: One or more checks failed.
cluster check: The greatest severity of all check failures was 3 (HIGH).
cluster check: Reports are in /var/cluster/logs/cluster_check/<Dec5>.
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.suncluster.txt
...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE: Global filesystem /etc/vfstab entries are not consistent across
all Oracle SolarisCluster 4.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
```

▼ 如何检查全局挂载点

`cluster` 命令会执行多项检查，以确定 `/etc/vfstab` 文件中是否存在与群集文件系统及其全局挂载点有关的配置错误。有关更多信息，请参见 [cluster\(1CL\)](#) 手册页。

注 - 在进行了影响到设备或卷管理组件的群集配置更改后，请运行 `cluster check`。

1. 在全局群集的某个活动成员节点上承担 **root** 角色。
从全局群集的节点执行此过程中的所有步骤。

```
% su
```

2. 检验全局群集配置。

```
phys-schost# cluster check
```

例 11 检查全局挂载点

以下示例显示，群集 `suncluster` 中的节点 `phys-schost-2` 缺少挂载点 `/global/schost-1`。报告将发送到输出目录 `/var/cluster/logs/cluster_check/<timestamp>/`。

```
phys-schost# cluster check -v1 -n phys-schost-1,phys-schost-2 -o /var/cluster//logs/cluster_check/Dec5/
```

```
cluster check: Requesting explorer data and node report from phys-schost-1.
cluster check: Requesting explorer data and node report from phys-schost-2.
cluster check: phys-schost-1: Explorer finished.
cluster check: phys-schost-1: Starting single-node checks.
cluster check: phys-schost-1: Single-node checks finished.
cluster check: phys-schost-2: Explorer finished.
cluster check: phys-schost-2: Starting single-node checks.
cluster check: phys-schost-2: Single-node checks finished.
cluster check: Starting multi-node checks.
cluster check: Multi-node checks finished.
cluster check: One or more checks failed.
cluster check: The greatest severity of all check failures was 3 (HIGH).
cluster check: Reports are in /var/cluster/logs/cluster_check/Dec5.
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.suncluster.txt

...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE: Global filesystem /etc/vfstab entries are not consistent across
all Oracle Solaris Cluster 4.x nodes.
```

```

ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.phys-schost-1.txt

...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 1398
SEVERITY : HIGH
FAILURE: An unsupported server is being used as an Oracle Solaris Cluster 4.x node.
ANALYSIS : This server may not been qualified to be used as an Oracle Solaris Cluster 4.x
node.
Only servers that have been qualified with Oracle Solaris Cluster 4.0 are supported as
Oracle Solaris Cluster 4.x nodes.
RECOMMEND: Because the list of supported servers is always being updated, check with
your Oracle representative to get the latest information on what servers
are currently supported and only use a server that is supported with Oracle Solaris Cluster 4.
x.
...
#

```

▼ 如何查看 Oracle Solaris Cluster 命令日志的内容

/var/cluster/logs/commandlog ASCII 文本文件包含在群集中执行的选定 Oracle Solaris Cluster 命令的记录。一旦设置群集，系统会自动启动对命令的日志记录，并在您关闭群集时结束。在所有已启动并以群集模式引导的节点上，系统都会记录命令。

不在该文件中记录的命令包括那些显示群集配置和当前状态的命令。

在该文件中记录的命令包括那些配置和更改群集当前状态的命令：

- claccess
- cldevice
- cldevicegroup
- clinterconnect
- clnasdevice
- clnode
- clquorum
- clreslogicalhostname

- clresource
- clresourcegroup
- clresourcetype
- clressharedaddress
- clsetup
- clsnmpghost
- clsnmpmib
- clnsmpuser
- cltelemetryattribute
- cluster
- clzonecluster

commandlog 文件中的记录可包含下列元素：

- 日期和时间戳
- 发出命令的主机的名称
- 命令的进程 ID
- 执行命令的用户的登录名
- 用户已执行的命令，包括所有选项和操作数

注 - 命令选项在 commandlog 文件中用引号括起，这样您就可以轻松找到它们，然后复制粘贴到 shell 中并在 shell 中执行。

- 已执行命令的退出状态

注 - 如果命令异常中止并产生未知结果，则 Oracle Solaris Cluster 软件不会在 commandlog 文件中显示退出状态。

默认情况下，系统每周定期对 commandlog 文件进行一次归档。要更改 commandlog 文件的归档策略，请在群集的每个节点上运行 crontab 命令。有关更多信息，请参见 [crontab\(1\)](#) 手册页。

在任意给定时刻，Oracle Solaris Cluster 软件最多可在每个群集节点上维护八个先前归档的 commandlog 文件。当周的 commandlog 文件名为 commandlog。时间最近的完整的周归档文件名为 commandlog.0。时间最早的完整的周归档文件名为 commandlog.7。

- 查看当周 commandlog 文件的内容，一次显示一屏。

```
phys-schost# more /var/cluster/logs/commandlog
```


例 12 查看 Oracle Solaris Cluster 命令日志的内容

以下示例显示了通过执行 `more` 命令显示的 `commandlog` 文件内容。

```
more -lines10 /var/cluster/logs/commandlog
11/11/2006 09:42:51 phys-schost-1 5222 root START - clsetup
11/11/2006 09:43:36 phys-schost-1 5758 root START - clrg add "app-sa-1"
11/11/2006 09:43:36 phys-schost-1 5758 root END 0
11/11/2006 09:43:36 phys-schost-1 5760 root START - clrg set -y
"RG_description=Department Shared Address RG" "app-sa-1"
11/11/2006 09:43:37 phys-schost-1 5760 root END 0
11/11/2006 09:44:15 phys-schost-1 5810 root START - clrg online "app-sa-1"
11/11/2006 09:44:15 phys-schost-1 5810 root END 0
11/11/2006 09:44:19 phys-schost-1 5222 root END -20988320
12/02/2006 14:37:21 phys-schost-1 5542 jbloggs START - clrg -c -g "app-sa-1"
-y "RG_description=Joe Bloggs Shared Address RG"
12/02/2006 14:37:22 phys-schost-1 5542 jbloggs END 0
```


Oracle Solaris Cluster 和 RBAC

本章介绍了与 Oracle Solaris Cluster 相关的基于角色的访问控制 (role-based access control, RBAC)，涵盖的主题包括：

- “通过 Oracle Solaris Cluster 设置和使用 RBAC” [59]
- “Oracle Solaris Cluster RBAC 权限配置文件” [59]
- “使用 Oracle Solaris Cluster 管理权限配置文件创建和分配 RBAC 角色” [61]
- “修改用户的 RBAC 属性” [62]

通过 Oracle Solaris Cluster 设置和使用 RBAC

使用下表确定设置和使用 RBAC 时要参考的文档。本章稍后介绍在 Oracle Solaris Cluster 软件中设置和使用 RBAC 所需执行的具体步骤。

任务	说明
详细了解 RBAC	《在 Oracle Solaris 11.3 中确保用户和进程的安全》中的 第 1 章, “关于使用权限控制用户和进程”
设置和管理元素以及使用 RBAC	《在 Oracle Solaris 11.3 中确保用户和进程的安全》中的 第 3 章, “在 Oracle Solaris 中指定权限”
详细了解 RBAC 元素和工具	《在 Oracle Solaris 11.3 中确保用户和进程的安全》中的 第 8 章, “Oracle Solaris 权限参考信息”

Oracle Solaris Cluster RBAC 权限配置文件

您在命令行发出的选定 Oracle Solaris Cluster 命令和选项使用 RBAC 进行授权。需要进行 RBAC 授权的 Oracle Solaris Cluster 命令和选项将需要使用以下一个或多个授权级别。Oracle Solaris Cluster RBAC 权限配置文件适用于全局群集中的节点。

`solaris.cluster.read`

对执行列出、显示和其他读取操作进行授权。

`solaris.cluster.admin`

对更改群集对象的状态的操作进行授权。

`solaris.cluster.modify`

对创建、删除以及更改群集对象的属性的操作进行授权。

有关 Oracle Solaris Cluster 命令所需的 RBAC 授权的更多信息，请参见该命令的手册页。

RBAC 权限配置文件包含一个或多个 RBAC 授权。您可以将这些权限配置文件分配给用户或角色，使其对 Oracle Solaris Cluster 具有不同级别的访问权限。Oracle 在 Oracle Solaris Cluster 软件中提供了以下权限配置文件。

权限配置文件	授权和安全属性	授予的权限
Oracle Solaris Cluster 命令	使用 <code>euclid=0</code> 安全属性运行的 Oracle Solaris Cluster 命令列表。	执行用于配置和管理群集的选定 Oracle Solaris Cluster 命令，包括所有 Oracle Solaris Cluster 命令的以下子命令： ■ <code>list</code> ■ <code>show</code> ■ <code>status</code> <code>scha_control</code> <code>scha_resource_get</code> <code>scha_resource_setstatus</code> <code>scha_resourcegroup_get</code> <code>scha_resourcetype_get</code>
基本 Oracle Solaris 用户	此现有的 Oracle Solaris 权限配置文件包含 Oracle Solaris 授权以及以下授权：	<code>solaris.cluster.read</code> – 执行 Oracle Solaris Cluster 命令的列出、显示和其他读取操作以及访问 Oracle Solaris Cluster Manager 浏览器界面。
群集操作	此权限配置文件专用于 Oracle Solaris Cluster 软件，并包含以下授权：	<code>solaris.cluster.read</code> – 执行列出、显示、导出、状态和其他读取操作以及访问 Oracle Solaris Cluster Manager 浏览器界面。
系统管理员	此现有的 Oracle Solaris 权限配置文件包含的授权与群集管理配置文件中包含的授权相同。	<code>solaris.cluster.admin</code> – 更改群集对象的状态。
群集管理	该权限配置文件包含的授权与群集操作配置文件中包含的授权相同。此外，还包含 <code>solaris.cluster.modify</code> 授权。	执行群集操作角色身份所能执行的所有操作，以及更改群集对象的属性。

使用 Oracle Solaris Cluster 管理权限配置文件创建和分配 RBAC 角色

使用此任务通过 Oracle Solaris Cluster 管理权限配置文件创建一个新的 RBAC 角色，并向用户分配此新角色。

▼ 如何从命令行创建角色

1. 选择创建角色的方法：

- 对于本地范围内的角色，请使用 `roleadd` 命令指定新的本地角色及其属性。有关更多信息，请参见 [roleadd\(1M\)](#) 手册页。
- 另外，对于本地范围内的角色，也可以编辑 `user_attr` 文件来添加 `type=role` 的用户。有关更多信息，请参见 [user_attr\(4\)](#) 手册页。
只有在紧急情况下才能使用此方法。
- 对于名称服务中的角色，请使用 `roleadd` 和 `rolemod` 命令来指定新角色及其属性。有关更多信息，请参见 [roleadd\(1M\)](#) 和 [rolemod\(1M\)](#) 手册页。

此命令需要由能够创建其他角色的 `root` 角色来执行才能通过验证。可以将 `roleadd` 命令应用于所有名称服务。此命令将作为 Solaris Management Console 服务器的客户机运行。

2. 启动和停止名称服务缓存守护进程。

重新启动名称服务缓存守护进程后，新角色才能生效。以 `root` 用户身份键入以下文本：

```
# /etc/init.d/nscd stop
# /etc/init.d/nscd start
```

例 13 使用 `smrole` 命令创建定制操作员角色

以下命令序列演示了如何使用 `smrole` 命令创建角色。在本示例中创建了操作员角色的新版本，此操作员角色分配有标准 "Operator"（操作员）权限配置文件和 "Media Restore"（介质恢复）权限配置文件。

```
% su primaryadmin
# /usr/sadm/bin/smrole add -H myHost -- -c "Custom Operator" -n oper2 -a johnDoe \
-d /export/home/oper2 -F "Backup/Restore Operator" -p "Operator" -p "Media Restore"
```

Authenticating as user: primaryadmin

Type `/?` for help, pressing `<enter>` accepts the default denoted by `[ ]`
Please enter a string value for: password :: `<键入 primaryadmin 密码>`

Loading Tool: com.sun.admin.usermgr.cli.role.UserMgrRoleCli from myHost

```
Login to myHost as user primaryadmin was successful.
Download of com.sun.admin.usermgr.cli.role.UserMgrRoleCli from myHost was successful.
```

```
Type /? for help, pressing <enter> accepts the default denoted by [ ]
Please enter a string value for: password :: <键入 oper2 密码>
```

```
# /etc/init.d/nscd stop
# /etc/init.d/nscd start
```

要查看新创建的角色（以及任何其他角色），请将 `smrole` 命令与 `list` 选项配合使用，如下所示：

```
# /usr/sadm/bin/smrole list --
```

```
Authenticating as user: primaryadmin
```

```
Type /? for help, pressing <enter> accepts the default denoted by [ ]
Please enter a string value for: password :: <键入 primaryadmin 密码>
```

```
Loading Tool: com.sun.admin.usermgr.cli.role.UserMgrRoleCli from myHost
Login to myHost as user primaryadmin was successful.
Download of com.sun.admin.usermgr.cli.role.UserMgrRoleCli from myHost was successful.
root                0                Super-User
primaryadmin        100             Most powerful role
sysadmin            101             Performs non-security admin tasks
oper2               102             Custom Operator
```

修改用户的 RBAC 属性

您可以使用用户帐户工具或命令行修改用户的 RBAC 属性。要修改用户的 RBAC 属性，请参见[如何从命令行修改用户的 RBAC 属性 \[63\]](#)。请选择以下过程之一。

- [如何使用用户帐户工具修改用户的 RBAC 属性 \[62\]](#)
- [如何从命令行修改用户的 RBAC 属性 \[63\]](#)

▼ 如何使用用户帐户工具修改用户的 RBAC 属性

开始之前 要修改用户的属性，必须以 `root` 用户身份或承担分配有 "System Administrator"（系统管理员）权限配置文件的角色来运行 "User Tool Collection"（用户工具集）。

1. 启动用户帐户工具。

要运行用户帐户工具，请按照《[在 Oracle Solaris 11.3 中确保用户和进程的安全](#)》中的“[使用所指定的管理权限](#)”中所述，启动 Solaris Management Console。打开 "User Tool Collection"（用户工具集），并单击 "User Accounts"（用户帐户）图标。

用户帐户工具启动后，现有用户帐户的图标会显示在视图窗格中。

2. 单击要更改的 "User Account" (用户帐户) 图标，然后从 "Action" (动作) 菜单中选择 "Properties" (属性) (或双击用户帐户图标)。
3. 单击对话框中待更改属性的相应选项卡，如下所示：
 - 要更改分配给用户的角色，请单击 "Roles" (角色) 选项卡，并将要更改的角色分配移到以下相应列中："Available Roles" (可用角色) 或 "Assigned Roles" (指定的角色)。
 - 要更改分配给用户的权限配置文件，请单击 "Rights" (权限) 选项卡，并将其移到以下相应列中："Available Rights" (可用的权限) 或 "Assigned Rights" (指定的权限)。

注 - 请不要将权限配置文件直接分配给用户。推荐的方法是要求用户必须通过承担角色来执行特权应用程序。此策略可防止用户滥用特权。

▼ 如何从命令行修改用户的 RBAC 属性

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 选择适当的命令：
 - 对于在本地范围内或在 LDAP 系统信息库中定义的用户，要更改分配给该用户的用户属性，请使用 `usermod` 命令。有关更多信息，请参见 [usermod\(1M\)](#) 手册页。
 - 要更改分配给本地范围内定义的用户授权、角色或权限配置文件，还可以编辑 `user_attr` 文件。
只有在紧急情况下才能使用此方法。
 - 要管理本地的或者某个名称服务（例如 LDAP 系统信息库）中的角色，请使用 `roleadd` 或 `rolemo` 命令。有关更多信息，请参见 [roleadd\(1M\)](#) 或 [rolemo\(1M\)](#) 手册页。

这些命令需要由能够更改用户配置文件的 `root` 角色来执行才能通过验证。可以将这些命令应用于所有名称服务。请参见《[在 Oracle Solaris 11.3 中管理用户帐户和用户环境](#)》中的“用于管理用户、角色和组的命令”。

无法修改 Oracle Solaris 11 附带的 "Forced Privilege" (强制特权) 和 "Stop Rights" (停止权限) 配置文件。

关闭和引导群集

本章介绍了关闭和引导全局群集、区域群集及单个群集节点的相关信息和具体过程。

- “关闭和引导群集概述” [65]
- “关闭和引导群集中的单个节点” [79]
- “修复已满的 /var 文件系统” [91]

有关本章中相关过程的概括性说明，请参见[如何以非群集模式引导节点 \[89\]](#)和表 4 “任务列表：关闭并引导节点”。

关闭和引导群集概述

Oracle Solaris Cluster `cluster shutdown` 命令以有序方式停止全局群集服务并完全关闭整个全局群集。您可以在移动全局群集的位置时使用 `cluster shutdown` 命令，或者在应用程序错误导致数据损坏时关闭全局群集。`clzonecluster halt` 命令停止在特定节点上运行的区域群集或所有已配置节点上的整个区域群集。（还可以在区域群集内使用 `cluster shutdown` 命令。）有关更多信息，请参见 [cluster\(1CL\)](#) 手册页。

在本章的操作过程中，`phys-schost#` 表示全局群集提示符。`clzonecluster` 交互式 shell 提示符为 `clzc:schost>`。

注 - 使用 `cluster shutdown` 命令可确保正确关闭整个全局群集。Oracle Solaris shutdown 命令与 `clnode evacuate` 命令一起使用可关闭单个节点。有关更多信息，请参见[如何关闭群集 \[66\]](#)、“关闭和引导群集中的单个节点” [79]或 [clnode\(1CL\)](#) 手册页。

还可以使用 Oracle Solaris Cluster Manager 浏览器界面清除节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

`cluster shutdown` 和 `clzonecluster halt` 命令可通过执行下列操作分别停止全局群集或区域群集中的所有节点：

1. 使所有正在运行的资源组脱机。
2. 卸载全局群集或区域群集的所有群集文件系统。
3. `cluster shutdown` 命令用于关闭全局群集或区域群集上的活动设备服务。

4. `cluster shutdown` 命令运行 `init 0`，并使群集上的所有节点均显示 `penBoot PROM ok` 提示符（在基于 SPARC 的系统上）或在 GRUB 菜单中显示消息 "Press any key to continue"（在基于 x86 的系统上）。有关基于 GRUB 的引导的更多信息，请参见《引导和关闭 Oracle Solaris 11.3 系统》中的“引导系统”。`clzonecluster halt` 命令执行 `zoneadm -z zone-cluster-name halt` 命令来停止（但不是关闭）区域群集的区域。

注 - 根据需要，您可以在非群集模式下引导节点，这样，节点便不是群集成员。非群集模式在安装群集软件或执行某些管理过程时很有用。有关更多信息，请参见[如何以非群集模式引导节点 \[89\]](#)。

表 3 任务列表：关闭和引导群集

任务	说明
停止群集。	如何关闭群集 [66]
通过引导所有节点来启动群集。节点必须具有到群集互连的有效连接才能获得群集成员的身份。	如何引导群集 [68]
重新引导群集。	如何重新引导群集 [72]

▼ 如何关闭群集

您可以关闭全局群集、一个区域群集或所有区域群集。



注意 - 不要在群集控制台上使用 `send brk` 来关闭全局群集节点或区域群集节点。群集内部不支持该命令。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- （仅限 x86）如果您的全局群集或区域群集正在运行 Oracle Real Application Clusters (RAC)，请关闭您要关闭的群集中的所有数据库实例。
有关关闭过程，请参阅 Oracle RAC 产品文档。
- 在群集中的任一节点上承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
从全局群集的节点执行此过程中的所有步骤。
- 关闭全局群集、一个区域群集或所有区域群集。

- 关闭全局群集。此操作还会关闭所有区域群集。

```
phys-schost# cluster shutdown -g0 -y
```

- 关闭特定的区域群集。

```
phys-schost# clzonecluster halt zone-cluster-name
```

- 关闭所有区域群集。

```
phys-schost# clzonecluster halt +
```

还可以在某个区域群集内使用 `cluster shutdown` 命令来关闭该区域群集。

4. 确认全局群集或区域群集上的所有节点都显示 `ok` 提示符（在基于 SPARC 的系统上）或显示 GRUB 菜单（在基于 x86 的系统上）。

除非所有节点均显示 `ok` 提示符（在基于 SPARC 的系统上）或均处于引导子系统中（在基于 x86 的系统上），否则请勿关闭任何节点的电源。

- 从仍然在群集中正常运行的某个全局群集节点上检查其他一个或多个全局群集节点的状态。

```
phys-schost# cluster status -t node
```

- 使用 `status` 子命令检验该区域群集是否已关闭。

```
phys-schost# clzonecluster status
```

5. 如有必要，关闭全局群集节点的电源。

例 14 关闭区域群集

以下示例关闭了一个名为 `sczone` 的区域群集。

```
phys-schost# clzonecluster halt SCZONE
Waiting for zone halt commands to complete on all the nodes of the zone cluster "sczone"...
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 2 of cluster 'sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 4 of cluster 'sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster 'sczone' died.
phys-schost#
```

例 15 SPARC: 关闭全局群集

以下示例显示了当系统停止正常的全局群集操作并关闭所有节点以显示 `ok` 提示符时控制台的输出。`-g 0` 选项表示将关闭宽限期设置为零，`-y` 选项表示在接收到要求确认的问题时自动回答 `yes`。全局群集中其他节点的控制台上也会显示关闭消息。

```
phys-schost# cluster shutdown -g0 -y
Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:
WARNING: CMM monitoring disabled.
phys-schost-1#
```

```
INIT: New run level: 0
The system is coming down. Please wait.
System services are now being stopped.
/etc/rc0.d/K05initrgm: Calling clnode evacuate
The system is down.
syncing file systems... done
Program terminated
ok
```

例 16 x86: 关闭全局群集

以下示例显示了当系统停止正常的全局群集操作并关闭所有节点时控制台的输出。在该示例中，没有在所有节点上均显示 ok 提示符。-g 0 选项表示将关闭宽限期设置为零，-y 选项表示在接收到要求确认的问题时自动回答 yes。全局群集中其他节点的控制台上也会显示关闭消息。

```
phys-schost# cluster shutdown -g0 -y
May  2 10:32:57 phys-schost-1 cl_runtime:
WARNING: CMM: Monitoring disabled.
root@phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
System services are now being stopped.
/etc/rc0.d/K05initrgm: Calling clnode evacuate
failfasts already disabled on node 1
Print services already stopped.
May  2 10:33:13 phys-schost-1 syslogd: going down on signal 15
The system is down.
syncing file systems... done
Type any key to continue
```

另请参见 要重新启动已经关闭的全局群集或区域群集，请参见[如何引导群集 \[68\]](#)。

▼ 如何引导群集

该过程说明了如何启动节点已经关闭的全局群集或区域群集。对于全局群集节点，系统显示 ok 提示符（在 SPARC 系统上）或消息 "Press any key to continue"（在基于 GRUB 的 x86 系统上）。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 要创建区域群集，请按照《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[创建和配置区域群集](#)”中的说明操作或使用 Oracle Solaris Cluster Manager 浏览器界面创建区域群集。

1. 将每个节点都引导到群集模式下。
从全局群集的节点执行此过程中的所有步骤。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot
```

- 在基于 x86 的系统上，运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

有关基于 GRUB 的引导的更多信息，请参见《[引导和关闭 Oracle Solaris 11.3 系统](#)》中的“引导系统”。

注 - 节点必须具有到群集互连的有效连接才能获得群集成员的身份。

- 如果您具有一个区域群集，便可以引导整个区域群集。

```
phys-schost# clzonecluster boot zone-cluster-name
```

- 如果您具有多个区域群集，便可以引导所有区域群集。使用加号 (+)，而不是 *zone-cluster-name*。

2. 验证引导节点时未发生错误，而且节点现在处于联机状态。

`cluster status` 命令报告全局群集节点的状态。

```
phys-schost# cluster status -t node
```

当您从全局群集节点运行 `clzonecluster status` 状态命令时，该命令将报告区域群集节点的状态。

```
phys-schost# clzonecluster status
```

注 - 如果节点的 `/var` 文件系统已满，可能无法在该节点上重新启动 Oracle Solaris Cluster。如果出现该问题，请参见[如何修复已满的 /var 文件系统 \[91\]](#)。有关更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

例 17 SPARC: 引导全局群集

以下示例显示了将节点 `phys-schost-1` 引导到全局群集后的控制台输出。全局群集中其他节点的控制台上会显示类似的消息。当某个区域群集的自动引导属性设置为 `true` 时，系统将在引导该计算机上的全局群集节点之后自动引导该区域群集节点。

当全局群集节点重新引导时，该计算机上的所有区域群集节点都将停止。在该全局群集节点重新启动之后，将引导同一计算机上自动引导属性设置为 `true` 的任何区域群集节点。

```

ok boot
Rebooting with command: boot
...
Hostname: phys-schost-1
Booting as part of a cluster
NOTICE: Node phys-schost-1 with votecount = 1 added.
NOTICE: Node phys-schost-2 with votecount = 1 added.
NOTICE: Node phys-schost-3 with votecount = 1 added.
...
NOTICE: Node phys-schost-1: attempting to join cluster
...
NOTICE: Node phys-schost-2 (incarnation # 937690106) has become reachable.
NOTICE: Node phys-schost-3 (incarnation # 937690290) has become reachable.
NOTICE: cluster has reached quorum.
NOTICE: node phys-schost-1 is up; new incarnation number = 937846227.
NOTICE: node phys-schost-2 is up; new incarnation number = 937690106.
NOTICE: node phys-schost-3 is up; new incarnation number = 937690290.
NOTICE: Cluster members: phys-schost-1 phys-schost-2 phys-schost-3.
...

```

例 18 x86: 引导群集

以下示例显示了将节点 phys-schost-1 引导到群集后的控制台输出。群集中其他节点的控制台上会显示类似的消息。

```

ATI RAGE SDRAM BIOS P/N GR-xlint.007-4.330
* BIOS Lan-Console 2.0
Copyright (C) 1999-2001 Intel Corporation
MAC ADDR: 00 02 47 31 38 3C
AMIBIOS (C)1985-2002 American Megatrends Inc.,
Copyright 1996-2002 Intel Corporation
SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064
2 X Intel(R) Pentium(R) III CPU family      1400MHz
Testing system memory, memory size=2048MB
2048MB Extended Memory Passed
512K L2 Cache SRAM Passed
ATAPI CD-ROM SAMSUNG CD-ROM SN-124

Press <F2> to enter SETUP, <F12> Network

Adaptec AIC-7899 SCSI BIOS v2.57S4
(c) 2000 Adaptec, Inc. All Rights Reserved.
Press <Ctrl><A> for SCSISelect(TM) Utility!

Ch B,  SCSI ID: 0 SEAGATE ST336605LC      160
SCSI ID: 1 SEAGATE ST336605LC      160
SCSI ID: 6 ESG-SHV SCA HSBP M18      ASYN
Ch A,  SCSI ID: 2 SUN StorEdge 3310      160
SCSI ID: 3 SUN StorEdge 3310      160

AMIBIOS (C)1985-2002 American Megatrends Inc.,

```

Copyright 1996-2002 Intel Corporation
SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064

2 X Intel(R) Pentium(R) III CPU family 1400MHz
Testing system memory, memory size=2048MB
2048MB Extended Memory Passed
512K L2 Cache SRAM Passed
ATAPI CD-ROM SAMSUNG CD-ROM SN-124

SunOS - Intel Platform Edition Primary Boot Subsystem, vsn 2.0

Current Disk Partition Information

Part#	Status	Type	Start	Length
1	Active	X86 BOOT	2428	21852
2		SOLARIS	24280	71662420
3		<unused>		
4		<unused>		

Please select the partition you wish to boot: * *

Solaris DCB

loading /solaris/boot.bin

SunOS Secondary Boot version 3.00

Solaris Intel Platform Edition Booting System

Autobooting from bootpath: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/
pci8086,341a@7,1/sd@0,0:a

If the system hardware has changed, or to boot from a different
device, interrupt the autoboot process by pressing ESC.

Press ESCape to interrupt autoboot in 2 seconds.

Initializing system

Please wait...

Warning: Resource Conflict - both devices are added

NON-ACPI device: ISY0050

Port: 3F0-3F5, 3F7; IRQ: 6; DMA: 2

ACPI device: ISY0050

Port: 3F2-3F3, 3F4-3F5, 3F7; IRQ: 6; DMA: 2

<<< Current Boot Parameters >>>

Boot path: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/pci8086,341a@7,1/
sd@0,0:a

Boot args:

Type	b [file-name] [boot-flags] <ENTER>	to boot with options
or	i <ENTER>	to enter boot interpreter
or	<ENTER>	to boot with defaults

```
<<< timeout in 5 seconds >>>

Select (b)oot or (i)nterpreter:
Size: 275683 + 22092 + 150244 Bytes
/platform/i86pc/kernel/unix loaded - 0xac000 bytes used
SunOS Release 5.9 Version Generic_112234-07 32-bit
Copyright 1983-2003 Sun Microsystems, Inc. All rights reserved.
Use is subject to license terms.
configuring IPv4 interfaces: e1000g2.
Hostname: phys-schost-1
Booting as part of a cluster
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) with votecount = 1 added.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) with votecount = 1 added.
NOTICE: CMM: Quorum device 1 (/dev/did/rdisk/dls2) added; votecount = 1, bitmask
of nodes with configured paths = 0x3.
NOTICE: clcomm: Adapter e1000g3 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 online
NOTICE: clcomm: Adapter e1000g0 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being constructed
NOTICE: CMM: Node phys-schost-1: attempting to join cluster.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being initiated
NOTICE: CMM: Quorum device /dev/did/rdisk/dls2: owner set to node 1.
NOTICE: CMM: Cluster has reached quorum.
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) is up; new incarnation number = 1068496374.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) is up; new incarnation number = 1068496374.
NOTICE: CMM: Cluster members: phys-schost-1 phys-schost-2.
NOTICE: CMM: node reconfiguration #1 completed.
NOTICE: CMM: Node phys-schost-1: joined cluster.
```

▼ 如何重新引导群集

要关闭全局群集，请运行 `cluster shutdown` 命令，然后在每个节点上使用 `boot` 命令引导全局群集。要关闭区域群集，请使用 `clzonecluster halt` 命令，然后使用 `clzonecluster boot` 命令引导该区域群集。您还可以使用 `clzonecluster reboot` 命令。有关更多信息，请参见 [cluster\(1CL\)](#)、[boot\(1M\)](#) 和 [clzonecluster\(1CL\)](#) 手册页。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 如果您的群集正在运行 Oracle RAC，请关闭您要关闭的群集中的所有数据库实例。有关关闭过程，请参阅 Oracle RAC 产品文档。
2. 在群集中的任一节点上承担可提供 `solaris.cluster.admin` RBAC 授权的角色。

从全局群集的节点执行此过程中的所有步骤。

3. 关闭群集。

■ 关闭全局群集。

```
phys-schost# cluster shutdown -g0 -y
```

■ 如果您具有区域群集，请从全局群集节点关闭该区域群集。

```
phys-schost# clzonecluster halt zone-cluster-name
```

将关闭所有节点。还可以在区域群集内使用 `cluster shutdown` 命令来关闭该区域群集。

注 - 节点必须具有到群集互连的有效连接才能获得群集成员的身份。

4. 引导每个节点。

除非在两次关闭操作之间更改了配置，否则，节点的引导顺序无关紧要。如果在两次关闭操作之间进行了配置更改，则首先启动具有最新配置的节点。

■ 对于基于 SPARC 的系统上的全局群集节点，请运行以下命令。

```
ok boot
```

■ 对于基于 x86 的系统上的全局群集节点，请运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris OS 条目，然后按 Enter 键。

有关基于 GRUB 的引导的更多信息，请参见《[引导和关闭 Oracle Solaris 11.3 系统](#)》中的“引导系统”。

■ 对于区域群集，请在全局群集的单个节点上键入以下命令来引导该区域群集。

```
phys-schost# clzonecluster boot zone-cluster-name
```

注 - 节点必须具有到群集互连的有效连接才能获得群集成员的身份。

当群集组件被激活时，所引导的节点的控制台上将显示消息。

5. 验证引导节点时未发生错误，而且节点现在处于联机状态。

■ `clnode status` 命令报告全局群集中节点的状态。

```
phys-schost# clnode status
```

■ 在全局群集节点上运行 `clzonecluster status` 命令将报告区域群集节点的状态。

```
phys-schost# clzonecluster status
```

还可以在区域群集内运行 `cluster status` 命令来查看节点的状态。

注 - 如果节点的 /var 文件系统已满，可能无法在该节点上重新启动 Oracle Solaris Cluster。如果出现该问题，请参见[如何修复已满的 /var 文件系统 \[91\]](#)。

例 19 重新引导区域群集

以下示例显示了如何停止和引导一个名为 *sparse-sczone* 的区域群集。您还可以使用 `clzonecluster reboot` 命令。

```
phys-schost# clzonecluster halt sparse-sczone
Waiting for zone halt commands to complete on all the nodes of the zone cluster "sparse-sczone"...
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 4 of cluster 'sparse-sczone' died.
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 2 of cluster 'sparse-sczone' died.
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster 'sparse-sczone' died.
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone' died.
phys-schost#
phys-schost# clzonecluster boot sparse-sczone
Waiting for zone boot commands to complete on all the nodes of the zone cluster "sparse-sczone"...
phys-schost# Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster 'sparse-sczone' joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 2 of cluster 'sparse-sczone' joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone' joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 4 of cluster 'sparse-sczone' joined.

phys-schost#
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---

Name           Node Name   Zone HostName   Status   Zone Status
-----
sparse-sczone  schost-1    sczone-1        Online   Running
                schost-2    sczone-2        Online   Running
                schost-3    sczone-3        Online   Running
                schost-4    sczone-4        Online   Running
phys-schost#
```

例 20 SPARC: 重新引导全局群集

以下示例显示了当停止正常的全局群集操作、所有节点均关闭以显示 ok 提示符并且全局群集重新启动时的控制台输出。-g 0 选项表示将宽限期设置为零，-y 选项表示在接收到要求确认的问题时自动回答 yes。全局群集中其他节点的控制台上也会显示关闭消息。

```
phys-schost# cluster shutdown -g0 -y
Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:
WARNING: CMM monitoring disabled.
phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
...
The system is down.
syncing file systems... done
Program terminated
ok boot
Rebooting with command: boot
...
Hostname: phys-schost-1
Booting as part of a cluster
...
NOTICE: Node phys-schost-1: attempting to join cluster
...
NOTICE: Node phys-schost-2 (incarnation # 937690106) has become reachable.
NOTICE: Node phys-schost-3 (incarnation # 937690290) has become reachable.
NOTICE: cluster has reached quorum.
...
NOTICE: Cluster members: phys-schost-1 phys-schost-2 phys-schost-3.
...
NOTICE: Node phys-schost-1: joined cluster
...
The system is coming up. Please wait.
checking ufs filesystems
...
reservation program successfully exiting
Print services started.
volume management starting.
The system is ready.
phys-schost-1 console login:
NOTICE: Node phys-schost-1: joined cluster
...
The system is coming up. Please wait.
checking ufs filesystems
...
reservation program successfully exiting
Print services started.
volume management starting.
The system is ready.
phys-schost-1 console login:
```

例 21 x86: 重新引导群集

以下示例显示了当系统停止正常群集操作、关闭所有节点，并重新启动群集时，控制台的输出。-g 0 选项表示将宽限期设置为零，-y 表示在接收到要求确认的问题时自动回答 yes。群集中其他节点的控制台上也会显示关闭消息。

```
# cluster shutdown -g0 -y
May  2 10:32:57 phys-schost-1 cl_runtime:
WARNING: CMM: Monitoring disabled.
root@phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
System services are now being stopped.
/etc/rc0.d/K05initrgm: Calling clnode evacuate
failfasts already disabled on node 1
Print services already stopped.
May  2 10:33:13 phys-schost-1 syslogd: going down on signal 15
The system is down.
syncing file systems... done
Type any key to continue

ATI RAGE SDRAM BIOS P/N GR-xlint.007-4.330
*                               BIOS Lan-Console 2.0
Copyright (C) 1999-2001 Intel Corporation
MAC ADDR: 00 02 47 31 38 3C
AMIBIOS (C)1985-2002 American Megatrends Inc.,
Copyright 1996-2002 Intel Corporation
SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064
2 X Intel(R) Pentium(R) III CPU family      1400MHz
Testing system memory, memory size=2048MB
2048MB Extended Memory Passed
512K L2 Cache SRAM Passed
ATAPI CD-ROM SAMSUNG CD-ROM SN-124

Press <F2> to enter SETUP, <F12> Network

Adaptec AIC-7899 SCSI BIOS v2.57S4
(c) 2000 Adaptec, Inc. All Rights Reserved.
Press <Ctrl><A> for SCSISelect(TM) Utility!

Ch B,  SCSI ID: 0 SEAGATE ST336605LC      160
SCSI ID: 1 SEAGATE ST336605LC      160
SCSI ID: 6 ESG-SHV SCA HSBP M18      ASYN
Ch A,  SCSI ID: 2 SUN StorEdge 3310      160
SCSI ID: 3 SUN StorEdge 3310      160

AMIBIOS (C)1985-2002 American Megatrends Inc.,
Copyright 1996-2002 Intel Corporation
SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064
```

2 X Intel(R) Pentium(R) III CPU family 1400MHz
 Testing system memory, memory size=2048MB
 2048MB Extended Memory Passed
 512K L2 Cache SRAM Passed
 ATAPI CD-ROM SAMSUNG CD-ROM SN-124

SunOS - Intel Platform Edition Primary Boot Subsystem, vsn 2.0

Current Disk Partition Information

Part#	Status	Type	Start	Length
1	Active	X86 BOOT	2428	21852
2		SOLARIS	24280	71662420
3		<unused>		
4		<unused>		

Please select the partition you wish to boot: * *

Solaris DCB

loading /solaris/boot.bin

SunOS Secondary Boot version 3.00

Solaris Intel Platform Edition Booting System

Autobooting from bootpath: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/
 pci8086,341a@7,1/sd@0,0:a

If the system hardware has changed, or to boot from a different
 device, interrupt the autoboot process by pressing ESC.

Press ESCape to interrupt autoboot in 2 seconds.

Initializing system

Please wait...

Warning: Resource Conflict - both devices are added

NON-ACPI device: ISY0050

Port: 3F0-3F5, 3F7; IRQ: 6; DMA: 2

ACPI device: ISY0050

Port: 3F2-3F3, 3F4-3F5, 3F7; IRQ: 6; DMA: 2

<<< Current Boot Parameters >>>

Boot path: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/pci8086,341a@7,1/
 sd@0,0:a

Boot args:

Type b [file-name] [boot-flags] <ENTER> to boot with options
 or i <ENTER> to enter boot interpreter
 or <ENTER> to boot with defaults

<<< timeout in 5 seconds >>>

Select (b)oot or (i)nterpreter: **b**

Size: 275683 + 22092 + 150244 Bytes

```
/platform/i86pc/kernel/unix loaded - 0xac000 bytes used
SunOS Release 5.9 Version Generic_112234-07 32-bit
Copyright 1983-2003 Sun Microsystems, Inc. All rights reserved.
Use is subject to license terms.
configuring IPv4 interfaces: e1000g2.
Hostname: phys-schost-1
Booting as part of a cluster
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) with votecount = 1 added.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) with votecount = 1 added.
NOTICE: CMM: Quorum device 1 (/dev/did/rdisk/dls2) added; votecount = 1, bitmask
of nodes with configured paths = 0x3.
NOTICE: clcomm: Adapter e1000g3 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 online
NOTICE: clcomm: Adapter e1000g0 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being constructed
NOTICE: CMM: Node phys-schost-1: attempting to join cluster.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being initiated
NOTICE: CMM: Quorum device /dev/did/rdisk/dls2: owner set to node 1.
NOTICE: CMM: Cluster has reached quorum.
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) is up; new incarnation number = 1068496374.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) is up; new incarnation number = 1068496374.
NOTICE: CMM: Cluster members: phys-schost-1 phys-schost-2.
NOTICE: CMM: node reconfiguration #1 completed.
NOTICE: CMM: Node phys-schost-1: joined cluster.
WARNING: mod_installdrv: no major number for rsmrdt
ip: joining multicasts failed (18) on clprivnet0 - will use link layer
broadcasts for multicast
The system is coming up. Please wait.
checking ufs filesystems
/dev/rdisk/clt0d0s5: is clean.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 online
NIS domain name is dev.eng.mycompany.com
starting rpc services: rpcbind keyserver ypbind done.
Setting netmask of e1000g2 to 192.168.255.0
Setting netmask of e1000g3 to 192.168.255.128
Setting netmask of e1000g0 to 192.168.255.128
Setting netmask of clprivnet0 to 192.168.255.0
Setting default IPv4 interface for multicast: add net 224.0/4: gateway phys-schost-1
syslog service starting.
obtaining access to all attached disks

*****
*
* The X-server can not be started on display :0...
*
*****
volume management starting.
Starting Fault Injection Server...
The system is ready.

phys-schost-1 console login:
```

关闭和引导群集中的单个节点

您可以关闭全局群集节点或区域群集节点。本节说明了如何关闭全局群集节点和区域群集节点。

要关闭全局群集节点，请将 `clnode evacuate` 命令与 Oracle Solaris `shutdown` 命令配合使用。仅当关闭整个全局群集时，才能使用 `cluster shutdown` 命令。

对于区域群集节点，请在全局群集中使用 `clzonecluster halt` 命令以关闭单个区域群集节点或整个区域群集。还可以使用 `clnode evacuate` 和 `shutdown` 命令来关闭区域群集节点。

有关更多信息，请参见 [clnode\(1CL\)](#)、[shutdown\(1M\)](#) 和 [clzonecluster\(1CL\)](#) 手册页。

在本章的操作过程中，`phys-schost#` 表示全局群集提示符。`clzonecluster` 交互式 shell 提示符为 `clzc:schost>`。

表 4 任务列表：关闭并引导节点

任务	工具	说明
停止节点。	对于全局群集节点，请使用 <code>clnode evacuate</code> 和 <code>shutdown</code> 命令。 对于区域群集节点，请使用 <code>clzonecluster halt</code> 命令。	如何关闭节点 [80]
启动节点。 节点必须具有到群集互连的有效连接才能获得群集成员的身份。	对于全局群集节点，请使用 <code>boot</code> 或 <code>b</code> 命令。 对于区域群集节点，请使用 <code>clzonecluster boot</code> 命令。	如何引导节点 [82]
停止并重新启动（重新引导）群集中的节点。 节点必须具有到群集互连的有效连接才能获得群集成员的身份。	对于全局群集节点，请使用 <code>clnode evacuate</code> 和 <code>shutdown</code> 命令，然后使用 <code>boot</code> 或 <code>b</code> 。 对于区域群集节点，请使用 <code>clzonecluster reboot</code> 命令。	如何重新引导节点 [86]
引导一个节点，使该节点不成为群集成员。	对于全局群集节点，请使用 <code>clnode evacuate</code> 和 <code>shutdown</code> 命令，然后使用 <code>boot -x</code> （在 SPARC 上）或 GRUB 菜单项编辑（在 x86 上）。 如果底层的全局群集是以非群集模式引导的，则区域群集节点也自动以非群集模式引导。	如何以非群集模式引导节点 [89]

▼ 如何关闭节点

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。



注意 - 不要在群集控制台上使用 send brk 来关闭全局群集或区域群集上的节点。群集内部不支持该命令。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面清除全局区域节点以及将所有资源组和设备组切换到下一个首选节点。您还可以关闭区域群集节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 如果您的群集正在运行 Oracle RAC，请关闭您要关闭的群集中的所有数据库实例。有关关闭过程，请参阅 Oracle RAC 产品文档。
2. 在要关闭的群集节点上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。从全局群集的节点执行此过程中的所有步骤。
3. 如果您想要停止特定的区域群集成员，请跳过步骤 4-6，然后从全局群集节点执行以下命令：

```
phys-schost# clzonecluster halt -n physical-name zone-cluster-name
```

如果您指定了特定的区域群集节点，将仅停止该节点。默认情况下，halt 命令停止所有节点上的区域群集。

4. 将所有资源组、资源和设备组从要关闭的节点切换到其他全局群集成员。在要关闭的全局群集节点上，键入以下命令。clnode evacuate 命令可将指定节点上的所有资源组和设备组切换到下一个首选节点。（还可以在区域群集节点内运行 clnode evacuate。）

```
phys-schost# clnode evacuate node
```

node 指定从中切换资源组和设备组的节点。

5. 关闭该节点。在要关闭的全局群集节点上执行 shutdownn 命令。

```
phys-schost# shutdown -g0 -y -i0
```

检验该全局群集节点是否显示 ok 提示符（在基于 SPARC 的系统上）或在 GRUB 菜单中显示消息 "Press any key to continue"（在基于 x86 的系统上）。

6. 如有必要，关闭节点电源。

例 22 SPARC: 关闭全局群集节点

以下示例显示了节点 `phys-schost-1` 关闭时的控制台输出。`-g0` 选项表示将宽限期设置为零，`-y` 选项表示在接收到要求确认的问题时自动回答 `yes`。全局群集中其他节点的控制台上也显示此节点的关闭消息。

```
phys-schost# clnode evacuate phys-schost-1
phys-schost# shutdown -g0 -y -i0
Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:
WARNING: CMM monitoring disabled.
phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
Notice: rgmd is being stopped.
Notice: rpc.pmfd is being stopped.
Notice: rpc.fed is being stopped.
umount: /global/.devices/node@1 busy
umount: /global/phys-schost-1 busy
The system is down.
syncing file systems... done
Program terminated
ok
```

例 23 x86: 关闭全局群集节点

以下示例显示了节点 `phys-schost-1` 关闭时的控制台输出。`-g0` 选项表示将宽限期设置为零，`-y` 选项表示在接收到要求确认的问题时自动回答 `yes`。全局群集中其他节点的控制台上也显示此节点的关闭消息。

```
phys-schost# clnode evacuate phys-schost-1
phys-schost# shutdown -g0 -y -i0
Shutdown started.    Wed Mar 10 13:47:32 PST 2004

Changing to init state 0 - please wait
Broadcast Message from root (console) on phys-schost-1 Wed Mar 10 13:47:32...
THE SYSTEM phys-schost-1 IS BEING SHUT DOWN NOW ! ! !
Log off now or risk your files being damaged

phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
System services are now being stopped.
/etc/rc0.d/K05initrgm: Calling clnode evacuate
failfasts disabled on node 1
Print services already stopped.
Mar 10 13:47:44 phys-schost-1 syslogd: going down on signal 15
umount: /global/.devices/node@2 busy
umount: /global/.devices/node@1 busy
The system is down.
syncing file systems... done
```

```
WARNING: CMM: Node being shut down.
Type any key to continue
```

例 24 关闭区域群集节点

以下示例显示了如何使用 `clzonecluster halt` 关闭一个名为 *sparse-sczone* 的区域群集中的节点。（还可以在区域群集节点中运行 `clnode evacuate` 和 `shutdown` 命令。）

```
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---

Name           Node Name   Zone HostName   Status   Zone Status
-----
sparse-sczone  schost-1    sczone-1        Online   Running
                schost-2    sczone-2        Online   Running
                schost-3    sczone-3        Online   Running
                schost-4    sczone-4        Online   Running

phys-schost#
phys-schost# clzonecluster halt -n schost-4 sparse-sczone
Waiting for zone halt commands to complete on all the nodes of the zone cluster "sparse-sczone"...
Sep  5 19:24:00 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone'
died.
phys-host#
phys-host# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---

Name           Node Name   Zone HostName   Status   Zone Status
-----
sparse-sczone  schost-1    sczone-1        Online   Running
                schost-2    sczone-2        Online   Running
                schost-3    sczone-3        Offline  Installed
                schost-4    sczone-4        Online   Running

phys-schost#
```

另请参见 要重新启动已关闭的全局群集节点，请参见[如何引导节点 \[82\]](#)。

▼ 如何引导节点

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面引导区域群集节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

如果要在全局群集或区域群集中关闭或重新引导其他活动节点，对于您要引导的节点，请等到多用户服务器里程碑进入联机状态。否则，该节点将无法从群集中要关闭或重新引导的其他节点那里接管服务。

注 - 节点的启动可能受法定配置的影响。在双节点群集中，必须配置一个法定设备，使群集中的法定总计数为三。每个节点应有一个法定计数，法定设备有一个法定计数。在这种情况下，当第一个节点关闭后，第二个节点仍然具有法定，并且以唯一的群集成员的身份运行。要使第一个节点作为群集节点返回群集，第二个节点必须启动并且正在运行。必须存在所需的群集法定计数（两个）。

如果您在来宾域运行 Oracle Solaris Cluster，则重新引导控制域或 I/O 域会对运行的来宾域产生影响，包括使域停止运行。您应该将工作负荷重新调整到其他节点，并在重新引导控制域或 I/O 域前停止运行 Oracle Solaris Cluster 的来宾域。

重新引导控制域或 I/O 域时，来宾域不接收或发送心跳信号。这会导致发生记忆分裂和群集重新配置。由于控制域或 I/O 域已重新引导，来宾域无法访问任何共享设备。其他群集节点会阻止该来宾域访问共享设备。当控制域或 I/O 域完成重新引导后，来宾域上的 I/O 操作会继续进行，但任何针对共享存储的 I/O 都会导致来宾域出现紧急情况，因为在群集重新配置过程中已将其阻隔在共享磁盘之外。如果来宾使用两个 I/O 域实现冗余，同时一次仅重新引导一个 I/O 域，则可以缓解该问题。

注 - 节点必须具有到群集互连的有效连接才能获得群集成员的身份。

1. 要启动已关闭的全局群集节点或区域群集节点，请引导该节点。
从全局群集的节点执行此过程中的所有步骤。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot
```

- 在基于 x86 的系统上，运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

当群集组件被激活时，所引导的节点的控制台上将显示消息。

- 如果您具有区域群集，可以指定要引导的节点。

```
phys-schost# clzonecluster boot -n node zone-cluster-name
```

2. 确认引导该节点时未发生错误，而且该节点现在处于联机状态。

- 运行 **cluster status** 命令可报告全局群集节点的状态。

```
phys-schost# cluster status -t node
```

- 从全局群集中的节点运行 **clzonecluster status** 命令可报告所有区域群集节点的状态。

```
phys-schost# clzonecluster status
```

当托管区域群集节点的节点以群集模式引导时，区域群集节点只能以群集模式引导。

注 - 如果节点的 /var 文件系统已满，可能无法在该节点上重新启动 Oracle Solaris Cluster。如果出现该问题，请参见[如何修复已满的 /var 文件系统 \[91\]](#)。

例 25 SPARC: 引导全局群集节点

以下示例显示了将节点 phys-schost-1 引导到全局群集后的控制台输出。

```
ok boot
Rebooting with command: boot
...
Hostname: phys-schost-1
Booting as part of a cluster
...
NOTICE: Node phys-schost-1: attempting to join cluster
...
NOTICE: Node phys-schost-1: joined cluster
...
The system is coming up. Please wait.
checking ufs filesystems
...
reservation program successfully exiting
Print services started.
volume management starting.
The system is ready.
phys-schost-1 console login:
```

例 26 x86: 引导群集节点

以下示例显示了将节点 phys-schost-1 引导到群集后的控制台输出。

```
<<< Current Boot Parameters >>>
Boot path: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/pci8086,341a@7,1/sd@0,0:a
Boot args:

Type    b [file-name] [boot-flags] <ENTER>   to boot with options
or      i <ENTER>                           to enter boot interpreter
or      <ENTER>                             to boot with defaults
```

<<< timeout in 5 seconds >>>

```
Select (b)oot or (i)nterpreter: Size: 276915 + 22156 + 150372 Bytes
/platform/i86pc/kernel/unix loaded - 0xac000 bytes used
SunOS Release 5.9 Version on81-feature-patch:08/30/2003 32-bit
Copyright 1983-2003 Sun Microsystems, Inc. All rights reserved.
Use is subject to license terms.
configuring IPv4 interfaces: e1000g2.
Hostname: phys-schost-1
Booting as part of a cluster
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) with votecount = 1 added.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) with votecount = 1 added.
NOTICE: CMM: Quorum device 1 (/dev/did/rdisk/dls2) added; votecount = 1, bitmask
of nodes with configured paths = 0x3.
WARNING: CMM: Initialization for quorum device /dev/did/rdisk/dls2 failed with
error EACCES. Will retry later.
NOTICE: clcomm: Adapter e1000g3 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 online
NOTICE: clcomm: Adapter e1000g0 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being constructed
NOTICE: CMM: Node phys-schost-1: attempting to join cluster.
WARNING: CMM: Reading reservation keys from quorum device /dev/did/rdisk/dls2
failed with error 2.
NOTICE: CMM: Cluster has reached quorum.
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) is up; new incarnation number =
1068503958.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) is up; new incarnation number =
1068496374.
NOTICE: CMM: Cluster members: phys-schost-1 phys-schost-2.
NOTICE: CMM: node reconfiguration #3 completed.
NOTICE: CMM: Node phys-schost-1: joined cluster.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 online
NOTICE: CMM: Retry of initialization for quorum device /dev/did/rdisk/dls2 was
successful.
WARNING: mod_installdr: no major number for rsmrtd
ip: joining multicasts failed (18) on clprivnet0 - will use link layer
broadcasts for multicast
The system is coming up. Please wait.
checking ufs filesystems
/dev/rdisk/clt0d0s5: is clean.
NIS domain name is dev.eng.mycompany.com
starting rpc services: rpcbind keyserver ypbind done.
Setting netmask of e1000g2 to 192.168.255.0
Setting netmask of e1000g3 to 192.168.255.128
Setting netmask of e1000g0 to 192.168.255.128
Setting netmask of clprivnet0 to 192.168.255.0
Setting default IPv4 interface for multicast: add net 224.0/4: gateway phys-schost-1
syslog service starting.
obtaining access to all attached disks
```

```
*****
*
* The X-server can not be started on display :0...
*
*****
volume management starting.
Starting Fault Injection Server...
The system is ready.

phys-schost-1 console login:
```

▼ 如何重新引导节点

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面重新引导区域群集节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。



注意 - 如果用于任一资源的方法超时且无法中止，则仅当资源的 Failover_mode 属性设置为 HARD 时才会重新引导节点。如果 Failover_mode 属性设置为任何其他值，将不会重新引导节点。

要在全局群集或区域群集中关闭或重新引导其他活动节点，对于您要重新引导的节点，请等到多用户服务器里程碑进入联机状态。否则，该节点将无法从群集中要关闭或重新引导的其他节点那里接管服务。

1. 如果全局群集或区域群集节点正在运行 Oracle RAC，请关闭您要关闭的节点上的所有数据库实例。
有关关闭过程，请参阅 Oracle RAC 产品文档。
2. 在要关闭的节点上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
从全局群集的节点执行此过程中的所有步骤。

3. 使用 `clnode evacuate` 和 `shutdown` 命令关闭全局群集节点。

使用在全局群集的某个节点上执行的 `clzonecluster halt` 命令关闭区域群集。（`clnode evacuate` 和 `shutdown` 命令也可用于区域群集。）

对于全局群集，请在节点上键入以下命令将其关闭。`clnode evacuate` 命令可将指定节点上的所有设备组切换到下一个首选节点。此外，该命令还可将所有资源组从指定节点的全局区域切换到位于其他节点的下一个首选全局区域。

注 - 要关闭单个节点，请使用 `shutdown -g0 -y -i6` 命令。要同时关闭多个节点，请使用 `shutdown -g0 -y -i0` 命令停止这些节点。停止所有节点后，在所有节点上使用 `boot` 命令以将它们引导回群集中。

- 在基于 SPARC 的系统上，运行以下命令重新引导单个节点。

```
phys-schost# clnode evacuate node
```

```
phys-schost# shutdown -g0 -y -i6
```

- 在基于 x86 的系统上，运行以下命令重新引导单个节点。

```
phys-schost# clnode evacuate node
```

```
phys-schost# shutdown -g0 -y -i6
```

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

- 指定要关闭和重新引导的区域群集节点。

```
phys-schost# clzonecluster reboot - node zone-cluster-name
```

注 - 节点必须具有到群集互连的有效连接才能获得群集成员的身份。

4. 确认引导该节点时未发生错误，而且该节点现在处于联机状态。

- 确认全局群集节点处于联机状态。

```
phys-schost# cluster status -t node
```

- 确认区域群集节点处于联机状态。

```
phys-schost# clzonecluster status
```

例 27 SPARC: 重新引导全局群集节点

以下示例显示了当节点 `phys-schost-1` 重新引导时的控制台输出。有关该节点的消息（例如关闭和启动通知）出现在全局群集中其他节点的控制台上。

```
phys-schost# clnode evacuate phys-schost-1
phys-schost# shutdown -g0 -y -i6
Shutdown started.   Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:

WARNING: CMM monitoring disabled.
phys-schost-1#
```

```
INIT: New run level: 6
The system is coming down. Please wait.
System services are now being stopped.
Notice: rgmd is being stopped.
Notice: rpc.pmfd is being stopped.
Notice: rpc.fed is being stopped.
umount: /global/.devices/node@1 busy
umount: /global/phys-schost-1 busy
The system is down.
syncing file systems... done
rebooting...
Resetting ...

'''
Sun Ultra 1 SBus (UltraSPARC 143MHz), No Keyboard
OpenBoot 3.11, 128 MB memory installed, Serial #5932401.
Ethernet address 8:8:20:99:ab:77, Host ID: 8899ab77.
...
Rebooting with command: boot
...
Hostname: phys-schost-1
Booting as part of a cluster
...
NOTICE: Node phys-schost-1: attempting to join cluster
...
NOTICE: Node phys-schost-1: joined cluster
...
The system is coming up. Please wait.
The system is ready.
phys-schost-1 console login:
```

例 28 重新引导区域群集节点

以下示例显示了如何重新引导区域群集中的节点。

```
phys-schost# clzonecluster reboot -n schost-4 sparse-sczone
Waiting for zone reboot commands to complete on all the nodes of the zone cluster
"sparse-sczone"...
Sep  5 19:40:59 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster
'sparse-sczone' died.
phys-schost# Sep  5 19:41:27 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster
'sparse-sczone' joined.

phys-schost#
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---
Name           Node Name     Zone HostName  Status   Zone Status
-----
sparse-sczone  schost-1      sczone-1       Online   Running
                schost-2      sczone-2       Online   Running
                schost-3      sczone-3       Online   Running
                schost-4      sczone-4       Online   Running
```



```
phys-schost#
```

▼ 如何以非群集模式引导节点

您可以在非群集模式下引导全局群集节点，这种情况下该节点不会成为群集的成员。当安装群集软件或执行某些管理过程（如更新节点）时，非群集模式很有用。区域群集节点不能处于与底层的全局群集节点的状态不同的引导状态。如果底层的全局群集节点是以非群集模式引导的，则区域群集节点也自动处于非群集模式。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在将要以非群集模式启动的群集上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。

从全局群集的节点执行此过程中的所有步骤。

2. 关闭区域群集节点或全局群集节点。

`clnode evacuate` 命令可将指定节点上的所有设备组切换到下一个首选节点。此外，该命令还可将所有资源组从指定节点上的全局区域切换到位于其他节点的下一个首选全局区域。

- 关闭特定的全局群集节点。

```
phys-schost# clnode evacuate node
```

```
phys-schost# shutdown -g0 -y
```

- 从全局群集节点关闭特定的区域群集节点。

```
phys-schost# clzonecluster halt -n node zone-cluster-name
```

还可以在区域群集内使用 `clnode evacuate` 和 `shutdown` 命令。

3. 确认全局群集节点显示 `ok` 提示符（在基于 Oracle Solaris 的系统上）或在 GRUB 菜单中显示 `Press any key to continue` 消息（在基于 x86 的系统上）。

4. 以非群集模式引导全局群集节点。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot -xs
```

- 在基于 x86 的系统上，运行以下命令。

- a. 在 GRUB 菜单中，使用方向键选择适当的 Oracle Solaris 条目，然后键入 **e** 编辑其命令。

GRUB 菜单随即显示。

有关基于 GRUB 的引导的更多信息，请参见《[引导和关闭 Oracle Solaris 11.3 系统](#)》中的“引导系统”。

- b. 在引导参数屏幕中，使用方向键选择内核项，然后键入 **e** 编辑该项。

GRUB 引导参数屏幕随即显示。

- c. 在命令中添加 **-x** 以指定将系统引导至非群集模式。

```
[ Minimal BASH-like line editing is supported. For the first word, TAB
lists possible command completions. Anywhere else TAB lists the possible
completions of a device/filename. ESC at any time exits. ]
```

```
grub edit> kernel$ /platform/i86pc/kernel/$ISADIR/unix -B $ZFS-BOOTFS -x
```

- d. 按 **Enter** 键接受所做更改，并返回到引导参数屏幕。

屏幕将显示编辑后的命令。

- e. 键入 **b** 将节点引导至非群集模式。

注 - 对内核引导参数命令所做的这一更改在系统引导后将失效。下次重新引导节点时，系统将引导至群集模式。如果希望引导至非群集模式，请重新执行上述步骤，并将 **-x** 选项添加到内核引导参数命令中。

例 29 SPARC: 在非群集模式下引导全局群集节点

以下示例显示了当节点 `phys-schost-1` 关闭并以非群集模式重新启动时的控制台输出。`-g0` 选项表示将宽限期设置为零，`-y` 选项表示在接收到要求确认的问题时自动回答 `yes`，而 `-i0` 选项则表示调用运行级别 0（零）。全局群集中其他节点的控制台上也显示此节点的关闭消息。

```
phys-schost# clnode evacuate phys-schost-1
phys-schost# cluster shutdown -g0 -y
Shutdown started.    Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:

WARNING: CMM monitoring disabled.
phys-schost-1#
...
rg_name = schost-sa-1 ...
offline node = phys-schost-2 ...
num of node = 0 ...
phys-schost-1#
INIT: New run level: 0
```

```
The system is coming down. Please wait.  
System services are now being stopped.  
Print services stopped.  
syslogd: going down on signal 15  
...  
The system is down.  
syncing file systems... done  
WARNING: node phys-schost-1 is being shut down.  
Program terminated  
  
ok boot -x  
...  
Not booting as part of cluster  
...  
The system is ready.  
phys-schost-1 console login:
```

修复已满的 /var 文件系统

Oracle Solaris 软件和 Oracle Solaris Cluster 软件都会将错误消息写入 /var/adm/messages 文件，因此随着时间的推移，/var 文件系统将会被填满。如果某个群集节点的 /var 文件系统已满，则 Oracle Solaris Cluster 在下次引导时可能无法在该节点上启动。此外，您可能无法登录到此节点。

▼ 如何修复已满的 /var 文件系统

如果某个节点报告 /var 文件系统已满而且继续运行 Oracle Solaris Cluster 服务，则请按以下过程来清理被占满的文件系统。有关更多信息，请参阅[《在 Oracle Solaris 11.3 中排除系统管理问题》](#)中的“系统消息格式”。

1. 在带有完整 /var 文件系统的群集节点上承担 root 角色。
2. 清理被占满的文件系统。
例如，删除该文件系统中包含的无关紧要的文件。

数据复制方法

本章介绍可与 Oracle Solaris Cluster 软件一起使用的的数据复制技术。数据复制是指将数据从主存储设备复制到备份设备（即辅助设备）中。如果主设备发生故障，您可从辅助设备中获取数据。数据复制有助于确保群集的高可用性和容灾性 (disaster tolerance)。

Oracle Solaris Cluster 软件支持以下类型的数据复制：

- 群集间 – 使用 Oracle Solaris Cluster Geographic Edition 进行灾难恢复
- 群集内 – 在校园群集内，可用作基于主机的镜像的替代方案

要执行数据复制，必须具有与待复制对象同名的设备组。一个设备一次只能属于一个设备组，因此，如果已经存在包含该设备的 Oracle Solaris Cluster 设备组，则必须将该组删除后才能将该设备添加到新的设备组。有关创建和管理 Solaris Volume Manager、ZFS 或原始磁盘设备组的说明，请参见“[管理设备组](#)” [112]。

您必须先了解基于主机和基于存储的数据复制，才能选择最适合您群集的复制方法。有关使用 Oracle Solaris Cluster Geographic Edition 框架管理数据复制以实现灾难恢复的更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)》。

本章包含下节：

- “[了解数据复制](#)” [93]
- “[在校园群集内使用基于存储的数据复制](#)” [95]

了解数据复制

Oracle Solaris Cluster 软件支持基于主机和基于存储的数据复制。

- 基于主机的数据复制使用软件在地理位置分散的群集之间实时复制磁盘卷。远程镜像复制则可将数据从主群集的主卷复制到分散在不同地理位置的辅助群集的主卷上。该软件使用远程镜像位图来跟踪主磁盘上的主卷与辅助磁盘上的主卷之间的差别。例如，Oracle Solaris 的 Availability Suite 功能就是一种用于在群集之间或在群集与群集外的主机之间进行复制的基于主机的复制软件。

由于基于主机的数据复制使用主机资源，而无需使用特殊的存储阵列，因此是一种花费不高的数据复制解决方案。如果数据库、应用程序或文件系统配置为允许运行

Oracle Solaris OS 的多个主机将数据写入到共享卷，则不支持这些数据库、应用程序或文件系统，例如 Oracle RAC。

- 有关在两个群集之间使用基于主机的数据复制的更多信息，请参见《[Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Solaris Availability Suite](#)》。
- 要查看不使用 Geographic Edition 框架的基于主机的复制示例，请参见附录 A，附录 A, 部署示例：使用 Availability Suite 软件在群集间配置基于主机的数据复制。
- 基于存储的数据复制使用存储控制器上的软件将数据复制工作从群集节点转移到存储设备上。该软件可释放一些节点处理能力以响应群集请求。例如，EMC SRDF 就是一种可以在群集之内或群集之间复制数据的基于存储的软件。基于存储的数据复制在校园群集配置中特别重要，可以简化所需的基础结构。
 - 有关在校园群集环境中使用基于存储的数据复制的更多信息，请参见“[在校园群集内使用基于存储的数据复制](#)” [95]。
 - 有关在两个或更多个群集之间使用基于存储的复制，以及可自动完成该过程的 Oracle Solaris Cluster Geographic Edition 产品的更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)》。

支持的数据复制方法

Oracle Solaris Cluster 软件支持群集之间或群集内的下列数据复制方法：

- 群集之间的复制 – 对于使用 Oracle Solaris Cluster Geographic Edition 软件的灾难恢复，您可以使用基于主机的或基于存储的复制方法在群集之间执行数据复制。通常，您可以选择基于主机的复制和基于存储的复制两者之一，而不使用两者的组合。您可以使用 Geographic Edition 软件来管理这两种类型的复制。

复制类型	数据复制产品
基于主机的复制	Oracle Solaris 的 Availability Suite 功能。
基于存储的复制	EMC Symmetrix Remote Data Facility (SRDF)。 Hitachi TrueCopy 和 Hitachi Universal Replicator，通过 Oracle Solaris Cluster Geographic Edition 框架。 Oracle ZFS Storage Appliance。

有关更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)》中的“Data Replication”。

如果要 not 通过 Oracle Solaris Cluster Geographic Edition 软件来使用基于主机的复制，请参见附录 A, 部署示例：使用 Availability Suite 软件在群集间配置基于主机的数据复制中的说明。

如果要不通过 Oracle Solaris Cluster Geographic Edition 软件来使用基于存储的复制，请参见复制软件的相关文档。

- 园区群集内的复制 – 此方法用作基于主机的镜像的替代方法。

复制类型	数据复制产品
基于存储的复制	EMC Symmetrix Remote Data Facility (SRDF)。 Hitachi TrueCopy 和 Hitachi Universal Replicator，通过 Oracle Solaris Cluster Geographic Edition 框架。

- 基于应用程序的复制 – Oracle Data Guard 是基于应用程序的复制软件的一个示例。此类型的软件仅用于使用 Geographic Edition 框架的灾难恢复，目的是复制单一实例 Oracle 数据库或 RAC 数据库。
有关更多信息，请参见 [《Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard》](#)。

在园区群集内使用基于存储的数据复制

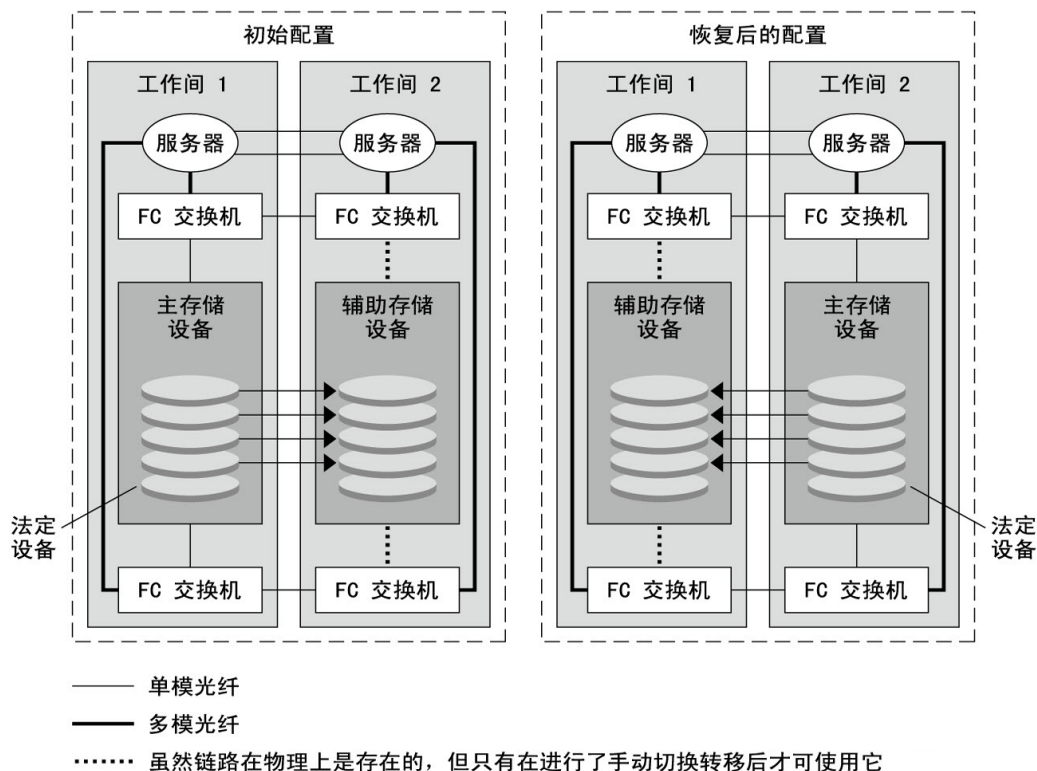
基于存储的数据复制使用安装在存储设备上的软件来管理称为园区群集的群集内复制。此类软件是特定存储设备所特有的，不用于灾难恢复。在配置基于存储的数据复制时，请参阅您的存储设备附带的文档。

根据您所用的软件，可以对基于存储的数据复制使用自动或手动的故障转移。Oracle Solaris Cluster 支持使用 EMC SRDF 软件对复制数据进行手动和自动故障转移。

本节介绍了在园区群集中使用的基于存储的数据复制。[图 1](#) 显示了在两个存储阵列间复制数据的双工作间配置样例。在该配置中，主存储阵列位于第一个工作间，并在此向两个工作间中的节点提供数据。主存储阵列还为辅助存储阵列提供要复制的数据。

注 - [图 1](#) 显示法定设备位于未复制的卷上。复制卷无法用作法定设备。

图 1 使用基于存储的数据复制的双工作间配置



Oracle Solaris Cluster 支持使用 EMC SRDF 的基于存储的同步复制。EMC SRDF 不支持异步复制。

请勿使用 EMC SRDF 的 Domino 模式或自适应复制模式。Domino 模式在目标不可用时，会使本地和目标 SRDF 卷对主机均不可用。自适应复制模式通常用于数据迁移和数据中心移动，不建议用于灾难恢复。

如果与远程存储设备失去联系，请确保主群集上运行的应用程序不会因为指定了设置为 never 或 async 的 Fence_level 而受到阻止。如果将 Fence_level 指定为 data 或 status，则当无法将更新内容复制到远程存储设备上时，主存储设备会拒绝更新。

在校园群集内使用基于存储的数据复制的要求和限制

为了确保数据完整性，请使用多路径和正确的 RAID 软件包。要使用基于存储的数据复制来实现群集配置，请注意下列事项：

- 如果为群集配置了自动故障转移，请使用同步复制。
有关配置群集以对复制卷进行自动故障转移的说明，请参见[“配置和管理基于存储的复制设备” \[100\]](#)。有关设计校园群集的要求的详细信息，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》中的“[Shared Data Storage](#)”。
- 某些应用程序特定数据可能不适合进行异步数据复制。请根据您的应用程序行为了解来确定在存储设备间复制应用程序特定数据的最佳方式。
- 节点到节点的距离受 Oracle Solaris Cluster 光纤通道和互连基础结构的限制。要获得有关当前限制和所支持技术的更多信息，请联系 Oracle 服务提供商。
- 不要将复制卷配置为法定设备。应使任何法定设备位于共享的非复制卷中，或使用法定服务器。
- 确保只有数据的主副本对群集节点可见。否则，卷管理器可能尝试同时访问数据的主副本和辅助副本。有关控制数据副本可见性的信息，请参阅存储阵列附带的相关文档。
- EMC SRDF 允许用户定义复制设备组。每个复制设备组需要具有相同名称的 Oracle Solaris Cluster 设备组。
- 对于使用 EMC SRDF 通过并发或级联 RDF 设备实现的三个站点或三个数据中心配置，必须在所有参与的群集节点上将以下条目添加到 Solutions Enabler SYMCLI 选项文件中：

```
SYMAPI_2SITE_CLUSTER_DG=device-group:rdf-group-number
```

此条目使群集软件能够自动在两个 SRDF 同步站点间移动应用程序。条目中的 *rdf-group-number* 代表 RDF 组，它将主机的本地 symmetrix 连接至第二个站点的 symmetrix。

有关三个数据中心配置的更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)》中的“[Three-Data-Center \(3DC\) Topologies](#)”。

- 在群集内执行复制时，SRDF 不支持 Oracle Real Application Clusters (Oracle RAC)。连接到非当前主副本的节点将不具有写入访问权限。复制设备无法支持任何要求群集的所有节点都具有直接写入权限的可伸缩应用程序。
- 不支持 Solaris Volume Manager for Sun Cluster 多属主磁盘集。
- 请勿在 EMC SRDF 中使用 Domino 模式或自适应复制模式。有关更多信息，请参见[“在校园群集内使用基于存储的数据复制” \[95\]](#)。

在校园群集内使用基于存储的数据复制时的手动恢复注意事项

与所有校园群集一样，使用基于存储的数据复制的群集通常在遇到单一故障时无需人为干预。但是，如果使用手动故障转移并且丢失了存放主存储设备的工作间（如[图 1](#) 中所示），则双节点群集中将出现问题。剩余的节点无法保留法定设备，也无法无法作为群集成员进行引导。在这种情况下，需要对群集进行以下手动干预：

1. 您的 Oracle 服务提供商必须重新配置剩余的节点，使其可作为群集成员进行引导。
2. 您或您的 Oracle 服务提供商必须将辅助存储设备的一个非复制卷配置为法定设备。
3. 您或您的 Oracle 服务提供商必须配置剩余的节点，使其将辅助存储设备作为主存储设备。这种重新配置可能涉及重新构建卷管理器卷、恢复数据或更改应用程序与存储卷的关联。

使用基于存储的数据复制时的最佳做法

当使用 EMC SRDF 软件进行基于存储的数据复制时，请使用动态设备而不是静态设备。静态设备更改主副本时需要几分钟，这会影响故障转移时间。

管理全局设备、磁盘路径监视和群集文件系统

本章介绍了有关管理全局设备、磁盘路径监视及群集文件系统的信息和具体过程。

- “管理全局设备和全局名称空间概述” [99]
- “配置和管理基于存储的复制设备” [100]
- “管理群集文件系统概述” [112]
- “管理设备组” [112]
- “管理存储设备的 SCSI 协议设置” [135]
- “管理群集文件系统” [140]
- “管理磁盘路径监视” [145]

有关本章中相关过程的概括性说明，请参见表 7 “任务列表：管理设备组”。

有关与全局设备、全局名称空间、设备组、磁盘路径监视以及群集文件系统相关的概念性信息，请参见《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》。

管理全局设备和全局名称空间概述

Oracle Solaris Cluster 设备组的管理依赖于群集上安装的卷管理器。Solaris Volume Manager 可以识别群集，因此您可以通过使用 Solaris Volume Manager `metaset` 命令添加、注册和删除设备组。有关更多信息，请参见 [metaset\(1M\)](#) 手册页。

Oracle Solaris Cluster 软件可自动为群集中的每个磁盘和磁带设备创建一个原始磁盘设备组。但是，群集设备组仍处于脱机状态，直到您将这些组作为全局设备访问。管理设备组或卷管理器磁盘组时，您需要在作为该组的主节点的群集节点上进行操作。

通常不需要管理全局设备名称空间。在安装过程中，系统会自动设置全局名称空间。而且，全局名称空间会在 Oracle Solaris OS 重新引导过程中自动更新。但是，如果需要更新全局名称空间，可从任意群集节点运行 `cldevice populate` 命令。该命令会导致在所有其他群集节点成员上以及将来可能加入群集的节点上更新全局名称空间。

Solaris Volume Manager 的全局设备许可

对于 Solaris Volume Manager 和磁盘设备，对全局设备权限所做的更改不会自动传播到群集中的所有节点。如果要更改全局设备的许可，必须手动更改群集中所有设备的许可。例如，如果要在全局设备 `/dev/global/dsk/d3s0` 上的权限更改为 644，必须在群集中的所有节点上发出以下命令：

```
# chmod 644 /dev/global/dsk/d3s0
```

动态重新配置全局设备

在对群集中的磁盘和磁带设备完成动态重新配置操作时，必须考虑以下问题。

- 针对 Oracle Solaris 动态重新配置功能介绍的所有要求、过程和限制也适用于 Oracle Solaris Cluster 动态重新配置支持。唯一的例外是操作系统的停止操作。因此，请先阅读关于 Oracle Solaris 动态重新配置功能的文档，然后再对 Oracle Solaris Cluster 软件使用动态重新配置功能。您应该特别注意那些在执行动态重新配置分离操作时影响非网络 IO 设备的问题。
- Oracle Solaris Cluster 拒绝在主节点中的活动设备上执行动态重新配置删除板操作。可以在主节点的非活动设备以及辅助节点的任意设备上执行动态重新配置操作。
- 执行动态重新配置操作之后，群集数据访问一如既往。
- Oracle Solaris Cluster 拒绝执行影响法定设备可用性的动态重新配置操作。有关更多信息，请参见[“动态重新配置法定设备” \[154\]](#)。



注意 - 如果当前的主节点在您对辅助节点执行动态重新配置操作时出现故障，则会影响群集的可用性。主节点将无处可转移故障，直到为其提供了一个新的辅助节点。

要对全局设备执行动态重新配置操作，须按所示顺序完成下列步骤。

表 5 任务列表：动态重新配置磁盘和磁带设备

任务	相关说明
1. 如果必须在当前主节点上执行影响活动设备组的动态重新配置操作，请先切换主节点和辅助节点，然后再对设备执行动态重新配置删除操作	如何切换设备组的主节点 [133]
2. 对要删除的设备执行动态重新配置删除操作	请检查系统附带的文档。

配置和管理基于存储的复制设备

您可以配置一个 Oracle Solaris Cluster 设备组，使其包含将使用基于存储的复制功能进行复制的设备。Oracle Solaris Cluster 软件支持使用 EMC Symmetrix Remote Data Facility 软件进行基于存储的复制。

在能够使用 EMC Symmetrix Remote Data Facility 软件复制数据之前，您必须熟悉基于存储的复制文档，并在系统上安装基于存储的复制产品以及最新的更新。有关安装基于存储的复制软件的信息，请参见产品文档。

基于存储的复制软件可将一对设备配置为副本，一个作为主副本，另一个作为辅助副本。在任意给定时间，与一组节点相连的设备将是主副本，与另一组节点相连的设备将是辅助副本。

在 Oracle Solaris Cluster 配置中，只要主副本所属的 Oracle Solaris Cluster 设备组发生移动，该主副本就会自动移动。因此，决不应在 Oracle Solaris Cluster 配置中直接移动主副本。正确的做法是，应通过移动关联的 Oracle Solaris Cluster 设备组来实现接管。



注意 - 您创建的 Oracle Solaris Cluster 设备组（Solaris Volume Manager 或原始磁盘）必须与复制的设备组同名。

管理 EMC Symmetrix Remote Data Facility 复制设备

下表列出了设置和管理 EMC Symmetrix Remote Data Facility (SRDF) 基于存储的复制设备时必须执行的任务。

表 6 任务列表：管理 EMC SRDF 基于存储的复制设备

任务	说明
在存储设备和节点上安装 SRDF 软件	EMC 存储设备随附的文档。
配置 EMC 复制组	如何配置 EMC SRDF 复制组 [101]
配置 DID 设备	如何使用 EMC SRDF 为复制配置 DID 设备 [103]
注册复制组	如何添加和注册设备组 (Solaris Volume Manager) [118]
检验配置	如何检验 EMC SRDF 复制全局设备组配置 [105]
校园群集的主工作间彻底故障后手动恢复数据	如何在主工作间彻底故障后恢复 EMC SRDF 数据 [110]

▼ 如何配置 EMC SRDF 复制组

- 开始之前
- 在配置 EMC Symmetrix Remote Data Facility (SRDF) 复制组之前，必须在所有群集节点上安装 EMC Solutions Enabler 软件。首先，在群集的共享磁盘上配置 EMC SRDF 设备组。有关如何配置 EMC SRDF 设备组的更多信息，请参见 EMC SRDF 产品文档。
 - 使用 EMC SRDF 时，请使用动态设备而不是静态设备。静态设备更改主副本时需要几分钟，这会影响故障转移时间。



注意 - 您创建的 Oracle Solaris Cluster 设备组 (Solaris Volume Manager 或原始磁盘) 必须与复制的设备组同名。

1. 在连接到存储阵列的所有节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 要使用并发 SRDF 或级联设备完成三站点或三数据中心实现，请设置 `SYMAPI_2SITE_CLUSTER_DG` 参数。

在所有参与的群集节点上向 Solutions Enabler 选项文件添加下列条目：

```
SYMAPI_2SITE_CLUSTER_DG=device-group:rdf-group-number
```

device-group 指定设备组的名称。

rdf-group-number 指定 RDF 组，用于将主机的本地 symmetrix 连接到第二个站点的 symmetrix。

此条目使群集软件能够自动在两个 SRDF 同步站点间移动应用程序。

有关三个数据中心配置的更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)》中的“[Three-Data-Center \(3DC\) Topologies](#)”。

3. 在配置有复制数据的每个节点上，搜索 Symmetrix 设备配置。
此过程可能需要几分钟时间。

```
# /usr/symcli/bin/symcfg discover
```

4. 如果尚未创建副本对，请立即创建。
使用 `symrdf` 命令创建副本对。有关创建副本对的说明，请参阅 SRDF 文档。

注 - 如果使用并发 RDF 设备完成三站点或三数据中心实现，请在所有 `symrdf` 命令中添加以下参数：

```
-rdfg rdf-group-number
```

在 `symrdf` 命令中指定 RDF 组编号可确保 `symrdf` 操作针对正确的 RDF 组执行。

5. 在配置有复制设备的每个节点上，检验数据复制设置是否正确。

```
# /usr/symcli/bin/symdg show group-name
```

6. 执行设备组交换。

- a. 检验主副本和辅助副本是否同步。

```
# /usr/symcli/bin/symrdf -g group-name verify -synchronized
```

- b. 使用 `symdg show` 命令确定哪个节点包含主副本，哪个节点包含辅助副本。

```
# /usr/symcli/bin/symdg show group-name
```

具有 RDF1 设备的节点包含主副本，具有 RDF2 设备状态的节点包含辅助副本。

- c. 启用辅助副本。

```
# /usr/symcli/bin/symrdf -g group-name failover
```

- d. 交换 RDF1 和 RDF2 设备。

```
# /usr/symcli/bin/symrdf -g group-name swap -refresh R1
```

- e. 启用副本对。

```
# /usr/symcli/bin/symrdf -g group-name establish
```

- f. 检验主节点和辅助副本是否同步。

```
# /usr/symcli/bin/symrdf -g group-name verify -synchronized
```

7. 在最初具有主副本的节点上重复步骤 5 的所有操作。

接下来的步骤 为 EMC SRDF 复制设备配置了设备组之后，必须配置复制设备所使用的设备标识符 (Device Identifier, DID) 驱动程序。

▼ 如何使用 EMC SRDF 为复制配置 DID 设备

此过程将配置复制设备所使用的设备标识符 (Device Identifier, DID) 驱动程序。确保指定的 DID 设备实例是彼此的副本，并且属于指定的复制组。

开始之前 `phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 确定与已配置的 RDF1 和 RDF2 设备相对应的 DID 设备。

```
# /usr/symcli/bin/symdg show group-name
```

注 - 如果系统未显示整个 Oracle Solaris 设备修补程序，请将环境变量 `SYMCLI_FULL_PDEVNAME` 设置为 1 并重新键入 `symdg -show` 命令。

3. 确定与 Oracle Solaris 设备相对应的 DID 设备。

```
# cldevice list -v
```

4. 对于每对匹配的 DID 设备，将实例组合到单个复制 DID 设备中。从 RDF2/辅助设备端运行以下命令。

```
# cldevice combine -t srdf -g replication-device-group \  
-d destination-instance source-instance
```

注 - SRDF 数据复制设备不支持 -t 选项。

-t replication-type

指定复制类型。对于 EMC SRDF，键入 **SRDF**。

-g replication-device-group

指定设备组的名称，如 `symdg show` 命令中所示。

-d destination-instance

指定与 RDF1 设备相对应的 DID 实例。

source-instance

指定与 RDF2 设备相对应的 DID 实例。

注 - 如果组合了错误的 DID 设备，请使用 `scdidadm` 命令的 `-b` 选项撤消这两个 DID 设备的组合。

```
# scdidadm -b device
```

-b device 组合实例时与 `destination_device` 相对应的 DID 实例。

5. 如果复制设备组的名称发生更改，请执行以下其他步骤。
如果复制设备组（以及相应的全局设备组）名称发生更改，则必须首先通过使用 `scdidadm -b` 命令删除现有信息来更新复制设备信息。最后一步是使用 `cldevice combine` 命令创建一个新的更新过的设备。

6. 检验是否已组合 DID 实例。

```
# cldevice list -v device
```

7. 检验是否已设置 SRDF 复制。

```
# cldevice show device
```

8. 在所有节点上，检验是否可访问所有组合 DID 实例的 DID 设备。


```
# cldevice list -v
```

接下来的步骤 配置完复制设备所使用的设备标识符 (device identifier, DID) 驱动程序之后，必须检验 EMC SRDF 复制全局设备组配置。

▼ 如何检验 EMC SRDF 复制全局设备组配置

开始之前 在检验全局设备组之前，必须先创建它。您可以使用 Solaris Volume Manager、ZFS 或原始磁盘中的设备组。有关更多信息，请参阅以下内容：

- [如何添加和注册设备组 \(Solaris Volume Manager\) \[118\]](#)
- [如何添加并注册设备组（原始磁盘） \[120\]](#)
- [如何添加并注册复制设备组 \(ZFS\) \[121\]](#)



注意 - 您创建的 Oracle Solaris Cluster 设备组（Solaris Volume Manager 或原始磁盘）必须与复制的设备组同名。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 检验主设备组是否对应于包含主副本的同一节点。

```
# symdg -show group-name
# cldevicegroup status -n nodename group-name
```

2. 尝试执行切换，确保设备组已正确配置且副本可在节点间移动。

如果设备组处于脱机状态，请使其联机。

```
# cldevicegroup switch -n nodename group-name
```

-n nodename 设备组切换到的节点。此节点成为新的主节点。

3. 通过比较以下命令的输出来检验切换操作是否已成功。

```
# symdg -show group-name
# cldevicegroup status -n nodename group-name
```

示例：为 Oracle Solaris Cluster 配置 SRDF 复制组

本示例用于完成在群集中设置 SRDF 复制时所需执行的特定于 Oracle Solaris Cluster 的步骤。本示例假定您已执行以下任务：

- 已针对阵列间的复制完成了 LUN 配对。
- 在存储设备和群集节点上安装 SRDF 软件。

例 30 创建副本对

本示例涉及一个四节点群集，其中两个节点连接到一个 Symmetrix，另外两个节点连接到第二个 Symmetrix。SRDF 设备组的名称为 dg1。

在所有节点上运行以下命令

```
# symcfg discover
! This operation might take up to a few minutes.
```

```
# symdev list pd

Symmetrix ID: 000187990182
```

	Device Name	Directors	Device			
						Cap
Sym	Physical	SA :P DA :IT	Config	Attribute	Sts	(MB)
0067	c5t600604800001879901*	16D:0 02A:C1	RDF2+Mir	N/Grp'd	RW	4315
0068	c5t600604800001879901*	16D:0 16B:C0	RDF1+Mir	N/Grp'd	RW	4315
0069	c5t600604800001879901*	16D:0 01A:C0	RDF1+Mir	N/Grp'd	RW	4315

... 在 RDF1 端的所有节点上，运行以下命令

```
# symdg -type RDF1 create dg1
# symld -g dg1 add dev 0067
```

在 RDF2 端的所有节点上，运行以下命令

```
# symdg -type RDF2 create dg1
# symld -g dg1 add dev 0067
```

例 31 检验数据复制设置

在群集的一个节点上执行以下命令。

```
# symdg show dg1

Group Name: dg1

Group Type : RDF1 (RDFA)
Device Group in GNS : No
Valid : Yes
Symmetrix ID : 000187900023
Group Creation Time : Thu Sep 13 13:21:15 2007
Vendor ID : EMC Corp
Application ID : SYMCLI

Number of STD Devices in Group : 1
```

```

Number of Associated GK's           : 0
Number of Locally-associated BCV's  : 0
Number of Locally-associated VDEV's : 0
Number of Remotely-associated BCV's (STD RDF): 0
Number of Remotely-associated BCV's (BCV RDF): 0
Number of Remotely-assoc'd RBCV's (RBCV RDF) : 0

```

Standard (STD) Devices (1):

```

{
-----
LdevName          PdevName          Sym      Att. Sts      Cap
Dev              Dev              Dev              (MB)
-----
DEV001            /dev/rdisk/c5t6006048000018790002353594D303637d0s2 0067  RW   4315
}

```

Device Group RDF Information

...

symrdf -g dg1 establish

Execute an RDF 'Incremental Establish' operation for device
group 'dg1' (y/[n]) ? y

An RDF 'Incremental Establish' operation execution is
in progress for device group 'dg1'. Please wait...

```

Write Disable device(s) on RA at target (R2).....Done.
Suspend RDF link(s).....Done.
Mark target (R2) devices to refresh from source (R1).....Started.
Device: 0067 ..... Marked.
Mark target (R2) devices to refresh from source (R1).....Done.
Merge device track tables between source and target.....Started.
Device: 0067 ..... Merged.
Merge device track tables between source and target.....Done.
Resume RDF link(s).....Started.
Resume RDF link(s).....Done.

```

The RDF 'Incremental Establish' operation successfully initiated for
device group 'dg1'.

#

symrdf -g dg1 query

```

Device Group (DG) Name      : dg1
DG's Type                   : RDF2
DG's Symmetrix ID           : 000187990182

```

	Target (R2) View					Source (R1) View				MODES	
	ST				LI	ST					
Standard	A				N	A					
Logical	T	R1 Inv	R2 Inv	K	T	R1 Inv	R2 Inv			RDF Pair	
Device Dev	E	Tracks	Tracks	S Dev	E	Tracks	Tracks	MDA		STATE	

```
-----
DEV001  0067 WD      0      0 RW 0067 RW      0      0 S..   Synchronized

Total      -----
MB(s)      0.0      0.0      0.0      0.0

Legend for MODES:

M(ode of Operation): A = Async, S = Sync, E = Semi-sync, C = Adaptive Copy
D(omino)           : X = Enabled, . = Disabled
A(daptive Copy)    : D = Disk Mode, W = WP Mode, . = ACp off

#
```

例 32 显示与所用磁盘相对应的 DID

对 RDF1 和 RDF2 端执行相同的过程。

您可以在 `dymdg show dg` 命令输出的 `PdevName` 字段下查看。

```
在 RDF1 端运行以下命令
# symdg show dg1

Group Name:  dg1

Group Type                : RDF1      (RDFA)
...
Standard (STD) Devices (1):
{
-----
LdevName      PdevName      Sym      Cap
Dev  Att.  Sts      (MB)
-----
DEV001      /dev/rdisk/c5t6006048000018790002353594D303637d0s2  0067      RW      4315
}

Device Group RDF Information
...
获取相应 DID
# cldevice list | grep c5t6006048000018790002353594D303637d0
217      pmoney1:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdisk/d217
217      pmoney2:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdisk/d217

列出相应 DID
# cldevice show d217

=== DID Device Instances ===

DID Device Name:      /dev/did/rdisk/d217
Full Device Path:      pmoney2:/dev/rdisk/
c5t6006048000018790002353594D303637d0
```

```

Full Device Path:                pmoney1:/dev/rdisk/
c5t600604800001879002353594D303637d0
Replication:                     none
default_fencing:                 global

#

    在 RDF2 端运行以下命令
# symdg show dg1

Group Name:  dg1

Group Type                        : RDF2      (RDFA)
...
Standard (STD) Devices (1):
{
-----
LdevName      PdevName      Sym      Cap
Dev  Att. Sts  (MB)
-----
DEV001        /dev/rdisk/c5t6006048000018799018253594D303637d0s2  0067    WD    4315
}

Device Group RDF Information
...
    获取相应 DID
# cldevice list | grep c5t6006048000018799018253594D303637d0
108      pmoney4:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdisk/d108
108      pmoney3:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdisk/d108

    列出相应 DID
# cldevice show d108

=== DID Device Instances ===

DID Device Name:                /dev/did/rdisk/d108
Full Device Path:                pmoney3:/dev/rdisk/c5t6006048000018799018253594D303637d0
Full Device Path:                pmoney4:/dev/rdisk/c5t6006048000018799018253594D303637d0
Replication:                     none
default_fencing:                 global

#

```

例 33 组合 DID 实例

从 RDF2 端，键入：

```
# cldevice combine -t srdf -g dg1 -d d217 d108
```

例 34 显示组合的 DID

从群集中的任一节点上，键入：

```
# cldevice show d217 d108
cldevice: (C727402) Could not locate instance "108".

=== DID Device Instances ===

DID Device Name:                /dev/did/rdisk/d217
Full Device Path:                pmoney1:/dev/rdisk/
c5t6006048000018790002353594D303637d0
Full Device Path:                pmoney2:/dev/rdisk/
c5t6006048000018790002353594D303637d0
Full Device Path:                pmoney4:/dev/rdisk/
c5t6006048000018799018253594D303637d0
Full Device Path:                pmoney3:/dev/rdisk/
c5t6006048000018799018253594D303637d0
Replication:                    srdf
default_fencing:                global
```

▼ 如何在主工作间彻底故障后恢复 EMC SRDF 数据

此过程可在校园群集的主工作间彻底故障、故障转移至辅助工作间并将主工作间重新联机之后，执行数据恢复。校园群集的主工作间是主节点和存储站点。工作间的彻底故障是指该工作间内的主机和存储设备同时发生故障。如果主工作间发生故障，Oracle Solaris Cluster 会自动故障转移至辅助工作间，使辅助工作间的存储设备可读写，并启用相应设备组和资源组的故障转移。

主工作间重新联机后，可手动从 SRDF 设备组恢复被写入至辅助工作间的数据，并重新同步数据。此过程可通过将原辅助工作间（此过程使用 `phys-campus-2` 作为辅助工作间）中的数据同步至原主工作间（`phys-campus-1`），来恢复 SRDF 设备组。此过程还会将 `phys-campus-2` 和 `phys-campus-1` 上的 SRDF 设备组类型分别更改为 RDF1 和 RDF2。

开始之前 在执行手动故障转移之前，必须先配置 EMC 复制组和 DID 设备，并注册 EMC 复制组。有关创建 Solaris Volume Manager 设备组的信息，请参见[如何添加和注册设备组 \(Solaris Volume Manager\) \[118\]](#)。

注 - 这些说明演示了一种在主工作间完成故障转移并重新联机后手动恢复 SRDF 数据的方法。有关其他方法，请查阅 EMC 文档。

登录校园群集的主工作间执行以下步骤。在上述过程中，`dg1` 为 SRDF 设备组名。发生故障时，此过程中的主工作间是 `phys-campus-1`，辅助工作间是 `phys-campus-2`。

1. 登录校园群集的主工作间，并承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 在主工作间中，使用 `symrdf` 命令查询 RDF 设备的复制状态并查看有关这些设备的信息。

```
phys-campus-1# symrdf -g dg1 query
```

提示 - 如果设备组处于 `split` 状态，说明未同步。

3. 如果 RDF 对的状态为 `split` 且设备组类型为 RDF1，则强制执行 SRDF 设备组的故障转移。

```
phys-campus-1# symrdf -g dg1 -force failover
```

4. 查看 RDF 设备的状态。

```
phys-campus-1# symrdf -g dg1 query
```

5. 故障转移后，可交换执行过故障转移的 RDF 设备上的数据。

```
phys-campus-1# symrdf -g dg1 swap
```

6. 查看有关 RDF 设备的状态及其他信息。

```
phys-campus-1# symrdf -g dg1 query
```

7. 在主工作间中建立 SRDF 设备组。

```
phys-campus-1# symrdf -g dg1 establish
```

8. 确认设备组处于已同步状态且设备组类型为 RDF2。

```
phys-campus-1# symrdf -g dg1 query
```

例 35 主站点故障转移后手动恢复 EMC SRDF 数据

本示例提供在校园群集的主工作间发生故障转移、辅助工作间进行接管并记录数据，随后主工作间重新联机之后，手动恢复 EMC SRDF 数据所必需的 Oracle Solaris Cluster 特定步骤。在本示例中，SRDF 设备组名为 `dg1`，标准逻辑设备为 `DEV001`。发生故障时，主工作间为 `phys-campus-1`，辅助工作间为 `phys-campus-2`。从校园群集的主工作间 `phys-campus-1` 中执行以下步骤。

```
phys-campus-1# symrdf -g dg1 query | grep DEV
DEV001 0012RW 0 0NR 0012RW 2031 0 S.. Split
```

```
phys-campus-1# symdg list | grep RDF
dg1 RDF1 Yes 00187990182 1 0 0 0 0
```

```
phys-campus-1# symrdf -g dg1 -force failover
...
```

```
phys-campus-1# symrdf -g dg1 query | grep DEV
DEV001 0012 WD 0 0 NR 0012 RW 2031 0 S.. Failed Over
```

```
phys-campus-1# symdg list | grep RDF
dg1 RDF1 Yes 00187990182 1 0 0 0 0
```

```
phys-campus-1# symrdf -g dg1 swap
...

phys-campus-1# symrdf -g dg1 query | grep DEV
DEV001  0012 WD  0  0 NR 0012 RW  0  2031 S.. Suspended

phys-campus-1# symdg list | grep RDF
dg1  RDF2  Yes  000187990182  1  0  0  0  0

phys-campus-1# symrdf -g dg1 establish
...

phys-campus-1# symrdf -g dg1 query | grep DEV
DEV001  0012 WD  0  0 RW 0012 RW  0  0 S.. Synchronized

phys-campus-1# symdg list | grep RDF
dg1  RDF2  Yes  000187990182  1  0  0  0  0
```

管理群集文件系统概述

群集文件系统管理不需要任何特殊 Oracle Solaris Cluster 命令。可以使用标准 Oracle Solaris 文件系统命令（如 `mount` 和 `newfs`），像管理其他任何 Oracle Solaris 文件系统一样来管理群集文件系统。可以通过为 `mount` 命令指定 `-g` 选项来挂载群集文件系统。群集文件系统使用 UFS，并且还可以在引导时自动挂载。只有从全局群集中的节点上才能看到群集文件系统。

注 - 群集文件系统在读取文件时，文件系统并不更新这些文件的访问时间。

群集文件系统限制

以下限制适用于群集文件系统的管理：

- 不支持对非空目录执行 `unlink` 命令。有关更多信息，请参见 [unlink\(1M\)](#) 手册页。
- 不支持 `lockfs -d` 命令。使用 `lockfs -n` 作为解决方法。
- 不能使用在重新挂载时添加的 `directio` 挂载选项重新挂载群集文件系统。

管理设备组

随着群集要求的变化，您可能需要在群集中添加、删除或修改设备组。Oracle Solaris Cluster 提供了一个名为 `clsetup` 的交互式界面，可使用该界面进行这些更改。`clsetup`

实用程序生成 `cluster` 命令。生成的命令显示在某些过程结尾部分的示例中。下表列出了用于管理设备组的各项任务，并提供了指向本节中相应操作过程的链接。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面使设备组联机和脱机。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

Oracle Solaris Cluster 软件可自动为群集中的每个磁盘和磁带设备创建一个原始磁盘设备组。但是，群集设备组仍处于脱机状态，直到您将这些组作为全局设备访问。



注意 - 如果其他节点是活动群集成员并且至少其中的一个节点拥有磁盘集，则请勿在群集之外引导的群集节点上运行 `metaset -s setname -f -t`。

表 7 任务列表：管理设备组

任务	说明
使用 <code>cldevice populate</code> 命令在不进行重新配置重新引导的情况下更新全局设备名称空间	如何更新全局设备名称空间 [114]
更改用于全局设备名称空间的 <code>lofi</code> 设备的大小	如何更改用于全局设备名称空间的 <code>lofi</code> 设备的大小 [114]
移动现有全局设备名称空间	如何将全局设备名称空间从专用分区迁移到 <code>lofi</code> 设备 [116] 如何将全局设备名称空间从 <code>lofi</code> 设备迁移到专用分区 [117]
使用 <code>metaset</code> 命令添加 Solaris Volume Manager 磁盘集并将其注册为设备组	如何添加和注册设备组 (Solaris Volume Manager) [118]
使用 <code>cldevicegroup</code> 命令添加并注册原始磁盘设备组	如何添加并注册设备组（原始磁盘） [120]
使用 <code>cldevicegroup</code> 命令为 ZFS 添加已命名的设备组	如何添加并注册复制设备组 (ZFS) [121]
使用 <code>metaset</code> 和 <code>metaclear</code> 命令从配置中删除 Solaris Volume Manager 设备组	“如何删除和取消注册设备组 (Solaris Volume Manager)” [124]
使用 <code>cldevicegroup</code> 、 <code>metaset</code> 和 <code>clsetup</code> 命令从所有设备组中删除一个节点	如何将节点从所有设备组中删除 [124]
使用 <code>metaset</code> 命令从 Solaris Volume Manager 设备组中删除一个节点	如何将节点从设备组中删除 (Solaris Volume Manager) [125]
使用 <code>cldevicegroup</code> 命令从原始磁盘设备组中删除一个节点	如何从原始磁盘设备组删除节点 [127]
使用 <code>clsetup</code> 命令生成 <code>cldevicegroup</code> 来更改设备组的属性	如何更改设备组属性 [128]
使用 <code>cldevicegroup show</code> 命令显示设备组及其属性	如何列出设备组配置 [132]
使用 <code>clsetup</code> 生成 <code>cldevicegroup</code> 来更改设备组所需的辅助节点数量	如何设置设备组所需的辅助节点数 [130]
使用 <code>cldevicegroup switch</code> 命令切换设备组的主节点	如何切换设备组的主节点 [133]
使用 <code>metaset</code> 命令将设备组置于维护状态	如何将设备组置于维护状态 [134]

▼ 如何更新全局设备名称空间

当添加新的全局设备时，请通过运行 `cldevice populate` 命令手动更新全局设备名称空间。

注 - 如果运行 `cldevice populate` 命令的节点当前不是群集成员，则该命令没有任何效果。如果未挂载 `/global/.devices/node@ nodeID` 文件系统，则该命令也没有任何效果。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 在群集中的每个节点上运行 `devfsadm` 命令。
您可以同时在群集的所有节点上运行此命令。有关更多信息，请参见 [devfsadm\(1M\)](#) 手册页。
3. 重新配置名称空间。

```
# cldevice populate
```

4. 在每个节点上，先验证 `cldevice populate` 命令是否已完成，然后再尝试创建磁盘集。
`cldevice` 命令会在所有节点上远程调用其自身，即使仅从一个节点上运行该命令也是如此。要确定 `cldevice populate` 命令是否已完成处理过程，请在群集的每个节点上运行以下命令。

```
# ps -ef | grep cldevice populate
```

例 36 更新全局设备名称空间

以下示例显示了成功运行 `cldevice populate` 命令后生成的输出。

```
# devfsadm
cldevice populate
Configuring the /dev/global directory (global devices)...
obtaining access to all attached disks
reservation program successfully exiting
# ps -ef | grep cldevice populate
```

▼ 如何更改用于全局设备名称空间的 `lofi` 设备的大小

如果在全局群集的一个或多个节点上为全局设备名称空间使用 `lofi` 设备，请执行此过程更改该设备的大小。

1. 在要为其 `lofi` 设备调整全局设备名称空间的节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。

2. 从节点清除服务，并将节点重新引导至非群集模式。
这样做可以确保您在执行此过程时不会从该节点为全局设备提供服务。有关说明，请参见[如何以非群集模式引导节点 \[89\]](#)。
3. 卸载全局设备文件系统，并分离其 lofi 设备。
全局设备文件系统在本地挂载。

```
phys-schost# umount /global/.devices/node\@`clinfo -n` > /dev/null 2>&1
```

确保 lofi 设备已分离

```
phys-schost# lofiadm -d /.globaldevices
```

如果该设备已分离，则该命令将不返回任何输出

注 - 如果使用 -m 选项挂载文件系统，则不会有任何条目添加到 mnttab 文件。umount 命令可能会报告类似如下的警告：

```
umount: warning: /global/.devices/node@2 not in mnttab  =====>>>not mounted
```

可以放心地忽略此警告。

4. 删除并重新创建所需大小的 /.globaldevices 文件。
以下示例显示了如何创建大小为 200 MB 的新 /.globaldevices 文件。

```
phys-schost# rm /.globaldevices
```

```
phys-schost# mkfile 200M /.globaldevices
```

5. 为全局设备名称空间创建新文件系统。

```
phys-schost# lofiadm -a /.globaldevices
```

```
phys-schost# newfs `lofiadm /.globaldevices` < /dev/null
```

6. 将节点引导到群集模式下。
全局设备现在位于新文件系统中。

```
phys-schost# reboot
```

7. 将希望在节点上运行的任何服务迁移到该节点。

迁移全局设备名称空间

可在回送文件接口 (lofi) 设备上创建名称空间，而不是在专用分区上创建全局设备名称空间。

注 - 支持 ZFS 用作根文件系统，但存在一个显著例外：如果使用引导磁盘的专用分区作为全局设备文件系统，则必须仅使用 UFS 作为其文件系统。全局设备名称空间要求在 UFS 文件系统中运行代理文件系统 (PxFS)。

但针对全局设备名称空间使用的 UFS 文件系统可与对根 (/) 文件系统及其他根文件系统（例如，/var 或 /home）使用的 ZFS 文件系统共存。另外，如果改为使用 lofi 设备来托管全局设备名称空间，则不会限制对根文件系统使用 ZFS。

以下程序介绍如何在专用分区与 lofi 设备之间移动现有全局设备名称空间：

- [如何将全局设备名称空间从专用分区迁移到 lofi 设备 \[116\]](#)
- [如何将全局设备名称空间从 lofi 设备迁移到专用分区 \[117\]](#)

▼ 如何将全局设备名称空间从专用分区迁移到 lofi 设备

1. 在要更改其名称空间位置的全局群集节点上，承担 root 角色。
2. 从节点清除服务，并将节点重新引导至非群集模式。
这样做可以确保您在执行此过程时不会从该节点为全局设备提供服务。有关说明，请参见[如何以非群集模式引导节点 \[89\]](#)。
3. 确保该节点上不存在名为 /.globaldevices 的文件。
如果存在该文件，请将其删除。

4. 创建 lofi 设备。

```
# mkfile 100m /.globaldevices# lofiadm -a /.globaldevices
# LOFI_DEV=`lofiadm /.globaldevices`
# newfs `echo ${LOFI_DEV} | sed -e 's/lofi/rlofi/g'` < /dev/null
# lofiadm -d /.globaldevices
```

5. 在 /etc/vfstab 文件中，注释掉全局设备名称空间条目。
该条目具有以 /global/.devices/node@nodeID 开头的挂载路径。
6. 卸载全局设备分区 /global/.devices/node@nodeID。
7. 禁用然后重新启用 globaldevices 和 scmountdev SMF 服务。

```
# svcadm disable globaldevices
# svcadm disable scmountdev
# svcadm enable scmountdev
# svcadm enable globaldevices
```

现已在 /.globaldevices 中创建 lofi 设备并挂载为全局设备文件系统。

8. 如果要将其他节点的全局设备名称空间从某一分区迁移到 lofi 设备，重复这些步骤即可。

9. 从一个节点填充全局设备名称空间。

```
# cldevice populate
```

请先在每个节点上检验命令是否已完成处理，然后再对群集执行其他操作。

```
# ps -ef | grep "cldevice populate"
```

全局设备名称空间现已驻留在 lofi 设备上。

10. 将希望在节点上运行的任何服务迁移到该节点。

▼ 如何将全局设备名称空间从 lofi 设备迁移到专用分区

1. 在要更改其名称空间位置的全局群集节点上，承担 root 角色。
2. 从节点清除服务，并将节点重新引导至非群集模式。
这样做可以确保您在执行此过程时不会从该节点为全局设备提供服务。有关说明，请参见[如何以非群集模式引导节点 \[89\]](#)。

3. 在节点的本地磁盘上，创建符合以下要求的新分区：

- 大小至少为 512M
- 使用 UFS 文件系统

4. 在 `/etc/vfstab` 文件中为新分区添加一个条目，使其挂载为全局设备文件系统。

- 确定当前节点的节点 ID。

```
# /usr/sbin/clinfo -n node-ID
```

- 使用以下格式在 `/etc/vfstab` 文件中创建新条目：

```
blockdevice rawdevice /global/.devices/node@nodeID ufs 2 no global
```

例如，如果您选择使用的分区是 `/dev/did/rdisk/d5s3`，则添加到 `/etc/vfstab` 文件中的新条目将如下所示：

```
/dev/did/dsk/d5s3 /dev/did/rdisk/d5s3 /global/.devices/node@3 ufs 2 no global
```

5. 卸载全局设备分区 `/global/.devices/node@ nodeID`。
6. 删除与 `/.globaldevices` 文件相关联的 lofi 设备。

```
# lofiadm -d /.globaldevices
```

7. 删除 /.globaldevices 文件。

```
# rm /.globaldevices
```

8. 禁用然后重新启用 globaldevices 和 scmountdev SMF 服务。

```
# svcadm disable globaldevices# svcadm disable scmountdev  
# svcadm enable scmountdev  
# svcadm enable globaldevices
```

该分区现已挂载为全局设备名称空间文件系统。

9. 如果要将其他节点的全局设备名称空间从 lofi 设备迁移到某一分区，重复这些步骤即可。

10. 引导至群集模式，然后填充全局设备名称空间。

- a. 从群集中的一个节点填充全局设备名称空间。

```
# cldevice populate
```

- b. 在对任意节点执行其他操作之前，请确保群集所有节点均已完成此过程。

```
# ps -ef | grep cldevice populate
```

全局设备名称空间现已驻留在专用分区上。

11. 将希望在节点上运行的任何服务迁移到该节点。

添加并注册设备组

您可以为 Solaris Volume Manager、ZFS 或原始磁盘添加并注册设备组。

▼ 如何添加和注册设备组 (Solaris Volume Manager)

使用 `metaset` 命令可创建 Solaris Volume Manager 磁盘集，并将磁盘集注册为 Oracle Solaris Cluster 设备组。注册磁盘集时，系统会将您指定给磁盘集的名称自动指定给设备组。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。



注意 - 您创建的 Oracle Solaris Cluster 设备组 (Solaris Volume Manager 或原始磁盘) 必须与复制的设备组同名。

1. 在与磁盘 (您要在这些磁盘上创建磁盘集) 相连的节点中的一个节点上, 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 添加 Solaris Volume Manager 磁盘集并将其注册为 Oracle Solaris Cluster 设备组。
要创建多所有者磁盘组, 请使用 `-M` 选项。

```
# metaset -s diskset -a -M -h nodelist
```

`-s diskset` 指定要创建的磁盘集。

`-a -h nodelist` 添加一组可以主控磁盘集的节点。

`-M` 将磁盘组指定为多所有者的。

注 - 如果运行 `metaset` 命令在群集上建立 Solaris Volume Manager 设备组, 则默认情况下, 无论该设备组中包含多少个节点, 都会生成一个辅助节点。创建了设备组之后, 您可以使用 `clsetup` 实用程序更改所需辅助节点数。有关磁盘故障转移的更多信息, 请参阅[如何设置设备组所需的辅助节点数 \[130\]](#)。

3. 如果正在配置一个复制设备组, 请为设备组设置复制属性。

```
# cldevicegroup sync devicegroup
```

4. 检验是否已添加设备组。
设备组名称与使用 `metaset` 命令指定的磁盘集名称相符。

```
# cldevicegroup list
```

5. 列出 DID 映射。

```
# cldevice show | grep Device
```

- 选择由将要控制或可能要控制磁盘集的群集节点共享的驱动器。
- 向磁盘集添加驱动器时, 请使用格式为 `/dev/did/rdisk/d N` 的完整 DID 设备名称。

在下面的示例中, DID 设备 `/dev/did/rdisk/d3` 的条目表明 `phys-schost-1` 和 `phys-schost-2` 正在共享该驱动器。

```
=== DID Device Instances ===
DID Device Name:                    /dev/did/rdisk/d1
Full Device Path:                   phys-schost-1:/dev/rdisk/c0t0d0
DID Device Name:                   /dev/did/rdisk/d2
Full Device Path:                   phys-schost-1:/dev/rdisk/c0t6d0
```

```
DID Device Name:                /dev/did/rdisk/d3
Full Device Path:                phys-schost-1:/dev/rdisk/clt1d0
Full Device Path:                phys-schost-2:/dev/rdisk/clt1d0
...
```

6. 将驱动器添加到磁盘集。

使用完整的 DID 路径名称。

```
# metaset -s setname -a /dev/did/rdisk/dN
```

-s *setname* 指定磁盘集的名称，该名称与设备组名称相同。

-a 给磁盘集添加驱动器。

注 - 在向磁盘集添加驱动器时，请不要使用低级别的设备名称 (cNtXdY)。因为低级别设备名称是本地名称，并且在群集中不是唯一的，使用该名称可能使元集无法切换。

7. 检验磁盘集和驱动器的状态。

```
# metaset -s setname
```

例 37 添加 Solaris Volume Manager 设备组

以下示例显示了如何使用磁盘驱动器 /dev/did/rdisk/d1 和 /dev/did/rdisk/d2 创建磁盘集和设备组，并检验设备组是否已创建。

```
# metaset -s dg-schost-1 -a -h phys-schost-1
```

```
# cldevicegroup list
dg-schost-1
```

```
# metaset -s dg-schost-1 -a /dev/did/rdisk/d1 /dev/did/rdisk/d2
```

▼ 如何添加并注册设备组（原始磁盘）

除其他卷管理器以外，Oracle Solaris Cluster 软件还支持使用原始磁盘设备组。对 Oracle Solaris Cluster 进行初始配置时，该软件会为群集中的每个原始设备自动配置设备组。请使用以下过程重新配置这些自动创建的设备组，以便在 Oracle Solaris Cluster 软件中使用。

基于以下原因新建原始磁盘类型的设备组：

- 您要将多个 DID 添加至设备组
- 您需要更改设备组的名称。
- 您要在不使用 `cldevicegroup` 命令的 `-v` 选项的情况下创建一组设备组



注意 - 如果要在复制的设备上创建设备组，则您创建的设备组（Solaris Volume Manager 或原始磁盘）必须与复制的设备组同名。

1. 识别要使用的设备，并取消配置为其预定义的设备组。

以下命令可删除为设备 dN 和 dX 预定义的设备组。

```
phys-schost-1# cldevicegroup disable dsk/dN dsk/dX
phys-schost-1# cldevicegroup offline dsk/dN dsk/dX
phys-schost-1# cldevicegroup delete dsk/dN dsk/dX
```

2. 创建包含所需设备的新原始磁盘设备组。

以下命令可创建全局设备组 *raw-disk-dg*，该设备组包含设备 dN 和 dX。

```
phys-schost-1# cldevicegroup create -n phys-schost-1,phys-schost-2 \
-t rawdisk -d dN,dX raw-disk-dg
```

3. 检验您创建的原始磁盘设备组。

```
phys-schost-1# cldevicegroup show raw-disk-dg
```

▼ 如何添加并注册复制设备组 (ZFS)

使用此过程创建由 HAStoragePlus 管理的复制的 ZFS 设备组。

要创建不使用 HAStoragePlus 的 ZFS 存储池 (zpool)，请转至[如何配置无 HAStoragePlus 的本地 ZFS 存储池 \[122\]](#)。

开始之前 要复制 ZFS，必须创建命名的设备组并列出于 zpool 的磁盘。一个设备一次只能属于一个设备组，因此，如果已经存在包含该设备的 Oracle Solaris Cluster 设备组，则必须将该组删除后才能将该设备添加到新的 ZFS 设备组。

您创建的 Oracle Solaris Cluster 设备组（Solaris Volume Manager 或原始磁盘）必须与复制的设备组同名。

1. 删除与 zpool 中设备相对应的默认设备组。

例如，如果 zpool mypool 含有两个设备 /dev/did/dsk/d2 和 /dev/did/dsk/d13，则必须删除 d2 和 d13 这两个默认设备组。

```
# cldevicegroup offline dsk/d2 dsk/d13
# cldevicegroup delete dsk/d2 dsk/d13
```

2. 创建命名的设备组，使其 DID 与[步骤 1](#)中删除的设备组相对应。

```
# cldevicegroup create -n pnode1,pnode2 -d d2,d13 -t rawdisk mypool
```

此操作创建名为 mypool 的设备组（与 zpool 同名），它管理原始设备 /dev/did/dsk/d2 和 /dev/did/dsk/d13。

3. 创建包含这些设备的 zpool。

```
# zpool create mypool mirror /dev/did/dsk/d2 /dev/did/dsk/d13
```

4. 创建用于管理复制设备（设备组中）迁移的资源组，其节点列表中只包含全局区域。

```
# clresourcegroup create -n pnode1,pnode2 migrate_srdfdg-rg
```

5. 在步骤 4 所创建的资源组中创建 hasp-rs 资源，将 globaldevicepaths 属性设置为类型为原始磁盘的设备组。

您在步骤 2 中创建了此设备。

```
# clresource create -t HASToragePlus -x globaldevicepaths=mypool \  
-g migrate_srdfdg-rg hasp2migrate_mypool
```

6. 将此资源组的 rg_affinities 属性的 +++ 值设置为在步骤 4 中创建的资源组。

```
# clresourcegroup create -n pnode1,pnode2 \  
-p RG_affinities=+++migrate_srdfdg-rg oracle-rg
```

7. 在步骤 4 或步骤 6 所创建的资源组中，为步骤 3 所创建的 zpool 创建一个 HASToragePlus 资源 (hasp-rs)。

将 resource_dependencies 属性设置为在步骤 5 中创建的 hasp-rs 资源。

```
# clresource create -g oracle-rg -t HASToragePlus -p zpools=mypool \  
-p resource_dependencies=hasp2migrate_mypool \  
-p ZpoolsSearchDir=/dev/did/dsk hasp2import_mypool
```

8. 需要设备组名称时，请使用这一新资源组名称。

▼ 如何配置无 HASToragePlus 的本地 ZFS 存储池

此过程介绍如何在本地设备上配置 ZFS 存储池 (zpool)，而不配置 HASToragePlus 资源。

注 - 要配置使用 HASToragePlus 资源的本地 zpool，请转至[如何添加并注册复制设备组 \(ZFS\) \[121\]](#)。

1. 列出 DID 映射并标识要使用的本地设备。

选择仅列出将使用新 zpool 的群集节点的设备。请注意 cNtXdY 设备名称和 /dev/did/rdisk/dN DID 设备名称。

```
phys-schost-1# cldevice show | grep Device
```

在下面的示例中，DID 设备 `/dev/did/rdisk/d1` 和 `/dev/did/rdisk/d2` 的条目表明仅 `phys-schost-1` 使用那些驱动器。对于此过程中的示例，将为群集节点 `phys-schost-1` 使用和配置具有设备名称 `c0t6d0` 的 DID 设备 `/dev/did/rdisk/d2`。

```
=== DID Device Instances ===
DID Device Name:                /dev/did/rdisk/d1
Full Device Path:                phys-schost-1:/dev/rdsk/c0t0d0
DID Device Name:                /dev/did/rdisk/d2
Full Device Path:                phys-schost-1:/dev/rdsk/c0t6d0
DID Device Name:                /dev/did/rdisk/d3
Full Device Path:                phys-schost-1:/dev/rdsk/clt1d0
Full Device Path:                phys-schost-2:/dev/rdsk/clt1d0
...
```

2. 确定您为 zpool 选择的 DID 设备的设备组名称。

下面的示例输出显示 `dsk/d2` 是 DID 设备 `/dev/did/rdisk/d2s2` 的设备组名称。设备组名称是 DID 设备名称的一部分，通常会存在这种关系，但不总是这样。此设备组在其节点列表中仅包含一个节点 `phys-schost-1`。

```
phys-schost-1# cldevicegroup show -v
...
Device Group Name:              dsk/d2
Type:                           Disk
failback:                       false
Node List:                      phys-schost-1
preferred:                      false
localonly:                      false
autogen:                        true
numsecondaries:                 1
device names:                   /dev/did/rdisk/d2s2
...
```

3. 为 DID 设备设置 `localonly` 属性。

指定在 [步骤 2](#) 中标识的设备组名称。如果要为设备禁用隔离，还要在命令中包括 `default_fencing=nofencing`。

```
phys-schost-1# cldevicegroup set -p localonly=true \
-p autogen=true [-p default_fencing=nofencing] dsk/d2
```

有关 `cldevicegroup` 属性的更多信息，请参见 [cldevicegroup\(1CL\)](#) 手册页。

4. 检验设备设置。

```
phys-schost-1# cldevicegroup show dsk/d2
```

5. 创建 zpool。

```
phys-schost-1# zpool create localpool c0t6d0
```

6. (可选) 创建 ZFS 数据集。

```
phys-schost-1# zfs create localpool/data
```

7. 检验新 zpool。

```
phys-schost-1# zpool list
```

维护设备组

您可针对设备组执行各种管理任务。其中一些任务还可以使用 Oracle Solaris Cluster Manager 浏览器界面执行。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

如何删除和取消注册设备组 (Solaris Volume Manager)

设备组是已经注册到 Oracle Solaris Cluster 的 Solaris Volume Manager 磁盘集。要删除一个 Solaris Volume Manager 设备组，请使用 `metaclear` 和 `metaset` 命令。这些命令可删除同名的设备组，并取消注册该磁盘组，使之不再是 Oracle Solaris Cluster 设备组。

有关删除磁盘集的步骤，请参阅《[Solaris Volume Manager 管理指南](#)》。

▼ 如何将节点从所有设备组中删除

使用此过程可将一个群集节点从所有将该节点列为潜在主节点的设备组中删除。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在要删除的作为所有设备组的潜在主节点的节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。

2. 确定成员是要删除的节点的一个或多个设备组。

在每个设备组的 Device group node list 中查找该节点名称。

```
# cldevicegroup list -v
```

3. 如果[步骤 2](#)中标识的设备组中有任何 `svm` 类型的设备组，请对每个该类型的设备组执行[如何将节点从设备组中删除 \(Solaris Volume Manager\) \[125\]](#)中的步骤。

4. 确定成员是要删除的节点的原始设备磁盘组。

```
# cldevicegroup list -v
```

5. 如果步骤 4 中所列设备组中有任何 **Disk** 或 **Local_Disk** 类型的设备组，请对所有这些设备组执行[如何从原始磁盘设备组删除节点 \[127\]](#)中的步骤。
6. 检验是否已将该节点从所有设备组的潜在主节点列表中删除。
如果该节点不再被列为任何设备组的潜在主节点，则以下命令不返回任何内容。

```
# cldevicegroup list -v nodename
```

▼ 如何将节点从设备组中删除 (Solaris Volume Manager)

使用此过程可将一个群集节点从 Solaris Volume Manager 设备组的潜在主节点列表中删除。对每个要从中删除该节点的设备组执行 `metaset` 命令。



注意 - 如果其他节点是活动群集成员并且至少其中的一个节点拥有磁盘集，则请勿在群集之外引导的群集节点上运行 `metaset -s setname -f -t`。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 检验该节点是否仍是设备组的成员，且该设备组是否是 Solaris Volume Manager 设备组。

设备组类型 `sds/svm` 指示 Solaris Volume Manager 设备组。

```
phys-schost-1% cldevicegroup show devicegroup
```

2. 确定哪个节点是设备组当前的主节点。

```
# cldevicegroup status devicegroup
```

3. 在当前拥有要修改的设备组的节点上，承担 **root** 角色。
4. 将该节点的主机名从设备组中删除。

```
# metaset -s setname -d -h nodelist
```

`-s setname` 指定设备组的名称。

`-d` 从设备组中删除以 `-h` 标识的节点。

`-h nodelist` 指定将要删除的一个或多个节点的节点名称。

注 - 完成更新可能需要几分钟。

如果该命令失败，请在命令中增加 `-f`（强制）选项。

```
# metaset -s setname -d -f -h nodelist
```

5. 对每个将要从中删除作为潜在主节点的节点的设备组执行[步骤 4](#)。

6. 检验该节点是否已从设备组中删除。

设备组名称与使用 `metaset` 命令指定的磁盘集名称相符。

```
phys-schost-1% cldevicegroup list -v devicegroup
```

例 38 从设备组中删除一个节点 (Solaris Volume Manager)

下面的示例显示了如何从设备组配置中删除主机名 `phys-schost-2`。本示例消除了 `phys-schost-2` 成为指定设备组的潜在主节点的可能性。可运行 `cldevicegroup show` 命令检验节点是否已删除。检查删除的节点是否不再显示在屏幕文本中。

为节点确定 *Solaris Volume Manager* 设备组

```
# cldevicegroup show dg-schost-1
```

```
=== Device Groups ===
```

```
Device Group Name:      dg-schost-1
Type:                   SVM
failback:               no
Node List:              phys-schost-1, phys-schost-2
preferenced:            yes
numsecondaries:         1
diskset name:           dg-schost-1
```

确定哪个节点是设备组当前的主节点

```
# cldevicegroup status dg-schost-1
```

```
=== Cluster Device Groups ===
```

```
--- Device Group Status ---
```

Device Group Name	Primary	Secondary	Status
dg-schost-1	phys-schost-1	phys-schost-2	Online

在当前拥有设备组的节点上，承担 *root* 角色

从设备组中删除主机名

```
# metaset -s dg-schost-1 -d -h phys-schost-2
```

检验是否删除了节点

```
phys-schost-1% cldevicegroup list -v dg-schost-1
```

```
=== Cluster Device Groups ===
```

```
--- Device Group Status ---
```

Device Group Name	Primary	Secondary	Status
-----	-----	-----	-----
dg-schost-1	phys-schost-1	-	Online

▼ 如何从原始磁盘设备组删除节点

使用此过程可将一个群集节点从原始磁盘设备组的潜在主节点列表中删除。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的某一节点（要删除的节点除外）上，承担可提供 `solaris.cluster.read` 和 `solaris.cluster.modify` RBAC 授权的角色。
2. 找出与正在删除的节点相连的设备组，并确定哪些是原始磁盘设备组。

```
# cldevicegroup show -n nodename -t rawdisk +
```

3. 禁用每个 Local_Disk 原始磁盘设备组的 `localonly` 属性。

```
# cldevicegroup set -p localonly=false devicegroup
```

有关 `localonly` 属性的更多信息，请参见 `cldevicegroup(1CL)` 手册页。

4. 检验是否已禁用与要删除的节点相连的所有原始磁盘设备组的 `localonly` 属性。
Disk 设备组类型指示已禁用该原始磁盘设备组的 `localonly` 属性。

```
# cldevicegroup show -n nodename -t rawdisk -v +
```

5. 将节点从步骤 2 所找出的所有原始磁盘设备组中删除。
必须为每一个与正在删除的节点相连接的原始磁盘设备组完成此步骤。

```
# cldevicegroup remove-node -n nodename devicegroup
```

例 39 从原始设备组中删除节点

本示例说明如何从原始磁盘设备组中删除节点 (phys-schost-2)。所有命令均是从该群集的另一节点 (phys-schost-1) 上运行的。

找出与正在删除的节点相连的设备组，
并确定哪些是原始磁盘设备组

```
phys-schost-1# cldevicegroup show -n phys-schost-2 -t rawdisk -v +
```

```

Device Group Name:                    dsk/d4
Type:                                Disk
failback:                            false
Node List:                           phys-schost-2
preferenced:                          false
localonly:                           false
autogen                              true
numsecondaries:                       1
device names:                         phys-schost-2
Device Group Name:                    dsk/d1
Type:                                SVM
failback:                            false
Node List:                           pbravel1, pbrave2
preferenced:                          true
localonly:                           false
autogen                              true
numsecondaries:                       1
diskset name:                         ms1
(dsk/d4) Device group node list:      phys-schost-2
(dsk/d2) Device group node list:      phys-schost-1, phys-schost-2
(dsk/d1) Device group node list:      phys-schost-1, phys-schost-2
    为节点上的每个本地磁盘禁用 localonly 标志
phys-schost-1# cldevicegroup set -p localonly=false dsk/d4

    检验是否禁用了 localonly 标志 phys-schost-1# cldevicegroup show -n phys-schost-2 -t
rawdisk +
(dsk/d4) Device group type:            Disk
(dsk/d8) Device group type:            Local_Disk

    从所有原始磁盘设备组中删除节点
phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d4
phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d2
phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d1

```

▼ 如何更改设备组属性

用于确立设备组的主所有权的方法基于一个名为 *preferenced* 的所有权首选属性的设置。如果未设置该属性，其他无主 (*unowned*) 设备组的主所有者便是第一个尝试访问该组中磁盘的节点。但是，如果设置了该属性，您必须指定节点尝试建立所有权时采用的首选顺序。

如果禁用 *preferenced* 属性，则 *failback* 属性也将自动被禁用。但是，如果尝试启用或重新启用 *preferenced* 属性，则可以选择启用或禁用 *failback* 属性。

如果启用或重新启用了 *preferenced* 属性，则需要重新排列主所有权首选列表中节点的顺序。

此过程使用 *clsetup* 实用程序设置或取消设置 Solaris Volume Manager 设备组的 *preferenced* 属性和 *failback* 属性。

开始之前 要执行此过程，您需要知道要更改其属性值的设备组的名称。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.read` 和 `solaris.cluster.modify` RBAC 授权的角色。

2. 启动 `clsetup` 实用程序。

```
# clsetup
```

这时将显示主菜单。

3. 要操作设备组，请键入与设备组和卷选项对应的编号。

这时将显示 "Device Groups" (设备组) 菜单。

4. 要更改设备组的主要属性，请键入用于更改 Solaris Volume Manager 设备组主要属性的选项对应的编号。

这时将显示 "Change Key Properties" (更改主要属性) 菜单。

5. 要更改某个设备组属性，请键入用于更改 `preferenced` 或 `failback` 属性的选项对应的编号。

按照说明为设备组设置 `preferenced` 和 `failback` 选项。

6. 检验设备组属性是否已更改。

通过以下命令查看所显示的磁盘设备组信息。

```
# cldevicegroup show -v devicegroup
```

例 40 更改设备组属性

以下示例显示了当 `clsetup` 对设备组 (`dg-schost-1`) 的属性值进行设置时所生成的 `cldevicegroup` 命令。

```
# cldevicegroup set -p preferenced=true -p failback=true -p numsecondaries=1 \
-p nodelist=phys-schost-1,phys-schost-2 dg-schost-1
# cldevicegroup show dg-schost-1
```

```
=== Device Groups ===
```

Device Group Name:	dg-schost-1
Type:	SVM
failback:	yes
Node List:	phys-schost-1, phys-schost-2

```
preferred:                yes
numsecondaries:           1
diskset names:            dg-schost-1
```

▼ 如何设置设备组所需的辅助节点数

`numsecondaries` 属性指定在主节点发生故障后设备组中可以控制该设备组的节点数。设备服务默认的辅助节点数为 1。您可将该值设置为一与设备组中有效非主提供节点的数目之间的任意整数。

该设置是平衡群集性能和可用性的一个重要因素。例如，增大所需的辅助节点数可以增大设备组在群集中同时发生多处故障时正常运行的机率。增大辅助节点数通常还会有规律地降低正常运行时的性能。一般情况下，辅助节点数越少，性能越好，但是可用性越差。但是，辅助节点数多并不一定会提高出现问题的文件系统或设备组的可用性。有关更多信息，请参阅《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》中的第 3 章, “[Key Concepts for System Administrators and Application Developers](#)”。

如果更改 `numsecondaries` 属性，则一旦此更改导致实际辅助节点数与所需辅助节点数不一致，系统将向设备组添加或从中删除辅助节点。

此过程使用 `clsetup` 实用程序为所有类型的设备组设置 `numsecondaries` 属性。有关配置任意设备组时的设备组选项的信息，请参阅 `cldevicegroup(1CL)` 手册页。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

1. 在群集的任一节点上承担可提供 `solaris.cluster.read` 和 `solaris.cluster.modify` RBAC 授权的角色。
2. 启动 `clsetup` 实用程序。

`clsetup`

这时将显示主菜单。
3. 要使用磁盘组，请选择 "Device Groups and Volumes"（设备组和卷）菜单项。
这时将显示 "Device Groups"（设备组）菜单。
4. 要更改设备组的主要属性，请选择 "Change Key Properties of a Device Group"（更改设备组的主要属性）菜单项。
这时将显示 "Change Key Properties"（更改主要属性）菜单。
5. 要更改所需的辅助节点数，请键入用于更改 `numsecondaries` 属性的选项对应的编号。
按照说明进行操作，并键入要为设备组配置的辅助节点数。此时，将执行相应的 `cldevicegroup` 命令并输出一条日志，然后返回到前一菜单。

6. 验证设备组的配置。

```
# cldevicegroup show dg-schost-1
=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      Local_Disk
failback:                  yes
Node List:                  phys-schost-1, phys-schost-2, phys-schost-3
preferenced:                yes
numsecondaries:             1
diskgroup names:           dg-schost-1
```

注 - 设备组配置更改包括添加或删除卷，以及更改现有卷的组、所有者或权限。配置更改后的注册将确保全局名称空间处于正确的状态。请参见[如何更新全局设备名称空间 \[114\]](#)。

7. 检验设备组属性是否已更改。

通过以下命令查看所显示的磁盘设备组信息。

```
# cldevicegroup show -v devicegroup
```

例 41 更改所需的辅助节点数 (Solaris Volume Manager)

以下示例显示了当 `clsetup` 为设备组 (`dg-schost-1`) 配置所需的辅助节点数时所生成的 `cldevicegroup` 命令。此示例假定磁盘组和卷是以前创建的。

```
# cldevicegroup set -p numsecondaries=1 dg-schost-1
# cldevicegroup show -v dg-schost-1

=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      SVM
failback:                  yes
Node List:                  phys-schost-1, phys-schost-2
preferenced:                yes
numsecondaries:             1
diskset names:             dg-schost-1
```

例 42 将所需的辅助节点数设置为默认值

以下示例显示如何使用空字符串值配置默认的辅助节点数。设备组将配置为使用该默认值，即使默认值改变。

```
# cldevicegroup set -p numsecondaries= dg-schost-1
# cldevicegroup show -v dg-schost-1

=== Device Groups ===
```

```
Device Group Name:      dg-schost-1
Type:                  SVM
failback:              yes
Node List:             phys-schost-1, phys-schost-2 phys-schost-3
preferenced:           yes
numsecondaries:        1
diskset names:         dg-schost-1
```

▼ 如何列出设备组配置

您无需成为 root 角色即可列出配置。但您需要具备 `solaris.cluster.read` 授权。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- 使用以下列表中的方法之一。

Oracle Solaris Cluster Manager 浏览器界面

有关更多信息，请参见[第 13 章 使用 Oracle Solaris Cluster Manager 浏览器界面](#)。

```
cldevicegroup show
```

使用 `cldevicegroup show` 可列出群集中所有设备组的配置。

```
cldevicegroup show devicegroup
```

使用 `cldevicegroup show devicegroup` 可列出单个设备组的配置。

```
cldevicegroup status devicegroup
```

使用 `cldevicegroup status devicegroup` 可确定单个设备组的状态。

```
cldevicegroup status +
```

使用 `cldevicegroup status +` 可确定群集中所有设备组的状态。

在这些命令中使用 `-v` 选项可获取更为详细的信息。

例 43 列出所有设备组的状态

```
# cldevicegroup status +

=== Cluster Device Groups ===

--- Device Group Status ---

Device Group Name    Primary    Secondary    Status
-----
```

dg-schost-1	phys-schost-2	phys-schost-1	Online
dg-schost-2	phys-schost-1	--	Offline
dg-schost-3	phys-schost-3	phy-shost-2	Online

例 44 列出特定设备组的配置

```
# cldevicegroup show dg-schost-1

=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      SVM
failback:                  yes
Node List:                  phys-schost-2, phys-schost-3
preferenced:                yes
numsecondaries:             1
diskset names:              dg-schost-1
```

▼ 如何切换设备组的主节点

此过程还可以用于启动不活动的设备组（使之联机）。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面使不活动的设备组联机。有关更多信息，请参见 Oracle Solaris Cluster Manager 联机帮助。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 使用 `cldevicegroup switch` 切换设备组的主节点。

```
# cldevicegroup switch -n nodename devicegroup
```

`-n nodename` 指定要切换到的节点的名称。此节点成为新的主节点。

`devicegroup` 指定要切换的设备组。

3. 检验设备组是否已切换到新的主节点上。
如果正确注册了设备组，则使用以下命令时将显示新设备组的信息。

```
# cldevice status devicegroup
```

例 45 切换设备组的主节点

以下示例显示了如何切换设备组的主节点以及如何检验此更改。

```
# cldevicegroup switch -n phys-schost-1 dg-schost-1

# cldevicegroup status dg-schost-1

=== Cluster Device Groups ===

--- Device Group Status ---

Device Group Name      Primary      Secondary      Status
-----
dg-schost-1            phys-schost-1  phys-schost-2  Online
```

▼ 如何将设备组置于维护状态

将设备组置于维护状态可防止在访问设备组中的某个设备时使设备组自动联机。如果修复过程要求停止所有 I/O 活动，直至修复完成，则在完成修复过程时您应当将设备组置于维护状态。此外，将设备组置于维护状态还可确保当系统在一个节点上修复磁盘集时，另一节点上的设备组不会联机，从而防止数据丢失。

有关如何恢复损坏的磁盘集的说明，请参见[“恢复损坏的磁盘集” \[243\]](#)。

注 - 在将设备组置于维护状态之前，必须停止对其设备的所有访问并且必须卸载所有依赖该设备的文件系统。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面使活动的设备组脱机。有关更多信息，请参见 Oracle Solaris Cluster Manager 联机帮助。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 使设备组处于维护状态。
 - a. 如果启用了设备组，请禁用设备组。

```
# cldevicegroup disable devicegroup
```
 - b. 使设备组脱机。

```
# cldevicegroup offline devicegroup
```

2. 如果所执行的修复过程需要磁盘集的所有权，请手动导入该磁盘集。

```
# metaset -C take -f -s diskset
```



注意 - 如果要获取 Solaris Volume Manager 磁盘集的所有权，必须在设备组处于维护状态时使用 `metaset -C take` 命令。使用 `metaset -t` 会将设备组联机，这是获取所有权的一部分。

3. 完成需要执行的修复过程。
4. 释放对磁盘集的所有权。



注意 - 使设备组脱离维护状态之前，必须先释放对该磁盘集的所有权。如果释放所有权失败，可导致数据丢失。

```
# metaset -C release -s diskset
```

5. 使设备组联机。

```
# cldevicegroup online devicegroup
# cldevicegroup enable devicegroup
```

例 46 将设备组置于维护状态

本示例说明了如何将设备组 `dg-schost-1` 置于维护状态，以及如何使该设备组脱离维护状态。

```
[使设备组处于维护状态。]
# cldevicegroup disable dg-schost-1
# cldevicegroup offline dg-schost-1
[如果需要，手动导入磁盘集。]
For Solaris Volume Manager:
# metaset -C take -f -s dg-schost-1
[完成所有必需的修复过程。]
[释放所有权。]
对于 Solaris Volume Manager :
# metaset -C release -s dg-schost-1
[使设备组联机。]
# cldevicegroup online dg-schost-1
# cldevicegroup enable dg-schost-1
```

管理存储设备的 SCSI 协议设置

安装 Oracle Solaris Cluster 软件时，系统会自动为所有存储设备分配 SCSI 预留空间。请执行以下过程检查设备的设置，并在必要时覆盖设备的设置：

- [如何显示所有存储设备的默认全局 SCSI 协议设置 \[136\]](#)
- [如何显示单个存储设备的 SCSI 协议 \[136\]](#)
- [如何更改所有存储设备的默认全局隔离协议设置 \[137\]](#)
- [如何更改单个存储设备的隔离协议 \[138\]](#)

▼ 如何显示所有存储设备的默认全局 SCSI 协议设置

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.read` RBAC 授权的角色。
2. 从任意节点上显示当前全局默认 SCSI 协议设置。

```
# cluster show -t global
```

有关更多信息，请参见 [cluster\(1CL\)](#) 手册页。

例 47 显示所有存储设备的默认全局 SCSI 协议设置

以下示例显示了群集上所有存储设备的 SCSI 协议设置。

```
# cluster show -t global
```

```
=== Cluster ===
```

Cluster Name:	racerrxx
clusterid:	0x4FE52C888
installmode:	disabled
heartbeat_timeout:	10000
heartbeat_quantum:	1000
private_netaddr:	172.16.0.0
private_netmask:	255.255.111.0
max_nodes:	64
max_privatenets:	10
udp_session_timeout:	480
concentrate_load:	False
global_fencing:	prefer3
Node List:	phys-racerrxx-1, phys-racerrxx-2

▼ 如何显示单个存储设备的 SCSI 协议

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.read` RBAC 授权的角色。
2. 从任意节点上显示存储设备的 SCSI 协议设置。

```
# cldevice show device
```

`device` 设备路径的名称或设备名称。

有关更多信息，请参见 [cldevice\(1CL\)](#) 手册页。

例 48 显示单个设备的 SCSI 协议

以下示例显示了设备 `/dev/rdisk/c4t8d0` 的 SCSI 协议。

```
# cldevice show /dev/rdisk/c4t8d0
```

```
=== DID Device Instances ===
```

DID Device Name:	/dev/did/rdsk/d3
Full Device Path:	phappy1:/dev/rdsk/c4t8d0
Full Device Path:	phappy2:/dev/rdsk/c4t8d0
Replication:	none
default_fencing:	global

▼ 如何更改所有存储设备的默认全局隔离协议设置

您可以针对连接到某个群集的所有存储设备全局性地打开或关闭隔离功能。如果单个存储设备的默认隔离值设置为 `pathcount`、`prefer3` 或 `nofencing`，则该设备的默认隔离设置将覆盖全局设置。如果存储设备的默认隔离值设置为 `global`，该存储设备将使用全局设置。例如，如果存储设备的默认设置为 `pathcount`，则当您执行以下过程将全局 SCSI 协议设置更改为 `prefer3` 时，该存储设备的设置不会更改。您必须执行[如何更改单个存储设备的隔离协议 \[138\]](#)中的过程来更改单个设备的默认设置。



注意 - 如果在错误的情况下关闭了隔离功能，则您的数据在应用程序故障转移过程中易于损坏。当您考虑关闭隔离功能时，请仔细分析此数据损坏的可能性。如果共享存储设备不支持 SCSI 协议，或者您想要允许从群集外部的本机访问群集的存储，则可以关闭隔离功能。

要更改某个法定设备的默认隔离设置，必须先取消配置该设备，更改其隔离设置，然后再重新配置该法定设备。如果您计划为包括法定设备在内的设备定期关闭和重新打开隔离功能，应考虑通过法定服务器服务来配置法定，以避免在法定操作中出现中断。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 为所有不是法定设备的存储设备设置隔离协议。

```
cluster set -p global_fencing={pathcount | prefer3 | nofencing | nofencing-noscrub}
```

```
-p global_fencing
```

为所有共享设备设置当前全局默认隔离算法。

```
prefer3
```

对具有两个以上路径的设备使用 SCSI-3 协议。

```
pathcount
```

根据连接到共享设备的 DID 路径的数目来确定隔离协议。pathcount 设置用于法定设备。

```
nofencing
```

通过设置所有存储设备的隔离状态来关闭隔离功能。

```
nofencing-noscrub
```

清理设备可确保设备清除所有持久的 SCSI 保留信息，并且允许从群集外部的系统访问存储。请仅对具有严重的 SCSI 保留问题的存储设备使用 nofencing-noscrub 选项。

例 49 为所有存储设备设置默认全局隔离协议设置

以下示例将群集中所有存储设备的隔离协议设置为 SCSI-3 协议。

```
# cluster set -p global_fencing=prefer3
```

▼ 如何更改单个存储设备的隔离协议

您还可以设置单个存储设备的隔离协议。

注 - 要更改某个法定设备的默认隔离设置，必须先取消配置该设备，更改其隔离设置，然后再重新配置该法定设备。如果您计划为包括法定设备在内的设备定期关闭和重新打开隔离功能，应考虑通过法定服务器服务来配置法定，以避免在法定操作中出现中断。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。



注意 - 如果在错误的情况下关闭了隔离功能，则您的数据在应用程序故障转移过程中易于损坏。当您考虑关闭隔离功能时，请仔细分析此数据损坏的可能性。如果共享存储设备不支持 SCSI 协议，或者您想要允许从群集外部的宿主访问群集的存储，则可以关闭隔离功能。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 设置存储设备的隔离协议。

```
cldevice set -p default_fencing ={pathcount | \
scsi3 | global | nofencing | nofencing-noscrub} device
```

-p default_fencing

修改设备的 default_fencing 属性。

pathcount

根据连接到共享设备的 DID 路径的数目来确定隔离协议。

scsi3

使用 SCSI-3 协议。

global

使用全局默认隔离设置。global 设置用于非法定设备。

nofencing

通过设置指定 DID 实例的隔离状态可关闭隔离功能。

nofencing-noscrub

清理设备可确保设备清除所有持久的 SCSI 保留信息，并且允许从群集外部的系统访问存储设备。请仅对具有严重的 SCSI 保留问题的存储设备使用 nofencing-noscrub 选项。

device

指定设备路径的名称或设备名称。

有关更多信息，请参见 [cluster\(1CL\)](#) 手册页。

例 50 设置单个设备的隔离协议

以下示例为设备 d5（由设备编号指定）设置了 SCSI-3 协议。

```
# cldevice set -p default_fencing=prefer3 d5
```

以下示例为 d11 设备关闭了默认隔离功能。

```
#cldevice set -p default_fencing=nofencing d11
```

管理群集文件系统

群集文件系统是通过全局方式使用的文件系统，可以从群集的任一节点对其进行读取或访问。

表 8 任务列表：管理群集文件系统

任务	说明
在初始 Oracle Solaris Cluster 安装后添加群集文件系统	如何添加群集文件系统 [140]
删除群集文件系统	如何删除群集文件系统 [142]
检查群集中各个节点上的全局挂载点是否一致	如何检查群集中的全局挂载点 [144]

▼ 如何添加群集文件系统

首次安装 Oracle Solaris Cluster 后，为所创建的每个群集文件系统执行此任务。



注意 - 请确保指定了正确的磁盘设备名称。创建群集文件系统会损坏磁盘上的所有数据。如果指定的设备名称不正确，则会擦除您可能并不打算删除的数据。

在添加其他群集文件系统之前，请确保具备以下先决条件：

- 已在群集的一个节点上建立 root 角色特权。
- 已在群集中安装并配置卷管理器软件。
- 存在可在其上创建群集文件系统的一个设备组（例如 Solaris Volume Manager 设备组）或块磁盘分片。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面将群集文件系统添加到区域群集。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

如果曾使用 Oracle Solaris Cluster Manager 安装数据服务，假如用以创建群集文件系统的共享磁盘充足，则系统中已存在一个或多个群集文件系统。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- 1. 在群集中的任一节点上承担 root 角色。

提示 - 要更快地创建文件系统，请在您要为其创建文件系统的全局设备的当前主节点上成为 root 角色。

- 2. 使用 newfs 命令创建一个 UFS 文件系统。



注意 - 创建文件系统时，会毁坏该磁盘上的所有数据。请确保指定的磁盘设备名称是正确的。如果指定的设备名称不正确，可能会删除不打算删除的数据。

phys-schost# newfs raw-disk-device

下表显示了 raw-disk-device 参数的名称示例。请注意，卷管理器的命名规则各不相同。

卷管理器	磁盘设备名称样例	描述
Solaris Volume Manager	/dev/md/nfs/rdisk/d1	nfs 磁盘集中的原始磁盘设备 d1
无	/dev/global/rdisk/d1s3	原始磁盘设备 d1s3

- 3. 在群集中的每个节点上，为群集文件系统创建一个挂载点目录。
每个节点上都需要一个挂载点，即使不在该节点上访问群集文件系统也是如此。

提示 - 为了便于管理，请在 /global/device-group/ 目录中创建挂载点。该位置允许您很容易地区别群集文件系统，这些文件系统从本地文件系统中全局可用。

phys-schost# mkdir -p /global/device-group/mount-point/

device-group

与包含该设备的设备组的名称相对应的目录名。

mount-point

要在其上挂载群集文件系统的目录的名称。

- 4. 在群集中的每个节点上，在 /etc/vfstab 文件中为挂载点添加一个条目。
有关详细信息，请参见 vfstab(4) 手册页。

- a. 在每个条目中，指定所用文件系统类型所需的挂载选项。
- b. 要自动挂载群集文件系统，请将 `mount at boot` 字段设置为 `yes`。
- c. 对于每个群集文件系统，请确保其 `/etc/vfstab` 条目中的信息在每个节点上是完全相同的。
- d. 请确保每个节点的 `/etc/vfstab` 文件中的条目都以相同顺序列出设备。
- e. 检查文件系统的引导顺序依赖性。

例如，考虑如下情形：phys-schost-1 将磁盘设备 d0 挂载到 `/global/oracle/` 上，phys-schost-2 将磁盘设备 d1 挂载到 `/global/oracle/logs/` 上。根据此配置，只有在 phys-schost-1 引导并挂载了 `/global/oracle/` 之后，phys-schost-2 才能引导并挂载 `/global/oracle/logs/`。

5. 在群集中的任一节点上，运行配置检查实用程序。

```
phys-schost# cluster check -k vfstab
```

配置检查实用程序将检验挂载点是否存在。该实用程序还将检验群集的所有节点上的 `/etc/vfstab` 文件条目是否正确。如果没有错误发生，则不返回任何输出。

有关更多信息，请参见 [cluster\(1CL\)](#) 手册页。

6. 从群集中的任何节点挂载群集文件系统。

```
phys-schost# mount /global/device-group/mountpoint/
```

7. 在群集的每个节点上，验证是否已挂载了群集文件系统。

可以使用 `df` 命令或 `mount` 命令列出已挂载的文件系统。有关更多信息，请参见 [df\(1M\)](#) 手册页或 [mount\(1M\)](#) 手册页。

▼ 如何删除群集文件系统

您只需卸载群集文件系统就可以将其删除。如果还要移除或删除数据，请从系统中删除底层的磁盘设备（或元设备或卷）。

注 - 当您运行 `cluster shutdown` 来停止整个群集时，作为系统关闭过程的一部分，群集文件系统会自动卸载。运行 `shutdown` 来停止单个节点时，将不会卸载群集文件系统。但是，如果只有正关闭的节点与磁盘相连，则对该磁盘上的群集文件系统进行的任何访问尝试均会导致出错。

在卸载群集文件系统之前，请确保具备以下先决条件：

- 已在群集的一个节点上建立 root 角色特权。
- 文件系统不能处于忙状态。如果有用户在文件系统的某个目录下工作，或有程序打开了该文件系统中的某个文件，则该文件系统被认为处于忙状态。这个用户或程序可能运行在群集中的任一节点上。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面删除区域群集文件系统。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担 root 角色。

2. 确定已安装的群集文件系统。

```
# mount -v
```

3. 在每个节点上，列出所有正在使用群集文件系统的进程，以便确定要停止哪些进程。

```
# fuser -c [ -u ] mountpoint
```

-c 报告有关用作文件系统的安装点的文件及安装的文件系统中任何文件的信息。

-u (可选) 显示每个进程 ID 的用户登录名称。

mountpoint 指定您要停止其进程的群集文件系统的名称。

4. 在每个节点上，停止群集文件系统的所有进程。

请使用停止进程的首选方法。根据需要，请使用以下命令强制终止与群集文件系统相关的进程。

```
# fuser -c -k mountpoint
```

系统将向每个使用群集文件系统的进程发出 SIGKILL 命令。

5. 在每个节点上，确保无任何进程正在使用群集文件系统。

```
# fuser -c mountpoint
```

6. 仅从一个节点卸载文件系统。

```
# umount mountpoint
```

mountpoint 指定要卸载的群集文件系统的名称。该名称既可以是安装群集文件系统的目录的名称，也可以是文件系统的设备名称路径。

7. (可选) 编辑 `/etc/vfstab` 文件以删除要被删除的群集文件系统的条目。
对于任何群集节点，只要其 `/etc/vfstab` 文件中有此群集文件系统的条目，就要在该群集节点上执行此步骤。
8. (可选) 删除磁盘设备 `group/metadevice/volume/plex`。
有关详细信息，请参阅卷管理器文档。

例 51 删除群集文件系统

以下示例删除了挂载在 Solaris Volume Manager 元设备或卷 `/dev/md/oracle/rdisk/d1` 上的 UFS 群集文件系统。

```
# mount -v
...
/global/oracle/d1 on /dev/md/oracle/dsk/d1 read/write/setuid/global/logging/largefiles
# fuser -c /global/oracle/d1
/global/oracle/d1: 4006c
# fuser -c -k /global/oracle/d1
/global/oracle/d1: 4006c
# fuser -c /global/oracle/d1
/global/oracle/d1:
# umount /global/oracle/d1

    在每个节点上，删除突出显示的条目
# pfedit /etc/vfstab
#device          device          mount  FS      fsck    mount  mount
#to mount        to fsck        point  type    pass   at boot options
#
/dev/md/oracle/dsk/d1 /dev/md/oracle/rdisk/d1 /global/oracle/d1 ufs 2 yes global,logging
```

保存并退出

要删除群集文件系统中的数据，请删除基础设备。有关详细信息，请参阅卷管理器文档。

▼ 如何检查群集中的全局挂载点

`cluster(1CL)` 实用程序检验 `/etc/vfstab` 文件中群集文件系统条目的语法。如果没有错误，则不返回任何内容。

注 - 进行了影响设备或卷管理组件的群集配置更改（如删除群集文件系统）后，请运行 `cluster check` 命令。

1. 在群集中的任一节点上承担 `root` 角色。
2. 检查群集全局安装。


```
# cluster check -k vfstab
```

管理磁盘路径监视

通过磁盘路径监视 (Disk Path Monitoring, DPM) 管理命令，可以接收辅助磁盘路径故障的通知。使用本节中的过程执行与监视磁盘路径关联的管理任务。

提供了以下附加信息：

主题	信息
有关磁盘路径监视守护进程的概念信息	《Oracle Solaris Cluster 4.3 Concepts Guide》中的第 3 章, “Key Concepts for System Administrators and Application Developers”
cldevice 命令选项和相关命令的描述	cldevice(1CL) 手册页
优化 scdpmd 守护进程	scdpmd.conf(4) 手册页
syslogd 守护进程报告的已记录错误	syslogd(1M) 手册页

注 - 当使用 cldevice 命令将 I/O 设备添加到节点中时，磁盘路径会自动添加到受监视的监视列表中。使用 Oracle Solaris Cluster 命令从节点删除设备时，也将自动取消监视磁盘路径。

表 9 任务列表：管理磁盘路径监视

任务	说明
监视磁盘路径。	如何监视磁盘路径 [146]
取消监视磁盘路径。	如何取消监视磁盘路径 [147]
显示节点故障磁盘路径的状态。	如何打印故障磁盘路径 [148]
从文件中监视磁盘路径。	如何从文件监视磁盘路径 [149]
启用或禁用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能。	如何启用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能 [150]
	如何禁用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能 [151]
解决错误的磁盘路径状态。如果引导时被监视 DID 设备不可用则可能会报告错误的磁盘路径状态，从而导致 DID 实例不上载到 DID 驱动程序。	如何解决磁盘路径状态错误 [148]

下一节介绍的过程将发出包含磁盘路径参数的 cldevice 命令。磁盘路径参数由一个节点名称和一个磁盘名称组成。节点名称不是必需的。如果不指定，它将采用默认值 all。

▼ 如何监视磁盘路径

执行此任务可以监视群集中的磁盘路径。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面启用磁盘路径监视。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 监视磁盘路径。

```
# cldevice monitor -n node disk
```

3. 检验是否已监视磁盘路径。

```
# cldevice status device
```

例 52 监视单个节点上的磁盘路径

以下示例监视单个节点的 `schost-1:/dev/did/rdisk/d1` 磁盘路径。只有节点 `schost-1` 上的 DPM 守护进程监视磁盘路径 `/dev/did/dsk/d1`。

```
# cldevice monitor -n schost-1 /dev/did/dsk/d1
# cldevice status d1
```

Device Instance	Node	Status

/dev/did/rdisk/d1	phys-schost-1	Ok

例 53 监视所有节点上的磁盘路径

以下示例监视所有节点的磁盘路径 `schost-1:/dev/did/dsk/d1`。DPM 可以在视 `/dev/did/dsk/d1` 为有效路径的所有节点上启动。

```
# cldevice monitor /dev/did/dsk/d1
# cldevice status /dev/did/dsk/d1
```

Device Instance	Node	Status

/dev/did/rdisk/d1	phys-schost-1	Ok

例 54 从 CCR 重新读取磁盘配置

以下示例强制守护进程从 CCR 重新读取磁盘配置并打印监视的磁盘路径及其状态。

```
# cldevice monitor +
# cldevice status
Device Instance      Node      Status
-----
/dev/did/rdisk/d1    schost-1  Ok
/dev/did/rdisk/d2    schost-1  Ok
/dev/did/rdisk/d3    schost-1  Ok
                        schost-2  Ok
/dev/did/rdisk/d4    schost-1  Ok
                        schost-2  Ok
/dev/did/rdisk/d5    schost-1  Ok
                        schost-2  Ok
/dev/did/rdisk/d6    schost-1  Ok
                        schost-2  Ok
/dev/did/rdisk/d7    schost-2  Ok
/dev/did/rdisk/d8    schost-2  Ok
```

▼ 如何取消监视磁盘路径

使用以下过程可以取消监视磁盘路径。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面禁用磁盘路径监视。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

- 1. 在群集中的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
- 2. 确定要取消监视的磁盘路径的状态。
`# cldevice status device`
- 3. 在每个节点上，取消监视相应的磁盘路径。
`# cldevice unmonitor -n node disk`

例 55 取消监视磁盘路径

以下示例取消监视 `schost-2:/dev/did/rdisk/d1` 磁盘路径并显示整个群集的磁盘路径及其状态。

```
# cldevice unmonitor -n schost2 /dev/did/rdisk/d1
# cldevice status -n schost2 /dev/did/rdisk/d1
```

Device Instance	Node	Status
-----	----	-----
/dev/did/rdisk/d1	schost-2	Unmonitored

▼ 如何打印故障磁盘路径

使用以下步骤可以打印群集的故障磁盘路径。

1. 在群集中的任一节点上承担 **root** 角色。
2. 打印整个群集中故障磁盘路径。

```
# cldevice status -s fail
```

例 56 打印故障磁盘路径

以下示例打印整个群集的故障磁盘路径。

```
# cldevice status -s fail
```

Device Instance	Node	Status
-----	----	-----
dev/did/dsk/d4	phys-schost-1	fail

▼ 如何解决磁盘路径状态错误

如果发生以下事件，重新联机时 DPM 可能不会更新故障路径的状态：

- 被监视路径故障导致重新引导节点。
- 在重新引导的节点重新联机前，被监视 DID 路径下的设备不会重新联机。

之所以会报告错误的磁盘路径状态，是因为引导时被监视的 DID 设备不可用，因而导致 DID 实例未上载到 DID 驱动程序。出现这种情况时，请手动更新 DID 信息。

1. 从一个节点上，更新全局设备名称空间。

```
# cldevice populate
```

2. 在每个节点上检验命令处理过程是否已完成，然后再继续进行下一步骤。

即使仅从一个节点运行，该命令也会以远程方式在所有的节点上执行。要确定该命令是否已完成处理过程，请在群集中的每个节点上运行以下命令。

```
# ps -ef | grep cldevice populate
```

3. 在 DPM 轮询时间帧中检验故障磁盘路径的状态现在是否为 Ok。

```
# cldevice status disk-device
```

Device Instance	Node	Status
-----	----	-----
dev/did/dsk/dN	phys-schost-1	Ok

▼ 如何从文件监视磁盘路径

使用以下步骤监视或取消监视文件的磁盘路径。

要使用文件来更改群集配置，必须首先导出当前配置。此导出操作会创建一个 XML 文件。您可稍后修改该文件来设置要更改的配置项。本过程中的说明描述了上述整个过程。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集中的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 将设备配置导出到一个 XML 文件。

```
# cldevice export -o configurationfile
```

-o 指定 XML 文件的文件名。
configurationfile

3. 修改配置文件，以便对设备路径进行监视。
找到要监视的设备路径，将 `monitored` 属性设置为 `true`。

4. 监视设备路径。

```
# cldevice monitor -i configurationfile
```

-i 指定已修改的 XML 文件的文件名。
configurationfile

5. 检验设备路径此时是否受监视。

```
# cldevice status
```

例 57 从文件监视磁盘路径

在下面的示例中，使用一个 XML 文件对节点 `phys-schost-2` 和设备 `d3` 之间的设备路径进行监视。deviceconfig XML 文件显示 `phys-schost-2` 与 `d3` 之间的路径当前未受监视。

导出当前群集配置

```
# cldevice export -o deviceconfig
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
...
<deviceList readonly="true">
<device name="d3" ctd="c1t8d0">
<devicePath nodeRef="phys-schost-1" monitored="true"/>
<devicePath nodeRef="phys-schost-2" monitored="false"/>
</device>
</deviceList>
</cluster>
```

通过将监视的属性设置为 `true` 来监视路径

```
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
...
<deviceList readonly="true">
<device name="d3" ctd="c1t8d0">
<devicePath nodeRef="phys-schost-1" monitored="true"/>
<devicePath nodeRef="phys-schost-2" monitored="true"/>
</device>
</deviceList>
</cluster>
```

读取文件并打开监视功能

```
# cldevice monitor -i deviceconfig
```

检验设备此时是否受监视

```
# cldevice status
```

另请参见 有关导出群集配置和使用生成的 XML 文件来设置群集配置的更多详细信息，请参见 [cluster\(1CL\)](#) 手册页和 [clconfiguration\(5CL\)](#) 手册页。

▼ 如何启用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能

如果启用了该功能，只要满足以下条件，节点便会自动重新引导：

- 节点上所有受监视的共享磁盘路径均发生故障。

- 至少有一个磁盘可从群集中的其他节点进行访问。

重新引导节点会将该节点管理的所有资源组和设备组在另一个节点上重新启动。

当节点自动重新引导后，如果该节点上所有受监视的共享磁盘路径仍不可访问，该节点不会再次自动重新引导。但是，如果节点重新引导后，有任何磁盘路径先是可用随后又出现故障，则该节点会再次自动重新引导。

如果启用 `reboot_on_path_failure` 属性，在判断是否需要重新引导节点时不会考虑本地磁盘路径的状态。仅受监视的共享磁盘会受影响。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面编辑

`reboot_on_path_failure` 节点属性。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 对于群集中的所有节点，请启用在节点的所有受监视共享磁盘路径均发生故障时自动重新引导的功能。

```
# clnode set -p reboot_on_path_failure=enabled +
```

▼ 如何禁用节点在所有受监视的共享磁盘路径均发生故障时自动重新引导的功能

如果禁用了该功能，当节点上所有受监视的共享磁盘路径均发生故障时，该节点不会自动重新引导。

1. 在群集中的任一节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 对于群集中的所有节点，请禁用节点的所有受监视共享磁盘路径均发生故障时自动重新引导的功能。

```
# clnode set -p reboot_on_path_failure=disabled +
```


管理法定

本章介绍有关在 Oracle Solaris Cluster 中管理法定设备和 Oracle Solaris Cluster 法定服务器的过程。有关法定概念的信息，请参见《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》中的“[Quorum and Quorum Devices](#)”。

- “[管理法定设备](#)” [153]
- “[管理 Oracle Solaris Cluster 法定服务器](#)” [169]

管理法定设备

法定设备是一个共享存储设备或法定服务器，由两个或多个节点共享，并提供用于建立法定的选票。本节介绍有关管理法定设备的过程。

可以使用 `clquorum` 命令执行所有法定设备管理过程。此外，还可以通过使用 `clsetup` 交互式实用程序或 Oracle Solaris Cluster Manager 浏览器界面完成一些过程。有关 Oracle Solaris Cluster Manager 登录说明，请参见《[Oracle Solaris Cluster 4.3 系统管理指南](#)》中的“[如何访问 Oracle Solaris Cluster Manager](#)”。

只要可能，本节中的法定过程都使用 `clsetup` 实用程序来说明。Oracle Solaris Cluster Manager 联机帮助介绍了如何使用 Oracle Solaris Cluster Manager 来执行法定过程。有关更多信息，请参见 `clquorum(1CL)` 和 `clsetup(1CL)` 手册页。

当您使用法定设备时，请牢记以下指导原则：

- 所有法定设备命令都必须在全局群集节点中运行。
- 如果 `clquorum` 命令中断或失败，群集配置数据库中的法定配置信息可能会变得不一致。出现这种不一致情况时，可重新运行该命令或运行 `clquorum reset` 命令重置法定配置。
- 为了最大限度地实现群集高可用性，请确保法定设备投的选票总数少于节点投的选票总数。否则，节点无法在所有法定设备都不可用时形成群集，即使所有节点都在正常运行也是如此。
- 请不要将当前配置为法定设备的磁盘添加到 Oracle Solaris ZFS 存储池中。如果将一个已配置的法定设备添加到 ZFS 存储池中，该磁盘会被重新标为 EFI 磁盘，法定配置信息将丢失，并且该磁盘将不会再向群集提供法定选票。磁盘一旦处于存储池中，

即可被配置为法定设备。或者，您也可以先取消配置磁盘，并将它添加到存储池中，然后将该磁盘重新配置为法定设备。

注 - `clsetup` 命令是一个访问其他 Oracle Solaris Cluster 命令的交互式接口。运行 `clsetup` 时，该命令会生成相应的特定命令，在此例中生成 `clquorum` 命令。这些生成的命令显示在这些过程结尾部分的示例中。

要查看法定配置，请使用 `clquorum show`。`clquorum list` 命令可显示群集中法定设备的名称。`clquorum status` 命令可提供状态和选票计数信息。

本节显示的多数示例均来自一个由三个节点组成的群集。

表 10 任务列表：法定管理

任务	相关说明
使用 <code>clsetup</code> 实用程序向群集添加法定设备	“添加法定设备” [155]
使用 <code>clsetup</code> 实用程序（生成 <code>clquorum</code> ）从群集中删除法定设备	如何删除法定设备 [160]
使用 <code>clsetup</code> 实用程序（生成 <code>clquorum</code> ）从群集中删除最后一个法定设备	如何从群集中删除最后一个法定设备 [161]
使用添加和删除过程替换群集中的法定设备	如何替换法定设备 [163]
使用添加和删除过程修改法定设备列表	如何修改法定设备节点列表 [164]
使用 <code>clsetup</code> 实用程序（生成 <code>clquorum</code> ）将法定设备置于维护状态	如何将法定设备置于维护状态 [164]
（在维护状态下，法定设备不参与选票来建立法定。）	
使用 <code>clsetup</code> 实用程序（生成 <code>clquorum</code> ）将法定设备配置重置为默认状态	如何使法定设备脱离维护状态 [166]
使用 <code>clquorum</code> 命令列出法定设备和选票计数	如何列出法定配置 [167]

动态重新配置法定设备

对群集中的法定设备完成动态重新配置时，必须考虑几个问题。

- 针对 Oracle Solaris 动态重新配置功能介绍的所有要求、过程和限制也适用于 Oracle Solaris Cluster 动态重新配置支持（操作系统停止操作除外）。因此，请先阅读关于 Oracle Solaris 动态重新配置功能的文档，然后再对 Oracle Solaris Cluster 软件使用动态重新配置功能。您应该特别注意那些在执行动态重新配置分离操作时影响非网络 IO 设备的问题。
- Oracle Solaris Cluster 拒绝在为法定设备配置了接口的情况下执行动态重新配置删除板操作。
- 如果动态重新配置操作会影响活动设备，Oracle Solaris Cluster 将拒绝此操作并标识出会受此操作影响的设备。

要删除法定设备，您必须按指示的顺序完成以下步骤。

表 11 任务列表：动态重新配置法定设备

任务	相关说明
1. 启用一个新的法定设备，以替换正要删除的设备。	“添加法定设备” [155]
2. 禁用要删除的法定设备。	如何删除法定设备 [160]
3. 对要删除的设备执行动态重新配置删除操作。	

添加法定设备

本节提供了添加法定设备的过程。在添加新法定设备之前，请确保群集中所有节点均处于联机状态。有关确定群集所需的法定投票计数数量、建议的法定配置和故障隔离的信息，请参见《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》中的“[Quorum and Quorum Devices](#)”。



注意 - 请不要将当前配置为法定设备的磁盘添加到 Solaris ZFS 存储池中。在将一个已配置的法定设备添加到 Solaris ZFS 存储池中的时候，该磁盘会被重新标为 EFI 磁盘，法定配置信息将丢失，并且该磁盘将不会再向群集提供法定选票。磁盘一旦处于存储池中，即可被配置为法定设备。您还可以先取消配置磁盘，并将它添加到存储池中，然后将该磁盘重新配置为法定设备。

Oracle Solaris Cluster 支持以下类型的法定设备：

- 以下各项中的共享 LUN：
 - 共享 SCSI 磁盘
 - 串行连接技术附件 (Serial Attached Technology Attachment, SATA) 存储
 - Oracle ZFS Storage Appliance
 - 支持的 NAS 设备
- Oracle Solaris Cluster 法定服务器

下节中将提供这些设备的添加过程：

- [如何添加共享磁盘法定设备 \[156\]](#)
- [如何添加法定服务器法定设备 \[158\]](#)

注 - 不能将复制磁盘配置为法定设备。如果尝试将一个复制磁盘添加为法定设备，您将收到以下错误消息，命令将退出并显示一个错误代码。

Disk-name is a replicated device. Replicated devices cannot be configured as quorum devices.

共享磁盘法定设备可为 Oracle Solaris Cluster 软件所支持的任何连接存储设备。共享磁盘连接到您的群集的两个或更多个节点。如果您打开隔离功能，则可以将双端口磁盘配置为使用 SCSI-2 或 SCSI-3（默认情况下为 SCSI-2）的法定设备。如果打开隔离功能并且您的共享设备连接到两个以上的节点，则可以将您的共享磁盘配置为使用

SCSI-3 协议（用于两个以上节点的默认协议）的法定设备。您可以使用 SCSI 覆盖标志使 Oracle Solaris Cluster 软件对双端口共享磁盘使用 SCSI-3 协议。

如果您为共享磁盘关闭隔离功能，那么可以将该磁盘配置为使用软件法定协议的法定设备。无论该磁盘是支持 SCSI-2 协议还是支持 SCSI-3 协议，都是如此。软件法定是 Oracle 的一种协议，用来模拟某种形式的 SCSI 永久组保留 (Persistent Group Reservation, PGR)。



注意 - 如果您使用的是不支持 SCSI 的磁盘（例如，SATA），则应该关闭 SCSI 隔离功能。

对于法定设备，您可以使用包含用户数据或属于某个设备组的磁盘。通过观察 `cluster show` 命令的输出中共享磁盘的 `access-mode` 值，可以查看由具有该共享磁盘的法定子系统所使用的协议。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面创建法定服务器设备或共享磁盘法定设备。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

有关以下过程所用命令的信息，请参见 [clsetup\(1CL\)](#) 和 [clquorum\(1CL\)](#) 手册页。

▼ 如何添加共享磁盘法定设备

Oracle Solaris Cluster 软件支持使用共享磁盘（包括 SCSI 和 SATA）设备作为法定设备。SATA 设备不支持 SCSI 保留项，您必须禁用 SCSI 保留隔离标志并使用软件法定协议将这些磁盘配置为法定设备。

要完成此过程，请使用磁盘驱动器的设备 ID (device ID, DID) 来标识该驱动器（设备 ID 由节点共享）。使用 `cldevice show` 命令可查看 DID 名称列表。有关其他信息，请参阅 [cldevice\(1CL\)](#) 手册页。在添加新法定设备之前，请确保群集中所有节点均处于联机状态。

使用此过程可配置 SCSI 或 SATA 设备。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 启动 `clsetup` 实用程序。

```
# clsetup
```

此时将显示 `clsetup` 主菜单。

3. 键入 "Quorum" (法定) 选项的编号。
这时将显示 "Quorum" (法定) 菜单。
4. 键入与添加法定设备选项对应的编号，然后在 `clsetup` 实用程序提示您确认要添加的法定设备时键入 `yes`。
`clsetup` 实用程序将询问您要添加哪种类型的法定设备。
5. 键入与共享磁盘法定设备选项对应的编号。
`clsetup` 实用程序将询问您要使用哪个全局设备。
6. 键入您正在使用的全局设备。
`clsetup` 实用程序将提示您确认将新的法定设备添加到指定的全局设备中。
7. 键入 `yes` 继续执行添加新法定设备的操作。
如果成功添加了新的法定设备，`clsetup` 实用程序会为此显示一条相应的消息。
8. 检验是否已添加法定设备。

```
# clquorum list -v
```

▼ 如何添加 Oracle ZFS Storage Appliance NAS 法定设备

在添加新法定设备之前，请确保群集中所有节点均处于联机状态。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面添加 Oracle ZFS Storage Appliance NAS 设备。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 有关设置 iSCSI 设备的说明，请参考随 Oracle ZFS Storage Appliance 提供的安装文档或者该设备的联机帮助。
2. 在每个群集节点上，搜索 iSCSI LUN 并将 iSCSI 访问列表设置为静态配置。

```
# iscsiadm modify discovery -s enable
```

```
# iscsiadm list discovery
Discovery:
Static: enabled
Send Targets: disabled
iSNS: disabled
```

```
# iscsiadm add static-config iqn.LUN-name,IP-address-of-NASdevice
# devfsadm -i iscsi
# cldevice refresh
```

3. 从一个群集节点上，为 iSCSI LUN 配置 DID。

```
# cldevice populate
```

4. 标识 DID 设备，该设备表示使用 iSCSI 刚刚配置到群集中的 NAS 设备 LUN。
使用 `cldevice show` 命令可查看 DID 名称列表。有关其他信息，请参阅 [cldevice\(1CL\)](#) 手册页。
5. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
6. 通过 `clquorum` 命令，使用[步骤 4](#)中指定的 DID 设备将 NAS 设备添加为法定设备。

```
# clquorum add d20
```

群集具有默认的规则来决定是使用 scsi-2、scsi-3 还是软件法定协议。有关更多信息，请参见 [clquorum\(1CL\)](#) 手册页。

▼ 如何添加法定服务器法定设备

开始之前 在添加 Oracle Solaris Cluster 法定服务器作为法定设备之前，必须先在主机上安装 Oracle Solaris Cluster 法定服务器软件，并启动和运行该法定服务器。有关安装法定服务器的信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[如何安装和配置 Oracle Solaris Cluster 法定服务器软件](#)”。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面创建法定服务器设备。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 确保所有 Oracle Solaris Cluster 节点都处于联机状态，并可与 Oracle Solaris Cluster 法定服务器进行通信。
 - a. 确保与群集节点直接相连的网络交换机满足以下条件之一：
 - 交换机支持快速生成树协议 (Rapid Spanning Tree Protocol, RSTP)。

- 交换机上已启用快速端口 (fast port) 模式。

必须具有上述某一项特性以确保群集节点与法定服务器之间的即时通信。如果通信因交换机而出现明显延迟，则群集会认为是缺少法定设备导致了通信不畅。

- b. 如果公共网络使用长度可变的子网，也称为无类别域间路由 (Classless Inter-Domain Routing, CIDR)，请在每个节点上修改以下文件。

如果您使用的是 RFC 791 中所定义的有类别子网，则无需执行这些步骤。

- i. 在 `/etc/inet/netmasks` 文件中，为群集所使用的每个公共子网添加一个相应的条目。

以下是一个包含了某个公共网络 IP 地址和网络掩码的条目示例：

```
10.11.30.0 255.255.255.0
```

- ii. 将 `netmask + broadcast +` 添加到每个 `/etc/hostname.adapter` 文件中主机名条目的后面。

```
nodename netmask + broadcast +
```

- c. 在群集的每个节点上，将法定服务器主机名添加到 `/etc/inet/hosts` 文件或 `/etc/inet/ipnodes` 文件中。

按如下所示，在文件中添加主机名到地址的映射。

```
ipaddress qshost1
```

`ipaddress` 正在运行法定服务器的计算机的 IP 地址。

`qshost1` 正在运行法定服务器的计算机的主机名。

- d. 如果使用了命名服务，请将法定服务器主机的名称到地址映射添加到名称服务数据库。

3. 启动 `clsetup` 实用程序。

```
# clsetup
```

此时将显示 `clsetup` 主菜单。

4. 键入 "Quorum" (法定) 选项的编号。
这时将显示 "Quorum" (法定) 菜单。
5. 键入与添加法定设备选项对应的编号。

然后键入 `yes` 确认添加法定设备。

`clsetup` 实用程序将询问您要添加哪种类型的法定设备。

6. 键入法定服务器法定设备对应的选项编号，然后键入 **yes** 确认您要添加法定服务器法定设备。
clsetup 实用程序将要求您提供新法定设备的名称。
7. 键入正在添加的法定设备的名称。
法定设备的名称可以是任一名称。该名称仅用于继续执行后续的管理命令。
clsetup 实用程序要求您提供法定服务器主机的名称。
8. 键入法定服务器所在主机的名称。
此名称指定了运行法定服务器的计算机的 IP 地址，或该计算机在网络中的主机名。
根据主机的 IPv4 或 IPv6 配置情况，必须在 `/etc/hosts` 文件或 `/etc/inet/ipnodes` 文件（或二者）中指定该计算机的 IP 地址。

注 - 指定的计算机必须能被所有群集节点访问，并且必须运行法定服务器。

clsetup 实用程序将提示您提供法定服务器的端口号。

9. 键入法定服务器用来与群集节点通信的端口号。
clsetup 实用程序将提示您确认添加新法定设备。
10. 键入 **yes** 继续执行添加新法定设备的操作。
如果成功添加了新的法定设备，clsetup 实用程序会为此显示一条相应的消息。
11. 检验是否已添加法定设备。

```
# clquorum list -v
```

删除或替换法定设备

本节提供了以下过程以删除或替换法定设备：

- [如何删除法定设备 \[160\]](#)
- [如何从群集中删除最后一个法定设备 \[161\]](#)
- [如何替换法定设备 \[163\]](#)

▼ 如何删除法定设备

删除法定设备后，它将不再参与建立法定的投票。请注意，所有由两个节点组成的群集均要求至少配置一个法定设备。如果这是群集的最后一个法定设备，`clquorum(1CL)` 将无法从配置中删除该设备。如果要删除某个节点，请删除连接到该节点的所有法定设备。

注 - 如果要删除的设备是群集中的最后一个法定设备，请参见[如何从群集中删除最后一个法定设备 \[161\]](#)中的过程。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面删除法定设备。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。

2. 确定要删除的法定设备。

```
# clquorum list -v
```

3. 执行 `clsetup` 实用程序。

```
# clsetup
```

这时将显示主菜单。

4. 键入 "Quorum" (法定) 选项的编号。

5. 键入与删除法定设备选项对应的编号。

回答删除过程中显示的问题。

6. 退出 `clsetup`。

7. 检验是否已删除法定设备。

```
# clquorum list -v
```

故障排除 如果在删除法定服务器法定设备时，群集与法定服务器主机之间的通信中断，则必须清除有关法定服务器主机的过时配置信息。有关执行此清除过程的说明，请参见[“清除过时的法定服务器群集信息” \[172\]](#)。

▼ 如何从群集中删除最后一个法定设备

此过程通过使用 `clquorum force` 选项 `-F` 从一个双节点群集删除最后一个法定设备。通常，应先删除故障设备，再添加替换法定设备。如果这不是双节点群集中的最后一个法定设备，请执行[如何删除法定设备 \[160\]](#)中的步骤。

添加法定设备涉及到节点重新配置，而这会涉及有故障的法定设备并会导致计算机出现紧急状态。使用 Force 选项可以删除有故障的法定设备，并且不会导致计算机出现紧急状态。使用 `clquorum` 命令可以从配置中删除设备。有关更多信息，请参见 [clquorum\(1CL\)](#) 手册页。删除故障的法定设备后，可使用 `clquorum add` 命令添加新设备。请参见“添加法定设备” [155]。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集中的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 使用 `clquorum` 命令删除法定设备。
如果法定设备发生故障，请使用 `-F Force` 选项删除故障设备。

```
# clquorum remove -F qd1
```

注 - 您也可以将要删除的节点设为维护状态，然后使用 `clquorum remove quorum` 命令删除法定设备。当群集处于安装模式时，`clsetup` 群集管理菜单选项不可用。有关更多信息，请参见[如何使节点进入维护状态](#) [216]以及 [clsetup\(1CL\)](#) 手册页。

3. 检验是否已删除法定设备。
4. 根据您删除最后一个法定设备的原因，继续执行以下步骤之一：

- 如果要替换已删除的法定设备，请完成以下子步骤：

- a. 添加新的法定设备。
有关添加新法定设备的说明，请参见“添加法定设备” [155]。
- b. 使群集脱离安装模式。

```
# cluster set -p installmode=disabled
```

- 如果要将群集缩减为单节点群集，请通过安装模式删除群集。

```
# cluster set -p installmode=disabled
```

例 58 删除最后一个法定设备

本示例阐述如何将群集设为维护模式并删除群集配置中最后一个法定设备。

将群集置于安装模式

```
# cluster set -p installmode=enabled
```

删除法定设备

```
# clquorum remove d3
```

检验是否已删除法定设备

```
# clquorum list -v
```

Quorum	Type
-----	----
scphyshost-1	node
scphyshost-2	node
scphyshost-3	node

▼ 如何替换法定设备

使用该过程用另一个法定设备替换现有的法定设备。您可以用类型相似的设备替换法定设备，例如可以用另一个 NAS 设备替换现有的 NAS 设备，还可以用不同类型的设备替换法定设备，例如用一个共享的磁盘替换 NAS 设备。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 配置新法定设备。
需要首先将一个法定设备添加到配置中，来替换旧的设备。有关将新的法定设备添加到群集中的信息，请参见[“添加法定设备” \[155\]](#)。
2. 删除要替换的法定设备。
有关从配置中删除旧的法定设备的信息，请参见[如何删除法定设备 \[160\]](#)。
3. 如果法定设备是故障磁盘，请替换该磁盘。
请参见磁盘盒硬件手册中的硬件过程。另请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》。

维护法定设备

本节提供了以下过程以维护法定设备：

- [如何修改法定设备节点列表 \[164\]](#)
- [如何将法定设备置于维护状态 \[164\]](#)
- [如何使法定设备脱离维护状态 \[166\]](#)
- [如何列出法定配置 \[167\]](#)
- [如何修复法定设备 \[168\]](#)

- “更改法定设备的默认超时时间” [169]

▼ 如何修改法定设备节点列表

您可以使用 `clsetup` 实用程序向现有法定设备的节点列表中添加节点或从中删除节点。要修改法定设备的节点列表，必须删除该法定设备，修改节点与删除的法定设备的物理连接，然后将该法定设备重新添加到群集配置中。添加法定设备时，`clquorum` 命令会自动为连接到磁盘的所有节点配置节点到磁盘路径。有关更多信息，请参见 [clquorum\(1CL\)](#) 手册页。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 确定要修改的法定设备的名称。

```
# clquorum list -v
```

3. 启动 `clsetup` 实用程序。

```
# clsetup
```

这时将显示主菜单。

4. 键入 "Quorum" (法定) 选项的编号。
这时将显示 "Quorum" (法定) 菜单。
5. 键入与删除法定设备选项对应的编号。
请按照说明进行操作。系统将询问您要删除的磁盘的名称。
6. 添加或删除节点与法定设备之间的连接。
7. 键入与添加法定设备选项对应的编号。
请按照说明进行操作。系统将询问您要作为法定设备使用的磁盘的名称。
8. 检验是否已添加法定设备。

```
# clquorum list -v
```

▼ 如何将法定设备置于维护状态

使用 `clquorum` 命令可将法定设备置于维护状态。有关更多信息，请参见 [clquorum\(1CL\)](#) 手册页。`clsetup` 实用程序目前没有此功能。

如果在较长的一段时间内不使用法定设备，请将其置于维护状态。这样，法定设备的法定选票计数设置为零，当设备正在维修时，将不会参与投票。在维护状态期间，法定设备的配置信息将被保留下来。

注 - 所有双节点群集均要求至少配置一个法定设备。如果这是双节点群集的最后一个法定设备，则 `clquorum` 无法将该设备置于维护状态。

要将群集节点置于维护状态，请参见[如何使节点进入维护状态 \[216\]](#)。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面禁用法定设备以将其置于维护状态。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

如果群集处于安装模式，则单击 "Reset Quorum Device"（复位法定设备）来退出安装模式。

- 1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
- 2. 将法定设备置于维护状态。

```
# clquorum disable device  
  
device           指定要更改的磁盘设备的 DID 名称，例如 d4。
```

- 3. 检验该法定设备当前是否处于维护状态。
处于维护状态的设备的输出应该为法定设备选票读取零。

```
# clquorum status device
```

例 59 将法定设备置于维护状态

以下示例显示了如何将法定设备置于维护状态，以及如何检验操作结果。

```
# clquorum disable d20  
# clquorum status d20  
  
=== Cluster Quorum ===  
  
--- Quorum Votes by Device ---  
  
Device Name      Present      Possible      Status  
-----
```

d20 1 1 Offline

另请参见 要重新启用法定设备，请参见[如何使法定设备脱离维护状态 \[166\]](#)。

要将某个节点置于维护状态，请参见[如何使节点进入维护状态 \[216\]](#)。

▼ 如何使法定设备脱离维护状态

如果某个法定设备处于维护状态，而您希望使该法定设备脱离维护状态并将法定选票数重置为默认值，请运行此过程。



注意 - 如果您既未指定 `globaldev` 选项，也未指定 `node` 选项，则会重置整个群集的法定计数。

配置法定设备时，Oracle Solaris Cluster 软件将 $N-1$ 作为选票数分配给法定设备，其中 N 是连接到法定设备的选票数。例如，连接到两个选票数非零的节点的法定设备的法定选票数为一（二减一）。

- 要使群集节点及其相关法定设备脱离维护状态，请参见[如何使节点脱离维护状态 \[218\]](#)。
- 要了解有关法定投票计数的更多信息，请参见《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》中的“[About Quorum Vote Counts](#)”。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面启用法定设备以使其脱离维护状态。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 重置法定计数。

```
# clquorum enable device
```

`device` 指定要重置的法定设备的 DID 名称，例如 `d4`。

3. 如果由于某个节点已处于维护状态而需要重置其法定计数，请重新引导该节点。
4. 检验法定选票数。

```
# clquorum show +
```

例 60 重新设置法定选票计数（法定设备）

以下示例将一个法定设备的法定计数重置为默认值并检验操作结果。

```
# clquorum enable d20
# clquorum show +

=== Cluster Nodes ===

Node Name:                phys-schost-2
Node ID:                   1
Quorum Vote Count:         1
Reservation Key:           0x43BAC41300000001

Node Name:                phys-schost-3
Node ID:                   2
Quorum Vote Count:         1
Reservation Key:           0x43BAC41300000002

=== Quorum Devices ===

Quorum Device Name:        d3
Enabled:                   yes
Votes:                     1
Global Name:               /dev/did/rdisk/d20s2
Type:                      shared_disk
Access Mode:               scsi3
Hosts (enabled):           phys-schost-2, phys-schost-3
```

▼ 如何列出法定配置

您无需成为 root 角色即可列出法定配置。您可以承担任何可提供 `solaris.cluster.read` RBAC 授权的角色。

注 - 在增加或减少连接到法定设备的节点数时，系统不会自动重新计算法定选票计数。如果删除了所有法定设备，然后将它们重新添加到配置中，则您可以重新建立正确的法定选票。对于双节点群集，请临时添加一个新的法定设备，然后删除原法定设备并将其添加回配置。然后，删除临时法定设备。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面查看法定配置。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

- 使用 `clquorum` 命令可列出法定配置。

```
% clquorum show +
```

例 61 列出法定配置

```
% clquorum show +
```

```
=== Cluster Nodes ===
```

```
Node Name:                phys-schost-2
Node ID:                   1
Quorum Vote Count:        1
Reservation Key:           0x43BAC41300000001
```

```
Node Name:                phys-schost-3
Node ID:                   2
Quorum Vote Count:        1
Reservation Key:           0x43BAC41300000002
```

```
=== Quorum Devices ===
```

```
Quorum Device Name:       d3
Enabled:                  yes
Votes:                    1
Global Name:              /dev/did/rdisk/d20s2
Type:                     shared_disk
Access Mode:              scsi3
Hosts (enabled):          phys-schost-2, phys-schost-3
```

▼ 如何修复法定设备

使用此过程可替换发生故障的法定设备。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 删除要替换的作为法定设备的磁盘设备。

注 - 如果要删除的设备是最后一个法定设备，则可能需要先添加另一个磁盘作为新的法定设备。此步骤可确保一旦在替换过程中出现故障，群集中仍存在有效的法定设备。有关添加新的法定设备的信息，请参见[“添加法定设备” \[155\]](#)。

要删除作为法定设备的磁盘设备，请参见[如何删除法定设备 \[160\]](#)。

2. 更换磁盘设备。

要更换磁盘设备，请参见硬件指南中针对磁盘盒的过程。另请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》。

3. 将更换的磁盘添加为新的法定设备。
要添加磁盘作为新的法定设备，请参见[“添加法定设备” \[155\]](#)。

注 - 如果已在[步骤 1](#)中添加了其他法定设备，现在就可以放心地删除它了。要删除法定设备，请参见[如何删除法定设备 \[160\]](#)。

更改法定设备的默认超时时间

在群集重新配置期间，默认有 25 秒的超时时间来完成法定操作。您可以按照《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的[“如何配置法定设备”](#)中的说明，将法定超时时间增大到较高的值。除了增大超时值，还可以切换到其他法定设备。

《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的[“如何配置法定设备”](#)中提供了其他故障排除信息。

注 - 对于 Oracle Real Application Clusters (Oracle RAC)，请勿更改默认的 25 秒法定超时时间。在某些记忆分裂方案中，较长的超时周期可能会导致 Oracle RAC VIP 故障转移因 VIP 资源超时而失败。

如果所用法定设备不适合使用默认的 25 秒超时，请使用其他法定设备。

管理 Oracle Solaris Cluster 法定服务器

Oracle Solaris Cluster 法定服务器提供一个法定设备（非共享存储设备）。本节介绍有关管理 Oracle Solaris Cluster 法定服务器的过程，其中包括：

- [“启动和停止法定服务器软件” \[170\]](#)
- [如何启动法定服务器 \[170\]](#)
- [如何停止法定服务器 \[170\]](#)
- [“显示有关法定服务器的信息” \[171\]](#)
- [“清除过时的法定服务器群集信息” \[172\]](#)

有关安装和配置 Oracle Solaris Cluster 法定服务器的信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的[“如何安装和配置 Oracle Solaris Cluster 法定服务器软件”](#)。

启动和停止法定服务器软件

这些过程介绍如何启动和停止 Oracle Solaris Cluster 软件。

默认情况下，这些过程会启动和停止单个默认法定服务器，除非您对法定服务器配置文件 `/etc/scqsd/scqsd.conf` 的内容进行了定制。默认法定服务器绑定在端口 9000 上，并使用 `/var/scqsd` 目录存储法定信息。

有关安装 Quorum Server 软件的信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[如何安装和配置 Oracle Solaris Cluster 法定服务器软件](#)”。有关更改法定超时时间值的信息，请参见“[更改法定设备的默认超时时间](#)”[169]。

▼ 如何启动法定服务器

1. 在要启动 Oracle Solaris Cluster 软件的主机上承担 `root` 角色。
2. 使用 `clquorumserver start` 命令启动该软件。

```
# clquorumserver start quorumserver
```

`quorumserver` 标识法定服务器。可以使用法定服务器所侦听的端口号。如果在配置文件中提供了实例名称，则也可以使用该名称。

要启动单个法定服务器，请提供实例名称或端口号。要启动所有法定服务器（配置了多个法定服务器时），请使用 `+` 操作数。

例 62 启动所有已配置的法定服务器

以下示例将启动所有已配置的法定服务器。

```
# clquorumserver start +
```

例 63 启动特定法定服务器

以下示例将启动侦听端口号 2000 的法定服务器。

```
# clquorumserver start 2000
```

▼ 如何停止法定服务器

1. 在要启动 Oracle Solaris Cluster 软件的主机上承担 `root` 角色。

2. 使用 `clquorumserver stop` 命令停止该软件。

```
# clquorumserver stop [-d] quorumserver
```

`-d` 控制法定服务器是否在您下一次引导计算机时启动。如果您指定了 `-d` 选项，则在计算机下一次引导时，法定服务器将不会启动。

`quorumserver` 标识法定服务器。可以使用法定服务器所侦听的端口号。如果在配置文件中提供了实例名称，则也可以使用该名称。
要停止单个法定服务器，请提供实例名称或端口号。要停止所有法定服务器（配置了多个法定服务器时），请使用 `+` 操作数。

例 64 停止所有已配置的法定服务器

以下示例将停止所有已配置的法定服务器。

```
# clquorumserver stop +
```

例 65 停止特定法定服务器

以下示例将停止侦听端口号 2000 的法定服务器。

```
# clquorumserver stop 2000
```

显示有关法定服务器的信息

可以显示有关法定服务器的配置信息。对于每个将法定服务器配置为法定设备的群集，该命令可以显示相应的群集名称、群集 ID、保留关键字列表和注册关键字列表。

▼ 如何显示有关法定服务器的信息

1. 在要显示法定服务器信息的主机上，承担 `root` 角色。

非 `root` 角色用户需要具有 `solaris.cluster.read` 基于角色的访问控制 (role-based access control, RBAC) 授权。有关 RBAC 权限配置文件的更多信息，请参见 [rbac\(5\)](#) 手册页。

2. 使用 `clquorumserver` 命令显示法定服务器的配置信息。

```
# clquorumserver show quorumserver
```

`quorumserver` 标识一个或多个法定服务器。可以使用实例名称或端口号指定法定服务器。要显示所有法定服务器的配置信息，请使用 `+` 操作数。

例 66 显示一个法定服务器的配置信息

下面的示例显示了使用端口 9000 的法定服务器的配置信息。该命令显示每个将法定服务器配置为法定设备的群集的信息。这些信息包括群集的名称和 ID 以及设备上保留项和注册项的列表。

在以下示例中，ID 为 1、2、3 和 4 的群集 bastille 节点在法定服务器上注册了自己的项。此外，由于节点 4 拥有法定设备保留关键字，因此其关键字显示在保留关键字列表中。

```
# clquorumserver show 9000

=== Quorum Server on port 9000 ===

--- Cluster bastille (id 0x439A2EFB) Reservation ---

Node ID:                4
Reservation key:         0x439a2efb00000004

--- Cluster bastille (id 0x439A2EFB) Registrations ---

Node ID:                1
Registration key:        0x439a2efb00000001

Node ID:                2
Registration key:        0x439a2efb00000002

Node ID:                3
Registration key:        0x439a2efb00000003

Node ID:                4
Registration key:        0x439a2efb00000004
```

例 67 显示多个法定服务器的配置信息

以下示例显示三个法定服务器 qs1、qs2 和 qs3 的配置信息。

```
# clquorumserver show qs1 qs2 qs3
```

例 68 显示所有正在运行的法定服务器的配置信息

以下示例显示所有正在运行的法定服务器的配置信息。

```
# clquorumserver show +
```

清除过时的法定服务器群集信息

要删除类型为 quorumserver 的法定设备，请使用 clquorum remove 命令（如[如何删除法定设备 \[160\]](#)所述）。在常规操作情况下，该命令也将删除有关法定服务器主机的法定

服务器信息。不过，如果群集与法定服务器主机之间的通信中断，则删除法定设备不会清除该信息。

在以下情况下，法定服务器群集信息将变为无效：

- 在未首先使用 `clquorum remove` 命令删除群集法定设备的情况下取消了对群集的授权
- 在法定服务器主机处于关闭状态时从群集中删除了 `quorum_server` 类型的法定设备



注意 - 如果尚未从群集中删除 `quorumserver` 类型的法定设备，则按照以下过程清除有效的法定服务器会影响群集法定。

▼ 如何清除法定服务器配置信息

开始之前 从群集中删除法定服务器法定设备，如[如何删除法定设备 \[160\]](#)所述。



注意 - 如果群集仍在使用该法定服务器，则执行该过程会影响群集法定。

1. 在法定服务器主机上承担 `root` 角色。
2. 使用 `clquorumserver clear` 命令清除配置文件。

```
# clquorumserver clear -c clustername -I clusterID quorumserver [-y]
```

<code>-c <i>clustername</i></code>	先前将法定服务器用作法定设备的群集的名称。 可以通过在群集节点上运行 <code>cluster show</code> 来获取群集名称。
<code>-I <i>clusterID</i></code>	群集 ID。 群集 ID 是一个 8 位十六进制数字。可以通过在群集节点上运行 <code>cluster show</code> 来获取群集 ID。
<code><i>quorumserver</i></code>	一个或多个法定服务器的标识符。 可以使用端口号或实例名称来标识法定服务器。端口号供群集节点用于与法定服务器进行通信。实例名称是在法定服务器配置文件 <code>/etc/scqsd/scqsd.conf</code> 中指定的。
<code>-y</code>	强制 <code>clquorumserver clear</code> 命令从配置文件中清除群集信息，而不先提示进行确认。 仅当确信要从法定服务器中删除过时的群集信息时，才使用该选项。

3. （可选）如果该服务器实例上未配置其他法定设备，请停止该法定服务器。

例 69 从法定服务器配置中清除过时的群集信息

本示例将从使用端口 9000 的法定服务器中删除有关名为 `sc-cluster` 的群集的信息。

```
# clquorumserver clear -c sc-cluster -I 0x4308D2CF 9000
```

```
The quorum server to be unconfigured must have been removed from the cluster.
```

```
Unconfiguring a valid quorum server could compromise the cluster quorum. Do you  
want to continue? (yes or no) y
```

管理群集互连和公共网络

本章提供管理 Oracle Solaris Cluster 互连和公共网络的软件过程。

群集互连和公共网络的管理由硬件和软件过程组成。通常，在初次安装和配置群集时，需要配置群集互连和公共网络，包括所有公共网络管理 (public network management, PNM) 对象。

PNM 对象包括 Internet 协议网络多路径 (Internet Protocol network multipathing, IPMP) 组，中继和数据链路多路径 (datalink multipathing, DLMP) 链路聚合以及由链路聚合直接支持的 VNIC。多路径功能随 Oracle Solaris 11 OS 自动安装，必须启用多路径功能才能使用该功能。如果以后需要改变群集互连网络配置，您可以使用本章中的软件过程。

有关在群集中配置 IP 网络多路径组的信息，请参见[“管理公共网络” \[188\]](#)一节。有关 IPMP 的信息，请参见《[在 Oracle Solaris 11.3 中管理 TCP/IP 网络、IPMP 和 IP 隧道](#)》中的第 3 章,“[管理 IPMP](#)”。有关链路聚合的信息，请参见《[在 Oracle Solaris 11.3 中管理网络数据链路](#)》中的第 2 章,“[使用链路聚合配置高可用性](#)”。

本章提供了有关以下主题的信息和过程。

- [“管理群集互连” \[175\]](#)
- [“管理公共网络” \[188\]](#)

有关本章中相关过程的概括性说明，请参见表 12 [“任务列表：管理群集互连”](#)。

有关群集互连和公共网络的背景信息和概述信息，请参阅《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》。

管理群集互连

本节提供了重新配置群集互连（例如群集传输适配器和群集传输电缆）的过程。这些过程要求安装 Oracle Solaris Cluster 软件。

在大多数情况下，可以使用 `clsetup` 实用程序来管理群集互连的群集传输。有关更多信息，请参见 [clsetup\(1CL\)](#) 手册页。所有群集互连命令都必须从全局群集节点中运行。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面执行这些任务中的一些任务。有关登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

有关群集软件安装过程，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》。有关维修群集硬件组件的过程，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》。

注 - 在群集互连过程中，只要适合，一般情况下都可以使用默认端口名。默认端口名称与用于控制电缆适配器端的那个节点的内部节点 ID 号相同。

表 12 任务列表：管理群集互连

任务	说明
使用 clsetup(1CL) 管理群集传输	如何访问群集配置实用程序 [31]
使用 <code>clinterconnect status</code> 检查群集互连的状态	如何检查群集互连的状态 [177]
使用 <code>clsetup</code> 添加群集传输电缆、传输适配器或交换机	如何添加群集传输电缆、传输适配器或传输交换机 [178]
使用 <code>clsetup</code> 删除群集传输电缆、传输适配器或传输交换机	如何删除传输电缆、传输适配器或传输交换机 [180]
使用 <code>clsetup</code> 启用群集传输电缆	如何启用群集传输电缆 [182]
使用 <code>clsetup</code> 禁用群集传输电缆	如何禁用群集传输电缆 [183]
确定传输适配器的实例编号	如何确定传输适配器的实例编号 [184]
更改现有群集的 IP 地址或地址范围	如何更改现有群集的专用网络地址或地址范围 [185]

动态重新配置群集互连

在对群集互连完成动态重新配置 (dynamic reconfiguration, DR) 操作时，必须考虑几个问题。

- 针对 Oracle Solaris 动态重新配置功能介绍的所有要求、过程和限制也适用于 Oracle Solaris Cluster 动态重新配置支持（操作系统停止操作除外）。因此，请先阅读关于 Oracle Solaris 动态重新配置功能的文档，然后再对 Oracle Solaris Cluster 软件使用动态重新配置功能。您应该特别注意那些在执行动态重新配置分离操作时影响非网络 IO 设备的问题。
- Oracle Solaris Cluster 软件拒绝对活动的专用互连接口进行动态重新配置删除板操作。
- 要对活动的群集互连执行动态重新配置，必须从群集中完全删除活动适配器。使用 `clsetup` 菜单或相应命令。



注意 - Oracle Solaris Cluster 软件要求每个群集节点与群集中其他节点之间至少有一个有效路径。如果某个专用互连接口支持到任何群集节点的最后一​​条路径，则请勿禁用它。

对公共网络接口执行动态重新配置操作时，请按所示顺序完成下列过程。

表 13 任务列表：动态重新配置公共网络接口

任务	说明
1. 从活动的互连中禁用并删除接口	“动态重新配置公共网络接口” [189]
2. 对公共网络接口执行动态重新配置操作	

▼ 如何检查群集互连的状态

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

执行此步骤不需要作为 root 角色登录。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面检查群集互连的状态。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 检查群集互连的状态。

```
% clinterconnect status
```

2. 有关常见状态消息，请参阅下表。

状态消息	说明和可能的操作
Path online	此路径当前工作正常。不需要执行任何操作。
Path waiting	当前正在初始化此路径。不需要执行任何操作。
Faulted	此路径当前不工作。如果路径在等待和联机状态之间，则这种情况是瞬态的。如果重新运行 clinterconnect status 后仍出现此消息，请采取更正措施。

例 70 检查群集互连的状态

以下示例说明了当前运行的群集互连的状态。

```
% clinterconnect status
-- Cluster Transport Paths --
      Endpoint                Endpoint                Status
      -----                -
Transport path: phys-schost-1:net0 phys-schost-2:net0 Path online
Transport path: phys-schost-1:net4 phys-schost-2:net4 Path online
Transport path: phys-schost-1:net0 phys-schost-3:net0 Path online
```

```
Transport path:  phys-schost-1:net4  phys-schost-3:net4  Path online
Transport path:  phys-schost-2:net0  phys-schost-3:net0  Path online
Transport path:  phys-schost-2:net4  phys-schost-3:net4  Path online
```

▼ 如何添加群集传输电缆、传输适配器或传输交换机

有关群集专用传输的要求的信息，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》中的“Interconnect Requirements and Restrictions”。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面向群集添加电缆、传输适配器和专用适配器。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 确保安装了物理群集传输电缆。
有关安装群集传输电缆的过程，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》。
2. 在群集中的任一节点上承担 root 角色。
3. 启动 clsetup 实用程序。

```
# clsetup
```

这时将显示主菜单。

4. 键入与显示群集互连菜单选项对应的编号。
5. 键入与添加传输电缆选项对应的编号。
按说明进行操作，并键入请求的信息。
6. 键入与向节点添加传输适配器选项对应的编号。
按说明进行操作，并键入请求的信息。

如果打算将以下任何适配器用于群集互联，请在各群集节点上的 /etc/system 文件中添加相关条目。此条目在下次引导系统后生效。

适配器	条目
nge	set nge:nge_taskq_disable=1

适配器	条目
e1000g	set e1000g:e1000g_taskq_disable=1

7. 键入与添加传输交换机选项对应的编号。
按说明进行操作，并键入请求的信息。
8. 检验是否添加了群集传输电缆、传输适配器或传输交换机。

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

例 71 检验添加群集传输电缆、传输适配器或传输交换机

以下示例说明了如何检验向节点添加传输电缆、传输适配器或传输交换机。该示例包含了数据链路提供者接口 (Data Link Provider Interface, DLPI) 传输类型的设置。

```
# clinterconnect show phys-schost-1:net5,hub2
===Transport Cables===
Transport Cable:          phys-schost-1:net5@0,hub2
Endpoint1:                phys-schost-2:net4@0
Endpoint2:                hub2@2
State:                    Enabled

# clinterconnect show phys-schost-1:net5
=== Transport Adepters for net5
Transport Adapter:                net5
Adapter State:                    Enabled
Adapter Transport Type:           dlpi
Adapter Property (device_name):   net6
Adapter Property (device_instance): 0
Adapter Property (lazy_free):     1
Adapter Property (dlpi_heartbeat_timeout): 10000
Adpater Property (dlpi_heartbeat_quantum): 1000
Adapter Property (nw_bandwidth):  80
Adapter Property (bandwidth):     70
Adapter Property (ip_address):    172.16.0.129
Adapter Property (netmask):       255.255.255.128
Adapter Port Names:               0
Adapter Port State (0):           Enabled

# clinterconnect show phys-schost-1:hub2

=== Transport Switches ===
Transport Switch:                hub2
Switch State:                    Enabled
Switch Type:                     switch
Switch Port Names:               1 2
Switch Port State(1):            Enabled
Switch Port State(2):            Enabled
```

接下来的步骤 要检查群集传输电缆的互连状态，请参见[如何检查群集互连的状态 \[177\]](#)。

▼ 如何删除传输电缆、传输适配器或传输交换机

可使用以下过程从节点配置中删除群集传输电缆、传输适配器以及传输交换机。禁用电缆后，电缆的两个端点仍处于已配置状态。如果适配器仍用作传输电缆的一个端点，则无法删除该适配器。



注意 - 每个群集节点至少需要一条通向群集中其他各节点的有效传输路径。任何两个节点之间都必须有传输路径。禁用电缆前，请务必检验节点的群集互连的状态。只有当您确认了某个电缆连接是冗余的之后，才能禁用它。也就是说，要确保有另外一个连接可用。禁用节点所剩的最后个工作电缆会使该节点脱离群集。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面从群集中删除电缆、传输适配器和专用适配器。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担 root 角色。
2. 检查其余群集传输路径的状态。

```
# clinterconnect status
```



注意 - 如果在尝试删除由两个节点所组成的群集中的一个节点时收到错误消息（如 path faulted），请先找出问题之所在，然后再继续执行此过程。这样的问题可能表明节点路径不可用。删除所剩的正常路径会使节点脱离群集，并可能导致群集重新配置。

3. 启动 clsetup 实用程序。

```
# clsetup
```

这时将显示主菜单。
4. 键入与访问群集互连菜单选项对应的编号。
5. 键入与禁用传输电缆选项对应的编号。
按说明进行操作，并键入请求的信息。您需要知道适用的节点名称、适配器名称和交换机名称。
6. 键入与删除传输电缆选项对应的编号。
按说明进行操作，并键入请求的信息。您需要知道适用的节点名称、适配器名称和交换机名称。

注 - 如果删除的是物理电缆，请断开端口与目标设备之间的电缆。

7. 键入与从节点删除传输适配器选项对应的编号。

按说明进行操作，并键入请求的信息。您需要知道适用的节点名称、适配器名称和交换机名称。

如果要从节点中删除物理适配器，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》以了解硬件维修过程。

8. 键入与删除传输交换机选项对应的编号。

按说明进行操作，并键入请求的信息。您需要知道适用的节点名称、适配器名称和交换机名称。

注 - 如果有任何端口在任何传输电缆上仍用作端点，则无法删除交换机。

9. 检验是否已删除电缆、适配器或交换机。

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

此命令的输出内容中不应出现已从相应节点删除的传输电缆或适配器。

例 72 检验删除传输电缆、传输适配器或传输交换机

以下示例说明了如何检验传输电缆、传输适配器或传输交换机的删除。

```
# clinterconnect show phys-schost-1:net5,hub2@0
===Transport Cables===
Transport Cable:                phys-schost-1:net5,hub2@0
Endpoint1:                      phys-schost-1:net5
Endpoint2:                      hub2@0
State:                          Enabled

# clinterconnect show phys-schost-1:net5
=== Transport Adepters for net5
Transport Adapter:              net5
Adapter State:                  Enabled
Adapter Transport Type:         dlpi
Adapter Property (device_name): net6
Adapter Property (device_instance): 0
Adapter Property (lazy_free):   1
Adapter Property (dlpi_heartbeat_timeout): 10000
Adapter Property (dlpi_heartbeat_quantum): 1000
Adapter Property (nw_bandwidth): 80
Adapter Property (bandwidth):   70
Adapter Property (ip_address):  172.16.0.129
Adapter Property (netmask):     255.255.255.128
```

```

Adapter Port Names:                                0
Adapter Port State (0):                            Enabled

# clinterconnect show hub2
=== Transport Switches ===
Transport Switch:                                hub2
State:                                            Enabled
Type:                                            switch
Port Names:                                    1 2
Port State(1):                                Enabled
Port State(2):                                Enabled

```

▼ 如何启用群集传输电缆

此选项用于启用现有群集传输电缆。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面启用电缆。单击 "Private Interconnects" (专用互连)，单击 "Cables" (电缆)，单击电缆编号以突出显示该电缆，然后单击 "Enable" (启用)。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担 root 角色。

2. 启动 clsetup 实用程序。

```
# clsetup
```

这时将显示主菜单。

3. 键入与访问群集互连菜单选项对应的编号。

4. 键入与启用传输电缆选项对应的编号。

出现提示后按说明操作。您需要提供正在尝试标识的电缆的一个端点的节点名称和适配器名称。

5. 检验是否已启用该电缆。

```
# clinterconnect show node:adapter,adapternode
```

输出内容将类似如下：

```
# clinterconnect show phys-schost-1:net5,hub2
```

```

Transport cable:  phys-schost-2:net0@0 ethernet-1@2    Enabled
Transport cable:  phys-schost-3:net5@1 ethernet-1@3    Enabled
Transport cable:  phys-schost-1:net5@0 ethernet-1@1    Enabled

```

▼ 如何禁用群集传输电缆

您可能需要禁用群集传输电缆以便临时关闭群集互连路径。这在排除群集互连故障或更换群集互连硬件时很有用。

禁用电缆后，电缆的两个端点仍处于已配置状态。如果适配器仍用作传输电缆的一个端点，则无法删除该适配器。



注意 - 每个群集节点至少需要一条通向群集中其他各节点的有效传输路径。任何两个节点之间都必须有传输路径。禁用电缆前，请务必检验节点的群集互连的状态。只有当您确认了某个电缆连接是冗余的之后，才能禁用它。也就是说，要确保有另外一个连接可用。禁用节点所剩的最后个工作电缆会使该节点脱离群集。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面禁用电缆。单击 "Private Interconnects" (专用互连)，单击 "Cables" (电缆)，单击电缆编号以突出显示该电缆，然后单击 "Disable" (禁用)。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在群集中的任一节点上承担 **root** 角色。
2. 禁用电缆前，请检查群集互连的状态。

```
# clinterconnect status
```



注意 - 如果在尝试删除由两个节点所组成的群集中的一个节点时收到错误消息（如 "path faulted"），请先找出问题之所在，然后再继续执行此过程。这样的问题可能表明节点路径不可用。删除所剩的正常路径会使节点脱离群集，并可能导致群集重新配置。

3. 启动 `clsetup` 实用程序。

```
# clsetup
```

这时将显示主菜单。

4. 键入与访问群集互连菜单选项对应的编号。

- 键入与禁用传输电缆选项对应的编号。

按提示进行操作，并提供请求的信息。这时将禁用此群集互连中的所有组件。您需要提供正在尝试标识的电缆的一个端点的节点名称和适配器名称。

- 检验是否已禁用电缆。

```
# clinterconnect show node:adapter,adapternode
```

输出内容将类似如下：

```
# clinterconnect show -p phys-schost-1:net5,hub2
Transport cable:  phys-schost-2:net0@0 ethernet-1@2    Disabled
Transport cable:  phys-schost-3:net5@1 ethernet-1@3    Enabled
Transport cable:  phys-schost-1:net5@0 ethernet-1@1    Enabled
```

▼ 如何确定传输适配器的实例编号

您需要确定传输适配器的实例编号，以确保通过 `clsetup` 命令添加和删除正确的传输适配器。适配器的名称是适配器类型和适配器的实例编号的组合。

- 根据槽号，查找适配器的名称。

下面的屏幕只是一个示例，反映的可能不是您的硬件的真实情况。

```
# prtdiag
...
===== IO Cards =====
Bus  Max
IO  Port Bus      Freq Bus  Dev,
Type  ID  Side Slot MHz  Freq Func State Name Model
-----
XYZ   8   B    2    33   33  2,0  ok   xyz11c8,0-xyz11c8,d665.11c8.0.0
XYZ   8   B    3    33   33  3,0  ok   xyz11c8,0-xyz11c8,d665.11c8.0.0
...
```

- 使用适配器的路径来查找适配器的实例编号。

下面的屏幕只是一个示例，反映的可能不是您的硬件的真实情况。

```
# grep sci /etc/path_to_inst
"/xyz@1f,400/pci11c8,0@2" 0 "ttt"
"/xyz@1f,4000.pci11c8,0@4 "ttt"
```

- 根据适配器的名称和槽号，查找适配器的实例编号。

下面的屏幕只是一个示例，反映的可能不是您的硬件的真实情况。

```
# prtconf
...
xyz, instance #0
xyz11c8,0, instance #0
```



```
xyz11c8,0, instance #1
...
```

▼ 如何更改现有群集的专用网络地址或地址范围

使用此过程可更改专用网络地址或/和所使用的网络地址的范围。要使用命令行执行此任务，请参见 [cluster\(1CL\)](#) 手册页。

开始之前 请确保启用了 root 角色对所有群集节点的远程 shell ([rsh\(1M\)](#)) 或安全 shell ([ssh\(1\)](#)) 访问权限。

1. 在每个群集节点上执行以下子步骤，将所有群集节点重新引导至非群集模式：
 - a. 在将以非群集模式启动的群集节点上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
 - b. 使用 `clnode evacuate` 和 `cluster shutdown` 命令关闭节点。
`clnode evacuate` 命令可将指定节点上的所有设备组切换到下一个首选节点。该命令还将指定节点上的所有资源组切换到下一个首选的节点。

```
# clnode evacuate node
# cluster shutdown -g0 -y
```

2. 从一个节点启动 `clsetup` 实用程序。
在非群集模式下运行时，`clsetup` 实用程序会为非群集模式操作显示主菜单。
3. 选择 "Change Network Addressing and Ranges for the Cluster Transport"（更改群集传输的网络寻址和范围）菜单项。
`clsetup` 实用程序显示当前的专用网络配置，然后询问您是否要更改此配置。
4. 要更改专用网络 IP 地址或 IP 地址范围，请键入 `yes`，然后按回车键。
`clsetup` 实用程序将显示默认的专用网络 IP 地址 `172.16.0.0`，并询问您是否接受此默认值。
5. 更改或接受此专用网络 IP 地址。
 - 要接受默认的专用网络 IP 地址并继续进行 IP 地址范围更改，请键入 `yes`，然后按回车键。
 - 要更改默认的专用网络 IP 地址，请执行以下步骤：
 - a. 对于 `clsetup` 实用程序询问的是否接受默认地址的问题，键入 `no` 作为响应，然后按回车键。

`clsetup` 实用程序将提示您输入新的专用网络 IP 地址。

- b. 键入新的 IP 地址，然后按回车键。

`clsetup` 实用程序会显示默认网络掩码，然后询问您是否接受该默认网络掩码。

6. 更改或接受默认的专用网络 IP 地址。

默认网络掩码为 255.255.240.0。此默认 IP 地址范围支持在群集中包含最多 64 个节点、12 个区域群集和 10 个专用网络。

- 要接受该默认 IP 地址范围，请键入 **yes**，然后按回车键。

- 要更改 IP 地址范围，请执行以下步骤：

- a. 对于 `clsetup` 实用程序询问的是否接受默认地址范围的问题，键入 **no** 作为响应，然后按回车键。

当您拒绝默认网络掩码时，`clsetup` 实用程序将提示您输入要在群集中配置的节点、专用网络和区域群集的数量。

- b. 提供您期望在群集中配置的节点、专用网络和区域群集的数目。

`clsetup` 实用程序将根据这些数字计算出两个网络掩码供选择：

- 第一个网络掩码是支持指定节点、专用网络和区域群数目的最小网络掩码。
- 第二个网络掩码可支持两倍于指定值的节点、专用网络和区域群集数目，从而适应未来可能出现的增长情况。

- c. 指定上述任一网络掩码，或另外指定一个可支持预期节点、专用网络和区域群集数目的网络掩码。

7. 对于 `clsetup` 实用程序询问的是否继续进行更新的问题，键入 **yes** 作为响应。

8. 完成后，退出 `clsetup` 实用程序。

9. 在每个群集节点上完成以下子步骤，将各个群集节点重新引导回群集模式：

- a. 引导节点。

- 在基于 SPARC 的系统上，运行以下命令。

`ok boot`

- 在基于 x86 的系统上，运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

10. 验证引导节点时未发生错误，而且节点现在处于联机状态。

```
# cluster status -t node
```

群集互连故障排除

本节提供了一个故障排除过程，用于禁用后再启用群集互连，例如群集传输适配器和传输电缆。

请勿使用 `ipadm` 命令来管理群集传输适配器。如果使用 `ipadm disable-if` 命令禁用了传输适配器，则必须使用 `clinterconnect` 命令禁用传输路径后再将其启用。

此过程要求安装 Oracle Solaris Cluster 软件。这些命令都必须在全局群集节点中运行。

▼ 如何启用群集互连

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面启用群集互连。单击 "Private Interconnects" (专用互连)，单击 "Cables" (电缆)，单击电缆编号以突出显示该电缆，然后单击 "Enable" (启用)。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 检查群集互连的状态。

```
% clinterconnect status
```

```
=== Cluster Transport Paths===
Endpoint1      Endpoint2      Status
-----
pnode1:net1    pnode2:net1    waiting
pnode1:net5    pnode2:net5    Path online
```

2. 禁用群集互连路径。

- a. 检查群集互连路径。

```
% clinterconnect show | egrep -ie "cable.*pnode1"
Transport Cable: pnode1:net5,switch2@1
Transport Cable: pnode1:net1,switch1@1
```

- b. 禁用群集互连路径。

```
% clinterconnect disable pnode1:net1,switch1@1
```

3. 启用群集互连路径。

```
% clinterconnect enable pnode1:net1,switch1@1
```

4. 检验是否已启用群集互连。

```
% clinterconnect status

=== Cluster Transport Paths===
Endpoint1          Endpoint2          Status
-----
pnode1:net1        pnode2:net1        Path online
pnode1:net5        pnode2:net5        Path online
```

管理公共网络

Oracle Solaris Cluster 软件支持将 IPMP、链路聚合和 VNIC 的 Oracle Solaris 软件实现用于公共网络。对群集环境而非群集环境而言，基本的公共网络管理是相同的。

多路径功能随 Oracle Solaris 11 OS 自动安装，必须启用多路径功能才能使用该功能。多路径管理在相应的 Oracle Solaris OS 文档中进行介绍。不过，在 Oracle Solaris Cluster 环境中管理 IPMP、链路聚合和 VNIC 之前请先查看下面的准则。

有关 IPMP 的信息，请参见《在 Oracle Solaris 11.3 中管理 TCP/IP 网络、IPMP 和 IP 隧道》中的第 3 章, “管理 IPMP”。有关链路聚合的信息，请参见《在 Oracle Solaris 11.3 中管理网络数据链路》中的第 2 章, “使用链路聚合配置高可用性”。

如何在群集中管理 IP 网络多路径组

在群集中执行 IPMP 过程之前，请考虑以下准则。

- 在配置使用 SUNW.SharedAddress 网络资源的可伸缩服务资源（在资源类型的资源类型注册文件中 SCALABLE=TRUE）时，可以配置 PNM 以监视群集节点上所有 IPMP 组的状态，以及 SUNW.SharedAddress 配置使用的 IPMP 组的状态。此配置允许在群集节点上的任何 IPMP 组发生故障时，重新启动服务并进行故障转移，以便最大程度地提高与群集节点并存于同一个子网上的网络客户机的服务可用性。例如：

```
# echo ssm_monitor_all > /etc/cluster/pnm/pnm.conf
```

重新引导该节点。

- 对于以太网适配器，local-mac-address? 变量必须具有 true 值。
- 您可以在群集中使用基于探测的 IPMP 组或基于链路的 IPMP 组。基于探测的 IPMP 组用于测试目标 IP 地址，并通过识别更多可能危及可用性的情况来提供最强大的保护。

如果要使用 iSCSI 存储作为法定设备，请确保正确配置基于探测的 IPMP 设备。如果 iSCSI 网络是仅包含群集节点和 iSCSI 存储设备的专用网络，并且在 iSCSI 网络上没有其他主机，则当除了其中一个群集节点以外的所有群集节点都发生故障时，可

能会破坏基于探测的 IPMP 机制。发生该问题的原因在于 iSCSI 网络上没有其他主机供 IPMP 探测，因此当群集中仅剩一个节点时，IPMP 将此视为网络故障。IPMP 会将 iSCSI 网络适配器脱机，从而使其余节点无法访问 iSCSI 存储，因此也无法访问法定设备。要解决此问题，可以在 iSCSI 网络中添加一个路由器，从而使群集以外的其他主机可以响应探测，防止 IPMP 将网络适配器脱机。或者，也可以将 IPMP 配置为基于链路的故障转移，而不是配置为基于探测的故障转移。

- 除非公共网络配置中有一个或多个非链路本地 IPv6 公共网络接口，否则 `scinstall` 实用程序会自动为群集中使用同一子网的每个公共网络适配器集配置一个多适配器 IPMP 组。这些组是基于链路的，具有传递式探测器。如果需要基于探测的故障探测，可添加测试地址。
- 同一个多路径组中的所有适配器的测试 IP 地址必须属于一个 IP 子网。
- 正常的应用程序不得使用测试 IP 地址，因为它们属于高度不可用地址。
- 未对多路径组的命名加以限制。不过，在配置资源组时，`netiflist` 命名惯例是多路径名称后接节点 ID 号或节点名称。例如，如果多路径组的名称为 `sc_ipmp0`，则 `netiflist` 的命名应为 `sc_ipmp0@1` 或 `sc_ipmp0@phys-schost-1`，其中适配器位于节点 ID 为 1 的节点 `phys-schost-1` 上。
- 在未将 IP 地址从要删除的适配器切换到组中的备用适配器（使用 `if_mpadm(1M)` 命令）之前，请不要取消配置（取消激活）或关闭 IP 网络多路径组的适配器。
- 在激活了 Oracle Solaris Cluster HA IP 地址的 IPMP 组中，请勿取消激活或删除网络接口。此 IP 地址可以属于逻辑主机名资源或共享地址资源。但是，如果使用 `ifconfig` 命令取消激活活动接口，Oracle Solaris Cluster 现在可识别此事件。如果 IPMP 组在过程中已变为不可用，它会将资源组故障转移到其他正常节点。如果 IPMP 组有效但缺少 HA IP 地址，Oracle Solaris Cluster 可能还重新启动资源组。出于多种原因，IPMP 组会变得不可用：IPv4 连接中断和/或 IPv6 连接中断。有关更多信息，请参见 `if_mpadm(1M)` 手册页。
- 避免在事先未将适配器从其各自的多路径组中删除的情况下，将其重新连接到其他子网上。
- 即使正在监视多路径组，也可以对适配器进行逻辑适配器操作。
- 您必须为群集中的每个节点至少维护一个公共网络连接。如果没有公共网络连接，就无法访问群集。
- 要查看某个群集上 IP 网络多路径组的状态，请使用 `ipmpstat -g` 命令。有关更多信息，请参见《在 Oracle Solaris 11.3 中管理 TCP/IP 网络、IPMP 和 IP 隧道》中的第 3 章，“管理 IPMP”。

有关群集软件安装过程，请参见《Oracle Solaris Cluster 4.3 软件安装指南》。有关维修公共网络硬件组件的过程，请参见《Oracle Solaris Cluster Hardware Administration Manual》。

动态重新配置公共网络接口

在对群集中的公共网络接口完成动态重新配置 (Dynamic Reconfiguration, DR) 操作时，必须考虑几个问题。

- 针对 Oracle Solaris 动态重新配置功能介绍的所有要求、过程和限制也适用于 Oracle Solaris Cluster 动态重新配置支持（操作系统停止操作除外）。因此，请先阅读关于 Oracle Solaris 动态重新配置功能的文档，然后再对 Oracle Solaris Cluster 软件使用动态重新配置功能。您应该特别注意那些在执行动态重新配置分离操作时影响非网络 IO 设备的问题。
- 只有公共网络接口不活动时，动态重新配置删除板操作才能成功。在删除活动的公共网络接口之前，使用 `if_mpadm` 命令将 IP 地址从要删除的适配器切换到多路径组中的另一个适配器。有关更多信息，请参见 [if_mpadm\(1M\)](#) 手册页。
- 在没有正确地禁用公共网络接口卡（作为活动网络接口适配器）的情况下，如果试图删除此公共网络接口卡，Oracle Solaris Cluster 将拒绝此操作并标识出会受此操作影响的接口。



注意 - 当多路径组中有两个适配器时，如果在对禁用的网络适配器执行动态重新配置删除操作时，另一个网络适配器出现故障，将会影响可用性。执行动态重新配置操作期间另一个适配器无法进行故障转移。

对公共网络接口执行动态重新配置操作时，请按所示顺序完成下列过程。

表 14 任务列表：动态重新配置公共网络接口

任务	说明
1. 使用 <code>if_mpadm</code> 命令将 IP 地址从要删除的适配器切换到多路径组中的另一个适配器	if_mpadm(1M) 手册页。 《在 Oracle Solaris 11.3 中管理 TCP/IP 网络、IPMP 和 IP 隧道》中的“如何将接口从一个 IPMP 组移至另一个 IPMP 组”
2. 使用 <code>ipadm</code> 命令将适配器从多路径组中删除	ipadm(1M) 手册页 《在 Oracle Solaris 11.3 中管理 TCP/IP 网络、IPMP 和 IP 隧道》中的“如何从 IPMP 组中删除接口”
3. 对公共网络接口执行动态重新配置操作	

管理群集节点

本章介绍如何向群集添加节点以及如何删除节点：

- “向群集或区域群集添加节点” [191]
- “恢复群集节点” [194]
- “从群集中删除节点” [198]

有关群集维护任务的信息，请参见第 9 章 管理群集。

向群集或区域群集添加节点

本节介绍如何向全局群集或区域群集添加节点。您可以在全局群集中托管区域群集的节点上创建一个新的区域群集节点，前提是该全局群集节点尚未托管该特定区域群集的节点。

注 - 添加的节点运行的 Oracle Solaris Cluster 软件版本必须与其加入的群集相同。

为每个区域群集节点指定 IP 地址和 NIC 是可选的。

注 - 如果不为每个区域群集节点配置 IP 地址，将出现以下两种情况：

1. 该特定区域群集将无法配置要在区域群集中使用的 NAS 设备。群集在与 NAS 设备通信时将使用区域群集节点的 IP 地址，所以缺失 IP 地址会阻止对隔离 NAS 设备的群集支持。
 2. 群集软件将激活所有 NIC 上的所有逻辑主机 IP 地址。
-

如果原始区域群集节点未指定 IP 地址或 NIC，则无需为新的区域群集节点指定该信息。

在本章中，`phys-schost#` 表示全局群集提示符。`clzonecluster` 交互式 shell 提示符为 `clzc:schost>`。

下表列出了向现有群集中添加节点时所要执行的任务。请按照显示的顺序执行这些任务。

表 15 任务列表：向现有的全局或区域群集添加节点

任务	说明
在节点上安装主机适配器并检验现有的群集互连是否支持该新节点	《Oracle Solaris Cluster Hardware Administration Manual》
添加共享存储	要手动添加共享存储，请按照 《Oracle Solaris Cluster Hardware Administration Manual》 中的说明操作。 您还可以使用 Oracle Solaris Cluster Manager 将共享存储设备添加到区域群集。在 Oracle Solaris Cluster Manager 中导航到区域群集的页面并单击 "Solaris Resources" (Solaris 资源) 选项卡。有关 Oracle Solaris Cluster Manager 登录说明，请参见 如何访问 Oracle Solaris Cluster Manager [266] 。
将节点添加到授权节点列表	<code>claccess allow -h node-being-added</code>
在新的群集节点上安装并配置软件	《Oracle Solaris Cluster 4.3 软件安装指南》 中的第 2 章, "在全局群集节点上安装软件"
向现有群集中添加新节点	如何向现有的群集或区域群集添加节点 [192]
如果该群集是在 Oracle Solaris Cluster Geographic Edition 伙伴关系中配置的，请将新节点配置为该配置中的积极参与者	《Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide》 中的 "How to Add a New Node to a Cluster in a Partnership"

▼ 如何向现有的群集或区域群集添加节点

在向现有全局群集或区域群集添加 Oracle Solaris 主机或虚拟机之前，请确保该节点已经正确安装和配置了所有必要的硬件，包括与专用群集互连之间的有效物理连接。

有关硬件安装信息，请参阅 [《Oracle Solaris Cluster Hardware Administration Manual》](#) 或服务器附带的硬件文档。

使用此过程可将计算机的节点名称添加到群集的授权节点列表中，从而使该计算机将自身安装到该群集中。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- 1. 在当前某个全局群集成员上承担 root 角色。从全局群集的节点中执行这些步骤。
- 2. 确保已正确完成表 15 “任务列表：向现有的全局或区域群集添加节点”的任务列表中列出的所有必要的先决硬件安装和配置任务。
- 3. 在此新群集节点上安装并配置软件。

使用 `scinstall` 实用程序完成新节点的安装和配置，如《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中所述。

4. 在新节点上使用 `scinstall` 实用程序在群集中配置该节点。
5. 要手动向区域群集中添加节点，必须指定 Oracle Solaris 主机和虚拟节点名称。
还必须指定要用于每个节点上的公共网络通信的网络资源。在以下示例中，区域名称为 `sczone`，`sc_ipmp0` 为 IPMP 组名称。

```
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-3
clzc:sczone:node>set hostname=hostname3
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname3
clzc:sczone:node:net>set physical=sc_ipmp0
clzc:sczone:node:net>end
clzc:sczone:node>end
clzc:sczone>exit
```

有关配置节点的详细说明，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[创建和配置区域群集](#)”。

6. 如果新的区域群集节点是 `solaris10` 标记类型，且区域群集上未安装 Oracle Solaris Cluster 软件，请提供 DVD 映像的路径并安装该软件。

```
# clzonecluster install-cluster -d dvd-image zone-cluster-name
```

7. 在配置该节点后，将节点重新引导至群集模式，然后在节点上安装区域群集。

```
# clzonecluster install zone-cluster-name
```

8. 要阻止向群集中添加任何新计算机，请通过 `clsetup` 实用程序键入相应选项的编号，以指示群集忽略添加新计算机的请求。

按回车键。

按照 `clsetup` 提示操作。该选项将通知群集忽略公共网络上尝试将自身添加到群集中的任何新计算机发出的所有请求。

9. 退出 `clsetup` 实用程序。

另请参见 [clsetup\(1CL\)](#) 手册页。

有关添加群集节点的完整任务列表，请参见表 15 “[任务列表：向现有的全局或区域群集添加节点](#)”“[任务列表：添加群集节点](#)”。

要向现有资源组添加节点，请参见《[Oracle Solaris Cluster 4.3 数据服务规划和管理指南](#)》。

恢复群集节点

可以使用统一归档文件恢复群集节点，使其与归档文件完全相同。在您恢复节点之前，必须先在群集节点上创建一个恢复归档文件。只能使用恢复归档文件；无法使用克隆归档文件恢复群集节点。有关创建恢复归档文件的说明，请参见下面的[步骤 1](#)。

此过程提示您输入群集名称、节点名称及其 MAC 地址以及统一归档文件的路径。对于指定的每个归档文件，scinstall 实用程序会验证归档文件的源节点名称是否与要恢复的节点相同。有关从统一归档文件恢复群集节点的说明，请参见[如何从统一归档文件恢复节点 \[194\]](#)。

▼ 如何从统一归档文件恢复节点

此过程在自动化安装程序服务器上使用 scinstall 实用程序的交互形式。您必须设置了 AI 服务器并且从 Oracle Solaris Cluster 系统信息库安装了 ha-cluster/system/install 软件包。归档文件的节点名称必须与要恢复的节点相同。

按照以下准则在此过程中使用交互式的 scinstall 实用程序：

- 交互式的 scinstall 使您可以提前键入。因此，如果未立即显示下一个菜单屏幕，请勿多次按回车键。
- 除非另外指明，否则按 Ctrl-D 键可返回到一系列相关问题的开始处或者返回到主菜单。
- 默认答案或先前会话的答案将显示在问题末尾的方括号 ([]) 中。按回车键即可输入方括号中的答复而无需键入。

1. 在全局群集中的某个节点上承担 **root** 角色，然后创建一个恢复归档文件。

```
phys-schost# archiveadm create -r archive-location
```

当您创建归档文件时，请排除位于共享存储上的 ZFS 数据集。如果您打算恢复共享存储上的数据，请使用传统方法。

有关使用 archiveadm 命令的更多信息，请参见 [archiveadm\(1M\)](#) 手册页。

2. 登录自动化安装程序服务器并承担 **root** 角色。
3. 启动 **scinstall** 实用程序。

```
phys-schost# scinstall
```

4. 键入用于恢复群集的选项号。

```
*** Main Menu ***
```

Please select from one of the following (*) options:

- * 1) Install, restore, or replicate a cluster from this Automated Installer server
- * 2) Securely install, restore, or replicate a cluster from this Automated Installer server
- * 3) Print release information for this Automated Installer install server
- * ?) Help with menu options
- * q) Quit

Option: 2

选择选项 1 将使用非安全 AI 服务器安装恢复群集节点。选择选项 2 将使用安全 AI 服务器安装恢复群集节点。

将显示定制自动化安装程序菜单或者定制安全自动化安装程序菜单。

5. 键入用于从统一归档文件恢复群集节点的选项号。
此时将显示 "Cluster Name" (群集名称) 屏幕。
6. 键入包含要恢复的节点的群集名称。
此时将显示 "Cluster Nodes" (群集节点) 屏幕。
7. 键入要从统一归档文件恢复的群集节点的名称。
每行键入一个节点名称。完成后, 按 Ctrl-D 组合键, 然后键入 yes 并按回车键来确认列表。如果要恢复群集中的所有节点, 则指定所有节点。
如果 scinstall 实用程序找不到节点的 MAC 地址, 请在系统提示时键入每个地址。
8. 键入恢复归档文件的完整路径。
用于恢复节点的归档文件必须是恢复归档文件。用于恢复特定节点的归档文件必须是在相同节点上创建的。为每个要恢复的群集节点重复此操作。
9. 对于每个节点, 确认您选择的选项, 以便 scinstall 实用程序执行必要的配置来从该 AI 服务器安装群集节点。
此实用程序还输出在 DHCP 服务器上添加 DHCP 宏的说明, 并为 SPARC 节点添加 (如果选择安全安装) 或清除安全密钥。请按照这些说明进行操作。
10. (可选) 要定制目标设备, 请针对每个节点更新 AI 清单。
AI 清单位于以下目录中:

```
/var/cluster/logs/install/autosinstall.d/ \
cluster-name/node-name/node-name_aimanifest.xml
```

- a. 要定制目标设备, 请更新清单文件中的 **target** 元素。
请根据您的希望如何使用受支持的条件为安装定位目标设备来更新清单文件中的 **target** 元素。例如, 您可以指定 **disk_name** 子元素。

注 - `scinstall` 假定清单文件中的现有引导磁盘将成为目标设备。要定制目标设备，请更新清单文件中的 `target` 元素。有关更多信息，请参见《[安装 Oracle Solaris 11.3 系统](#)》中的 [部分 III](#), “使用安装服务器安装,”和 [ai_manifest\(4\)](#) 手册页。

- b. 为每个节点运行 `installadm` 命令。

```
# installadm update-manifest -n cluster-name-{sparc|i386} \
-f /var/cluster/logs/install/autoscinstall.d/cluster-name/node-name/node-
name_aimanifest.xml \
-m node-name_manifest
```

请注意，SPARC 和 i386 是群集节点的体系结构。

11. 如果使用的是群集管理控制台，请为群集中的每个节点显示一个控制台屏幕。

- 如果您的管理控制台上安装并配置了 `pconsole` 软件，则可使用 `pconsole` 实用程序显示各个控制台屏幕。

以 `root` 角色使用以下命令启动 `pconsole` 实用程序：

```
adminconsole# pconsole host[:port] [...] &
```

`pconsole` 实用程序还将打开一个主窗口，您可以从该主窗口将您输入的内容同时发送到每个控制台窗口。

- 如果未使用 `pconsole` 实用程序，请分别连接到每个节点的控制台。

12. 关闭然后引导各个节点以启动 AI 安装。

这将以默认配置安装 Oracle Solaris 软件。

注 - 如果要定制 Oracle Solaris 安装，不能使用此方法。如果您选择 Oracle Solaris 交互式安装，则会绕过自动化安装程序并且不会安装和配置 Oracle Solaris Cluster 软件。

要在安装期间定制 Oracle Solaris，请改为按照《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“如何安装 Oracle Solaris 软件”中的说明进行操作，然后按照《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“如何安装 Oracle Solaris Cluster 软件包”中的说明安装并配置群集。

- SPARC：

- a. 关闭各个节点。

```
phys-schost# cluster shutdown -g 0 -y
```

- b. 使用以下命令引导节点

```
ok boot net:dhcp - install
```

注 - 命令中破折号 (-) 的两端需加空格。

■ x86

- a. 重新引导该节点。

```
# reboot -p
```

- b. 在 PXE 引导期间，按下 Ctrl-N 组合键。
此时将显示 "GRUB" 菜单。

- c. 立即选择 "Automated Install"（自动安装）项。

注 - 如果在 20 秒内没有选择 "Automated Install"（自动安装）项，系统将使用默认的交互式文本安装程序方法继续进行安装，此方法不会安装和配置 Oracle Solaris Cluster 软件。

在安装完成之后，每个节点将自动重新引导以加入群集。节点恢复到创建归档文件时的相同状态。安装 Oracle Solaris Cluster 时的输出信息记录在每个节点的 /var/cluster/logs/install/sc_ai_config.log 文件中。

- 13. 从一个节点上，验证是否所有节点都已加入群集。

```
phys-schost# clnode status
```

输出类似于以下内容。

```
=== Cluster Nodes ===
```

```
--- Node Status ---
```

Node Name	Status
phys-schost-1	Online
phys-schost-2	Online
phys-schost-3	Online

有关更多信息，请参见 [clnode\(1CL\)](#) 手册页。

从群集中删除节点

本节提供了有关如何在全局群集或区域群集中删除节点的说明。您还可以从全局群集中删除特定的区域群集。下表列出了从现有群集中删除节点时所要执行的任务。请按照显示的顺序执行这些任务。



注意 - 对于 RAC 配置，如果仅使用此过程删除某个节点，则该删除操作可能会导致该节点在重新引导期间出现紧急情况。有关如何从 RAC 配置删除节点的说明，请参见《适用于 Oracle Real Application Clusters 的 Oracle Solaris Cluster 数据服务指南》中的“如何从选定的节点删除 Support for Oracle RAC”。完成该过程后，按以下相应步骤删除 RAC 配置的节点。

表 16 任务列表：删除节点

任务	说明
将所有资源组和设备组移出要删除的节点。如果您有一个区域群集，请登录到区域群集中，并在要卸载的物理节点上清除区域群集节点。然后，从区域群集中删除节点，再关闭物理节点。	<code>clnode evacuate node</code> 如何从区域群集中删除节点 [198]
如果受影响的物理节点发生故障，则只需从群集中删除该节点即可。	
通过检查允许的主机检验该节点是否可以删除。	<code>claccess show</code>
如果节点无法通过 <code>claccess show</code> 命令列出，则无法删除该节点。授予该节点访问群集配置的权限。	<code>claccess allow -h node-to-remove</code>
从所有设备组中删除该节点。	如何将节点从设备组中删除 (Solaris Volume Manager) [125]
删除与要删除的节点连接的所有法定设备。	如何删除法定设备 [160]
如果您要从双节点群集中删除节点，则此步骤是可选的。	如何从群集中删除最后一个法定设备 [161]
注意，尽管在下一步中删除存储设备之前，您必须先删除法定设备，但是可以在之后立即重新添加该法定设备。	
将要删除的节点置于非群集模式。	如何使节点进入维护状态 [216]
从群集软件配置中删除节点。	如何从群集软件配置中删除节点 [199]
(可选) 从群集节点卸载 Oracle Solaris Cluster 软件。	如何从群集节点卸载 Oracle Solaris Cluster 软件 [220]

▼ 如何从区域群集中删除节点

您可以通过以下方式从区域群集删除节点：停止该节点并将其卸载，然后从配置中删除该节点。如果以后您决定将该节点添加回区域群集中，请按照表 15 “任务列表：向现有的全局或区域群集添加节点”中的说明操作。下面的大部分步骤都是从该全局群集节点中执行的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面关闭区域群集节点，但不删除区域群集节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在全局群集中的某个节点上承担 **root** 角色。
2. 通过指定要删除的区域群集节点及其区域群集来关闭该节点。

```
phys-schost# clzonecluster halt -n node zone-cluster-name
```

还可以在区域群集内使用 `clnode evacuate` 和 `shutdown` 命令。

3. 将节点从区域群集中的所有资源组中删除。

```
phys-schost# clresourcegroup remove-node -n zone-hostname -Z zone-cluster-name rg-name
```

如果您使用步骤 2 的注释中介绍的过程，则资源组应该能够自动删除，因此您可以跳过此步骤。

4. 卸载区域群集节点。

```
phys-schost# clzonecluster uninstall -n node zone-cluster-name
```

5. 从配置中删除该区域群集节点。

使用以下命令：

```
phys-schost# clzonecluster configure zone-cluster-name
clzc:sczone> remove node physical-host=node
clzc:sczone> exit
```

注 - 如果要删除的区域群集节点位于无法访问或者无法加入群集的系统上，请使用 `clzonecluster` 交互式 shell 删除该节点：

```
clzc:sczone> remove -F node physical-host=node
```

如果您使用此方法删除最后一个区域群集节点，则系统会提示删除整个区域群集。如果您选择不删除，则不会删除最后一个节点。此删除具有与 `clzonecluster delete -F zone-cluster-name` 相同的效果。

6. 检验该节点是否已从区域群集中删除。

```
phys-schost# clzonecluster status
```

▼ 如何从群集软件配置中删除节点

执行此过程可从全局群集中删除节点。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 继续执行此过程之前，请确保已将节点从所有资源组、设备组和法定设备配置中删除，并将该节点置于维护状态。
2. 在要删除的节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。从全局群集的节点执行此过程中的所有步骤。
3. 将您要删除的全局群集节点引导到非群集模式下。
对于区域群集节点，在执行此步骤之前，请按照[如何从区域群集中删除节点 \[198\]](#)中的说明操作。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot -x
```

- 在基于 x86 的系统上，运行以下命令。

```
shutdown -g -y -i0
```

```
Press any key to continue
```

- a. 在 GRUB 菜单中，使用方向键选择适当的 Oracle Solaris 条目，然后键入 **e** 编辑其命令。
有关基于 GRUB 的引导的更多信息，请参见《[引导和关闭 Oracle Solaris 11.3 系统](#)》中的“引导系统”。

- b. 在引导参数屏幕中，使用方向键选择内核项，然后键入 **e** 编辑该项。

- c. 在命令中添加 **-x** 以指定将系统引导至非群集模式。

```
[ Minimal BASH-like line editing is supported. For the first word, TAB
lists possible command completions. Anywhere else TAB lists the possible
completions of a device/filename. ESC at any time exits. ]
```

```
grub edit> kernel$ /platform/i86pc/kernel/#ISADIR/unix -B $ZFS-BOOTFS -x
```

- d. 按 **Enter** 键接受所做更改，并返回到引导参数屏幕。
屏幕将显示编辑后的命令。

- e. 键入 **b** 将节点引导至非群集模式。

对内核引导参数命令所做的这一更改在系统引导后将失效。下次重新引导节点时，系统将引导至群集模式。如果希望引导至非群集模式，请重新执行上述步骤，并将 **--x** 选项添加到内核引导参数命令中。

4. 将节点从群集中删除。

a. 从活动节点运行以下命令：

```
phys-schost# clnode clear -F nodename
```

如果资源组具有 `rg_system=true` 设置，则必须将其更改为 `rg_system=false`，这样才能成功执行 `clnode clear -F` 命令。在运行 `clnode clear -F` 之后，将资源组重置回 `rg_system=true`。

b. 从要删除的节点运行以下命令：

```
phys-schost# clnode remove -F
```

注 - 如果要删除的节点不可用或无法再引导，则在任何活动的群集节点上运行以下命令。

```
# clnode clear -F node-to-be-removed
```

通过运行 `clnode status nodename` 检验该节点是否已删除。

如果您要删除群集中的最后一个节点，则该节点必须处于非群集模式下并且该群集中未剩下任何活动节点。

5. 通过另一个群集节点检验是否已删除该节点。

```
phys-schost# clnode status nodename
```

6. 完成节点删除操作。

- 如果打算从已删除的节点中卸载 Oracle Solaris Cluster 软件，请转至[如何从群集节点卸载 Oracle Solaris Cluster 软件 \[220\]](#)。
您还可以选择从群集删除节点，同时卸载 Oracle Solaris Cluster 软件。转到不包含任何 Oracle Solaris Cluster 文件的目录，然后键入 `scinstall -r`。
- 如果不打算从已删除的节点中卸载 Oracle Solaris Cluster 软件，可以通过删除硬件连接从群集中物理删除该节点。
有关说明，请参见 [《Oracle Solaris Cluster Hardware Administration Manual》](#)。

例 73 从群集软件配置中删除节点

此示例说明了如何从群集中删除节点 `phys-schost-2`。应当在您要从群集中删除的节点 (`phys-schost-2`) 上以非群集模式运行 `clnode remove` 命令。

从群集中删除节点：

```
phys-schost-2# clnode remove
phys-schost-1# clnode clear -F phys-schost-2
```

```
验证节点删除：
phys-schost-1# clnode status
-- Cluster Nodes --
           Node name           Status
           -----           -
Cluster node: phys-schost-1   Online
```

另请参见 要从已删除的节点中卸载 Oracle Solaris Cluster 软件，请参见[如何从群集节点卸载 Oracle Solaris Cluster 软件 \[220\]](#)。

有关硬件过程，请参见《[Oracle Solaris Cluster Hardware Administration Manual](#)》。

有关删除群集节点的完整任务列表，请参见表 16 “[任务列表：删除节点](#)”。

要向现有的群集添加节点，请参见[如何向现有的群集或区域群集添加节点 \[192\]](#)。

▼ 如何在节点连接多于两个的群集中删除阵列与单个节点之间的连接

使用此过程可在具有三节点或四节点连通性的群集中从单个群集节点分离存储阵列。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 备份与要删除的存储阵列相关联的所有数据库表、数据服务和卷。
2. 确定要断开连接的节点上正在运行的资源组和设备组。

```
phys-schost# clresourcegroup status
phys-schost# cldevicegroup status
```

3. 如果需要，将所有资源组和设备组移出要断开连接的节点。



注意（仅限 SPARC） - 如果群集正在运行 Oracle RAC 软件，请先关闭在节点上运行的 Oracle RAC 数据库实例，然后再将组从节点移出。有关说明，请参见《[Oracle Database Administration Guide](#)》。

```
phys-schost# clnode evacuate node
```

clnode evacuate 命令可将指定节点上的所有设备组切换到下一个首选节点。该命令还将指定节点上的所有资源组切换到下一个首选的节点。

4. 使设备组处于维护状态。

有关将设备组置于维护状态的过程，请参见[如何使节点进入维护状态 \[216\]](#)。

5. 从设备组中删除节点。

如果您使用的是原始磁盘，请使用 `cldevicegroup(1CL)` 命令删除设备组。

6. 对于每一个包含 `HAStoragePlus` 资源的资源组，请从该资源组的节点列表中删除该节点。

```
phys-schost# clresourcegroup remove-node -n node + | resourcegroup
```

有关更改资源组的节点列表的更多信息，请参见《[Oracle Solaris Cluster 4.3 数据服务规划和管理指南](#)》。

注 - 执行 `clresourcegroup` 命令时，资源类型、资源组和资源属性的名称均区分大小写。

7. 如果要删除的存储阵列是节点上连接的最后一个存储阵列，应断开节点与该存储阵列连接的集线器或交换机之间的光缆。

否则，应跳过此步骤。

8. 如果要从将断开连接的节点上移除主机适配器，请关闭该节点的电源。

如果要从将断开连接的节点上移除主机适配器，请跳到[步骤 11](#)。

9. 从节点上移除主机适配器。

有关移除主机适配器的操作过程，请参见节点的相关文档。

10. 打开节点的电源，但不引导该节点。

11. 如果安装了 Oracle RAC 软件，请将 Oracle RAC 软件包从要断开连接的节点中删除。

```
phys-schost# pkg uninstall /ha-cluster/library/ucmm
```



注意（仅限 SPARC） - 如果不从已断开连接的节点上删除 Oracle RAC 软件，当该节点重新加入群集时，将出现紧急情况并可能导致数据不可用。

12. 以群集模式引导节点。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot
```

- 在基于 x86 的系统上，运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

13. 在节点上，通过更新 `/devices` 和 `/dev` 条目来更新设备名称空间。

```
phys-schost# devfsadm -C
```

```
phys-schost# cldevice refresh
```

14. 使设备组重新联机。

有关将设备组置于联机状态的信息，请参见[如何使节点脱离维护状态 \[218\]](#)。

▼ 如何纠正错误消息

要纠正尝试执行任何群集节点的删除过程时所产生的错误消息，请执行以下过程。

1. 尝试将该节点重新加入全局群集。

此操作过程仅适用于全局群集。

```
phys-schost# boot
```

2. 该节点是否成功加入群集？

- 如果不是，请继续执行[步骤 2b](#)。
- 如果是，请执行以下步骤从设备组中删除该节点。
 - a. 如果该节点重新加入群集成功，请从其余的设备组中删除该节点。
请按照[如何将节点从所有设备组中删除 \[124\]](#)中的过程操作。
 - b. 从所有设备组中删除该节点后，请返回到[如何从群集节点卸载 Oracle Solaris Cluster 软件 \[220\]](#)，并重复执行其中的过程。

3. 如果该节点未能重新加入群集，请将该节点的 `/etc/cluster/ccr` 文件重命名为您所选的任何其他名称，例如 `ccr.old`。

```
# mv /etc/cluster/ccr /etc/cluster/ccr.old
```

4. 返回到[如何从群集节点卸载 Oracle Solaris Cluster 软件 \[220\]](#)，并重复执行其中的过程。

管理群集

本章介绍会影响整个全局群集或某个区域群集的管理过程：

- “[管理群集概述](#)” [205]
- “[执行区域群集管理任务](#)” [231]
- “[用于测试的故障排除过程](#)” [241]

有关在群集中添加或删除节点的更多信息，请参见[第 8 章 管理群集节点](#)。

管理群集概述

本节介绍了如何对整个全局群集或区域群集执行管理任务。下表列出了这些管理任务及相关过程。通常在全局区域中执行群集管理任务。要管理区域群集，必须以群集模式启动至少一台将托管该区域群集的计算机。不需要启动并运行所有区域群集节点；当目前不在群集中的节点重新加入该群集时，Oracle Solaris Cluster 将重放任何配置更改。

注 - 默认情况下，电源管理功能处于禁用状态，因此不会干扰群集。如果您对单节点群集启用电源管理功能，该群集仍会运行，但可能有几秒钟不可用。电源管理功能会尝试关闭节点，但不会成功。

在本章中，`phys-schost#` 表示全局群集提示符。`clzonecluster` 交互式 shell 提示符为 `clzc:schost>`。

表 17 任务列表：管理群集

任务	说明
在群集中添加或删除节点	第 8 章 管理群集节点
更改群集的名称	如何更改群集名称 [206]
列出节点的 ID 及其相应的节点名称	如何将节点 ID 映射到节点名称 [207]
允许或拒绝新节点添加到群集中	如何使用对新群集节点的验证 [208]
使用 NTP 更改群集的时间	如何在群集中重置时间 [209]
关闭节点以显示 OpenBoot PROM ok 提示符 (在基于 SPARC 的系统上) 或在 GRUB 菜单	如何在节点上显示 OpenBoot PROM (OBP) [211]

任务	说明
中显示消息 "Press any key to continue" (在基于 x86 的系统上)。	
添加或更改专用主机名	如何更改节点专用主机名 [212] 还可以使用 Oracle Solaris Cluster Manager 浏览器界面 向全局群集或区域群集添加逻辑主机名。单击 "Tasks" (任务)，然后单击 "Logical Hostname" (逻辑主机 名) 以启动向导。有关 Oracle Solaris Cluster Manager 登录说明，请参见 如何访问 Oracle Solaris Cluster Manager [266] 。
使群集节点进入维护状态	如何使节点进入维护状态 [216]
重命名节点	如何重命名节点 [214]
使群集节点脱离维护状态	如何使节点脱离维护状态 [218]
从群集节点卸载群集软件	如何从群集节点卸载 Oracle Solaris Cluster 软件 [220]
添加和管理 SNMP 事件 MIB	如何启用 SNMP 事件 MIB [224] 如何在节点上添加 SNMP 用户 [228]
配置每个节点的负载限制	如何在节点上配置负载限制 [230]
移动区域群集；为应用程序准备区域群集，删除区域群集	“执行区域群集管理任务” [231]

▼ 如何更改群集名称

根据需要，您可以在初次安装后更改群集的名称。



注意 - 如果群集为 Oracle Solaris Cluster Geographic Edition 伙伴关系成员，则不要执行此过程。应改为按照《[Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide](#)》中的“[Renaming a Cluster That Is in a Partnership](#)”中的过程执行操作。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在全局群集中的任一节点上承担 **root** 角色。
2. 启动 **clsetup** 实用程序。

```
phys-schost# clsetup
```

这时将显示主菜单。

3. 要更改群集名称，请键入与其他群集属性选项对应的编号。

这时将显示 "Other Cluster Properties" (其他群集属性) 菜单。

4. 从菜单进行选择并按屏幕提示操作。
5. 如果想让 Oracle Solaris Cluster 的服务标签反映新的群集名称，请删除现有的 Oracle Solaris Cluster 标签并重新启动群集。

要删除 Oracle Solaris Cluster 服务标签实例，请针对群集中的所有节点完成以下子步骤。

- a. 列出所有的服务标签。

```
phys-schost# stclient -x
```

- b. 查找 Oracle Solaris Cluster 服务标签实例编号，然后运行以下命令。

```
phys-schost# stclient -d -i service_tag_instance_number
```

- c. 重新引导群集中的所有节点。

```
phys-schost# reboot
```

例 74 更改群集的名称

以下示例显示了由 clsetup 实用程序生成的 cluster 命令将群集名称更改为 dromedary。

```
phys-schost# cluster rename -c dromedary
```

有关更多信息，请参见 [cluster\(1CL\)](#) 和 [clsetup\(1CL\)](#) 手册页。

▼ 如何将节点 ID 映射到节点名称

在 Oracle Solaris Cluster 安装期间，会自动为每个节点分配一个唯一的节点 ID 号。该 ID 号是按节点首次加入群集的顺序分配的。节点 ID 号一经指定，便不能再更改。节点 ID 号经常在错误消息中使用，标识与消息有关的群集节点。请按照此过程来确定节点 ID 和节点名称之间的映射。

您无需成为 root 角色即可列出全局群集或区域群集的配置信息。此过程的其中一个步骤是从全局群集的一个节点上执行的。另一个步骤是从区域群集节点执行的。

1. 使用 **clnode** 命令列出全局群集的群集配置信息。

```
phys-schost# clnode show | grep Node
```

有关更多信息，请参见 [clnode\(1CL\)](#) 手册页。

2. (可选) 列出区域群集的节点 ID。

区域群集节点具有与它所运行在全局群集节点相同的节点 ID。

```
phys-schost# zlogin sczone clnode -v | grep Node
```

例 75 将节点 ID 映射到节点名称

下面的示例显示了全局群集的节点 ID 分配。

```
phys-schost# clnode show | grep Node
=== Cluster Nodes ===
Node Name:    phys-schost1
Node ID:      1
Node Name:    phys-schost2
Node ID:      2
Node Name:    phys-schost3
Node ID:      3
```

▼ 如何使用对新群集节点的验证

使用 Oracle Solaris Cluster 可以确定新节点是否可将自身添加到全局群集中以及要使用的验证的类型。您可以允许任何新的节点通过公共网络加入群集、拒绝新节点加入群集或指定可以加入群集的具体节点。

新节点可以通过使用标准 UNIX 或者 Diffie-Hellman (DES) 验证来进行验证。如果选择的是 DES 验证，还必须在节点加入前配置所有需要的加密密钥。有关更多信息，请参见 [keyserv\(1M\)](#) 和 [publickey\(4\)](#) 手册页。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在全局群集中的任一节点上承担 **root** 角色。
2. 启动 **clsetup** 实用程序。

```
phys-schost# clsetup
```

这时将显示主菜单。

3. 要使用群集验证，请键入与新节点选项对应的编号。
这时将显示 "New Nodes" (新节点) 菜单。
4. 从菜单进行选择并按屏幕提示操作。

例 76 防止将新计算机添加到全局群集中

clsetup 实用程序生成 claccess 命令。下面的示例显示了可防止将新计算机添加到群集的 claccess 命令。

```
phys-schost# claccess deny -h hostname
```

例 77 允许将所有新计算机添加到全局群集中

clsetup 实用程序生成 claccess 命令。下面的示例显示了使所有新计算机都可以添加到群集的 claccess 命令。

```
phys-schost# claccess allow-all
```

例 78 指定要添加到全局群集中的新计算机

clsetup 实用程序生成 claccess 命令。下面的示例显示了使单个新计算机可以添加到群集的 claccess 命令。

```
phys-schost# claccess allow -h hostname
```

例 79 将验证设置为标准 UNIX

clsetup 实用程序生成 claccess 命令。下面的示例显示了使加入群集的新节点重置为标准 UNIX 验证的 claccess 命令。

```
phys-schost# claccess set -p protocol=sys
```

例 80 将验证设置为 DES

clsetup 实用程序生成 claccess 命令。下面的示例显示了对加入群集的新节点使用 DES 验证的 claccess 命令。

```
phys-schost# claccess set -p protocol=des
```

如果采用 DES 验证，您还必须配置所有必要的加密密钥，然后才能将节点加入群集。有关更多信息，请参见 [keyserv\(1M\)](#) 和 [publickey\(4\)](#) 手册页。

▼ 如何在群集中重置时间

Oracle Solaris Cluster 软件使用 NTP 来维护群集节点间的时间同步。节点进行时间同步时，全局群集会根据需要自动进行调整。有关更多信息，请参见《[Oracle Solaris Cluster 4.3 Concepts Guide](#)》和位于 <http://download.oracle.com/docs/cd/E19065-01/>

[servers.10k/](#) 的《*Network Time Protocol's User's Guide*》（《网络时间协议用户指南》）。



注意 - 如果使用的是 NTP，请不要在群集处于打开和运行状态时调整群集时间。不要使用 `date`、`rdate` 或 `svcadm` 命令以交互方式调整时间，也不要再在 `cron` 脚本中调整时间。有关更多信息，请参见 [date\(1\)](#)、[rdate\(1M\)](#)、[svcadm\(1M\)](#) 或 [cron\(1M\)](#) 手册页。`ntpd(1M)` 手册页是在 `service/network/ntp` Oracle Solaris 11 软件包中提供的。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在全局群集中的任一节点上承担 `root` 角色。
2. 关闭全局群集。

```
phys-schost# cluster shutdown -g0 -y -i0
```

3. 检验该节点是否显示 `ok` 提示符（在基于 SPARC 的系统上）或在 GRUB 菜单中显示消息 "Press any key to continue"（在基于 x86 的系统上）。
4. 以非群集模式引导节点。

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot -x
```

- 在基于 x86 的系统上，运行以下命令。

```
# shutdown -g -y -i0
```

```
Press any key to continue
```

- a. 在 GRUB 菜单中，使用方向键选择适当的 Oracle Solaris 条目，然后键入 `e` 编辑其命令。

GRUB 菜单随即显示。

有关基于 GRUB 的引导的更多信息，请参见《[引导和关闭 Oracle Solaris 11.3 系统](#)》中的“引导系统”。

- b. 在引导参数屏幕中，使用方向键选择内核项，然后键入 `e` 编辑该项。
GRUB 引导参数屏幕随即显示。
- c. 在命令中添加 `-x` 以指定将系统引导至非群集模式。

```
[ Minimal BASH-like line editing is supported. For the first word, TAB
```

lists possible command completions. Anywhere else TAB lists the possible completions of a device/filename. ESC at any time exits.]

```
grub edit> kernel$ /platform/i86pc/kernel/$ISADIR/unix _B $ZFS-BOOTFS -x
```

- d. 按 Enter 键接受所做更改，并返回到引导参数屏幕。

屏幕将显示编辑后的命令。

- e. 键入 **b** 将节点引导至非群集模式。

注 - 对内核引导参数命令所做的这一更改在系统引导后将失效。下次重新引导节点时，系统将引导至群集模式。如果希望引导至非群集模式，请重新执行上述步骤，并将 **-x** 选项添加到内核引导参数命令中。

5. 在单个节点上，通过运行 **date** 命令设置时间。

```
phys-schost# date HHMM.SS
```

6. 在其他计算机上，通过运行 **rddate(1M)** 命令使其时间与上述节点的时间同步。

```
phys-schost# rdate hostname
```

7. 引导每个节点以重新启动该群集。

```
phys-schost# reboot
```

8. 检验是否所有群集节点都已进行了更改。

在每个节点上，运行 **date** 命令。

```
phys-schost# date
```

▼ SPARC: 如何在节点上显示 OpenBoot PROM (OBP)

如果需要配置或更改 OpenBoot™ PROM 设置，请使用此过程。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 连接到要关闭的节点上的控制台。

```
# telnet tc_name tc_port_number
```

tc_name 指定终端集中器的名称。

`tc_port_number` 指定终端集中器上的端口号。端口号取决于配置。通常，端口 2 和 3（5002 和 5003）供站点上所安装的第一个群集使用。

2. 依次执行 `clnode evacuate` 命令和 `shutdown` 命令使群集节点正常关机。

`clnode evacuate` 命令可将指定节点上的所有设备组切换到下一个首选节点。该命令还将全局群集的指定节点中的所有资源组切换到下一个首选的节点。

```
phys-schost# clnode evacuate node
# shutdown -g0 -y
```



注意 - 不要在群集控制台上使用 `send brk` 来关闭群集节点。

3. 执行 OBP 命令。

▼ 如何更改节点专用主机名

使用此过程可在安装完成后更改群集节点的专用主机名。

首次安装群集时，系统会指定默认专用主机名。默认专用主机名的格式为：`clusternodenodeid-priv`，例如：`clusternode3-priv`。只有当专用主机名已在域中使用，您才能更改它。



注意 - 不要尝试给新的专用主机名分配 IP 地址。群集软件将对其进行分配。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在群集的所有节点上，禁用所有数据服务资源或其他可能缓存有专用主机名的应用程序。

```
phys-schost# clresource disable resource[,...]
```

要禁用的应用程序应包括：

- HA-DNS 和 HA-NFS 服务（如果已配置）
- 已通过定制操作配置为使用该专用主机名的任何应用程序
- 任何正由客户机通过专用互连使用的应用程序

有关使用 `clresource` 命令的信息，请参见 [clresource\(1CL\)](#) 手册页和《[Oracle Solaris Cluster 4.3 数据服务规划和管理指南](#)》。

2. 如果 NTP 配置文件引用了您要更改的专用主机名，请在群集的每个节点上关闭 NTP 守护进程。

可使用 `svcadm` 命令关闭 NTP 守护进程。有关 NTP 守护进程的更多信息，请参见 [svcadm\(1M\)](#) 手册页。

```
phys-schost# svcadm disable ntp
```

3. 运行 `clsetup` 实用程序以更改相应节点的专用主机名。
请仅从群集中的某一个节点运行此实用程序。有关更多信息，请参见 [clsetup\(1CL\)](#) 手册页。

注 - 选择新的专用主机名时，请确保该名称对于群集节点是唯一的。

也可以运行 `clnode` 命令以代替 `clsetup` 实用程序来更改专用主机名。在下面的示例中，群集节点名称为 `phys-schost-1`。在运行下面的 `clnode` 命令后，转到 [步骤 6](#)。

```
phys-schost# clnode set -p privatehostname=New-private-nodename phys-schost-1
```

4. 在 `clsetup` 实用程序中，键入专用主机名选项对应的编号。
5. 在 `clsetup` 实用程序中，键入与更改专用主机名选项对应的编号。
在系统进行提示时回答问题。系统会要求您提供要更改专用主机名的节点的名称 (`clusternodeid-priv`) 以及新的专用主机名。
6. 刷新名称服务高速缓存。
请在群集每个节点上执行此步骤。刷新可防止群集应用程序和数据服务尝试访问旧的专用主机名。

```
phys-schost# nscd -i hosts
```

7. 如果您在 NTP 配置或 `include` 文件中更改了专用主机名，请在每个节点上更新 NTP 文件。
如果您在 NTP 配置文件 (`/etc/inet/ntp.conf`) 中更改了某个专用主机名，并且 NTP 配置文件 (`/etc/inet/ntp.conf.include`) 中存在对等主机条目或存在指向对等主机的 `include` 文件的指针，请在每个节点上更新该文件。如果您在 NTP `include` 文件中更改了专用主机名，请在每个节点上更新 `/etc/inet/ntp.conf.sc` 文件。
 - a. 使用您选择的编辑工具。
如果在安装时执行此步骤，还要记得删除所配置的节点的名称。通常，各个群集节点上的 `ntp.conf.scr` 文件完全相同。
 - b. 检验是否能从所有群集节点成功 ping 到新的专用主机名。
 - c. 重新启动 NTP 守护进程。

请在群集的每个节点上执行此步骤。

使用 `svcadm` 命令重新启动 NTP 守护进程。

```
# svcadm enable svc:network/ntp:default
```

8. 启用在[步骤 1](#)中禁用的所有数据服务资源和其他应用程序。

```
phys-schost# clresource enable resource[,...]
```

有关使用 `clresource` 命令的信息，请参见 [clresource\(1CL\)](#) 手册页和《[Oracle Solaris Cluster 4.3 数据服务规划和管理指南](#)》。

例 81 更改专用主机名

以下示例在节点 `phys-schost-2` 上将专用主机名从 `clusternode2-priv` 更改为 `clusternode4-priv`。在每个节点上执行此操作。

Disable all applications and data services as necessary

```
phys-schost-1# svcadm disable ntp
```

```
phys-schost-1# clnode show | grep node
```

```
...
private hostname:                clusternode1-priv
private hostname:                clusternode2-priv
private hostname:                clusternode3-priv
...
```

```
phys-schost-1# clsetup
```

```
phys-schost-1# nscd -i hosts
```

```
phys-schost-1# pfedit /etc/inet/ntp.conf.sc
```

```
...
peer clusternode1-priv
peer clusternode4-priv
peer clusternode3-priv
phys-schost-1# ping clusternode4-priv
```

```
phys-schost-1# svcadm enable ntp
```

Enable all applications and data services disabled at the beginning of the procedure

▼ 如何重命名节点

您可以更改 Oracle Solaris Cluster 配置中某个节点的名称。必须先重命名 Oracle Solaris 主机名，才能重命名节点。使用 `clnode rename` 命令重命名节点。

以下说明适用于全局群集中运行的任何应用程序。

1. 在全局群集中，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. （可选）如果要重命名作为伙伴关系成员的 Oracle Solaris Cluster Geographic Edition 群集中的节点，请确定是否切换保护组。

如果您执行重命名过程时所在的群集是保护组的主群集，并且您希望将保护组中的应用程序联机，则可在重命名过程中将保护组切换到辅助群集。

有关 Geographic Edition 群集和节点的更多信息，请参见《[Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide](#)》中的第 5 章, “Administering Cluster Partnerships”。

3. 重命名 Oracle Solaris 主机名。

完成《[在 Oracle Solaris 11.3 中管理系统信息、进程和性能](#)》中的“如何更改系统标识”中的步骤，除了不要执行过程结尾的重新引导操作。

而是在完成这些步骤后执行群集关闭。

4. 将所有群集节点引导至非群集模式。

```
ok> boot -x
```

5. 在非群集模式下，在您重命名了 Oracle Solaris 主机名的节点上，重命名节点并在每个重命名后的主机上运行 `cmd` 命令。

一次重命名一个节点。

```
# clnode rename -n new-node old-node
```

6. 更新群集上运行的应用程序中现有的对以前主机名的任何引用。

7. 检查命令消息和日志文件以确认节点已被重命名。

8. 将所有节点重新引导至群集模式。

```
# sync;sync;sync;reboot
```

9. 检验节点是否显示新名称。

```
# clnode status -v
```

10. 更新 Geographic Edition 配置以使用新的群集节点名称。

保护组以及数据复制产品使用的配置信息可能会指定节点名称。

11. 您可以选择更改逻辑主机名资源的 `hostnamelist` 属性。

有关此可选步骤的说明，请参见[如何更改现有 Oracle Solaris Cluster 逻辑主机名资源使用的逻辑主机名 \[216\]](#)。

▼ 如何更改现有 Oracle Solaris Cluster 逻辑主机名资源使用的逻辑主机名

您可以选择在按照[如何重命名节点 \[214\]](#)中的步骤重命名节点之前或之后来更改逻辑主机名资源的 `HostnameList` 属性。此步骤是可选的。

1. 在全局群集中，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. (可选) 您可以更改任何现有 Oracle Solaris Cluster 逻辑主机名资源使用的逻辑主机名。

以下步骤说明了如何配置 `apache-lh-res` 资源以使用新的逻辑主机名，这些步骤必须在群集模式下执行。

- a. 在群集模式中，使包含逻辑主机名的 Apache 资源组脱机。

```
# clresourcegroup offline apache-rg
```

- b. 禁用 Apache 逻辑主机名资源。

```
# clresource disable apache-lh-res
```

- c. 提供新主机名列表。

```
# clresource set -p HostnameList=test-2 apache-lh-res
```

- d. 在 `hostnameList` 属性中更改应用程序对以前条目的引用来引用新条目。

- e. 启用新的 Apache 逻辑主机名资源。

```
# clresource enable apache-lh-res
```

- f. 使 Apache 资源组联机。

```
# clresourcegroup online -eM apache-rg
```

- g. 通过运行以下命令检查客户机，确认应用程序正确启动。

```
# clresource status apache-rs
```

▼ 如何使节点进入维护状态

如果要使某个全局群集节点在很长一段时间内停止服务，请将该节点置于维护状态。这样，在维护节点时，该节点不参与法定计数。要将某个节点置于维护状态，必须使用

`clnode evacuate` 和 `shutdown` 命令关闭该节点。有关更多信息，请参见 [clnode\(1CL\)](#) 和 [cluster\(1CL\)](#) 手册页。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面清除节点以及将所有资源组和设备组切换到下一个首选节点。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

当群集节点关闭并置于维护状态后，配置到该节点端口的所有法定设备的法定选票计数均会减 1。当该节点脱离维护模式并恢复联机状态后，节点和法定设备的选票计数会递增 1。

注 - Oracle Solaris `shutdown` 命令关闭单个节点，而 `cluster shutdown` 命令关闭整个群集。

从仍是群集成员的一个节点使用 `clquorum disable` 命令可将另一个群集节点置于维护状态。有关更多信息，请参见 [clquorum\(1CL\)](#) 手册页。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在要置于维护状态的全局群集节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。

2. 从节点中清除任何资源组和设备组。

`clnode evacuate` 命令可将指定节点上的所有资源组和设备组切换到下一个首选节点。

```
phys-schost# clnode evacuate node
```

3. 关闭已清除的节点。

```
phys-schost# shutdown -g0 -y -i0
```

4. 在群集中的另一个节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色，并将[步骤 3](#)中关闭的节点置于维护状态。

```
phys-schost# clquorum disable node
```

5. 检验全局群集节点现在是否处于维护状态。

```
phys-schost# clquorum status node
```

如果节点置于维护状态，则它的 Status 值应为 `offline`，`Present` 和 `Possible` 法定选票数均应为 0（零）。

例 82 将全局群集节点置于维护状态

以下示例将一个群集节点置于维护状态并检验结果。clnode status 的输出内容显示，phys-schost-1 的 Node votes 值为 0（零），状态为 Offline。Quorum Summary 也应显示选票计数已减少。根据您的配置，Quorum Votes by Device 的输出信息可能也会表明某些法定磁盘设备已脱机。

```
[On the node to be put into maintenance state:]
phys-schost-1# clnode evacuate phys-schost-1
phys-schost-1# shutdown -g0 -y -i0

[On another node in the cluster:]
phys-schost-2# clquorum disable phys-schost-1
phys-schost-2# clquorum status phys-schost-1

-- Quorum Votes by Node --
```

Node Name	Present	Possible	Status
phys-schost-1	0	0	Offline
phys-schost-2	1	1	Online
phys-schost-3	1	1	Online

另请参见 要使节点恢复联机状态，请参见[如何使节点脱离维护状态 \[218\]](#)。

▼ 如何使节点脱离维护状态

使用以下过程可使全局群集节点恢复联机状态，并将法定选票计数重置为默认值。对于群集节点，默认法定计数为 1。对于法定设备，默认法定计数为 $N-1$ ，其中 N 是具有指向该法定设备的端口且选票计数不为零的节点的数目。

当节点置于维护状态后，其法定选票计数会减 1。所有配置了到该节点的端口的法定设备也将减少其法定选票计数。重置法定选票计数并使节点脱离维护状态后，该节点的法定选票计数和法定设备选票计数均会加 1。

只要在全局群集节点已置于维护状态的情况下运行此过程，即可使该节点脱离维护状态。



注意 - 如果您既未指定 globaldev 选项，也未指定 node 选项，则会重置整个群集的法定计数。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在全局群集中的任一不处于维护状态的节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 根据全局群集配置中的节点数，执行以下步骤之一：
 - 如果群集配置中有两个节点，请转到[步骤 4](#)。
 - 如果群集配置中有两个以上的节点，请转到[步骤 3](#)。
3. 如果要脱离维护状态的节点将具有法定设备，请从某个不处于维护状态的节点上重置群集法定计数。
重新引导处于维护状态的节点之前，必须先从其他任一节点上重置法定计数，否则该节点可能会挂起，等待达到法定数目。

```
phys-schost# clquorum reset
```

4. 引导要脱离维护状态的节点。
5. 检验法定选票计数。

```
phys-schost# clquorum status
```

如果节点已脱离维护状态，它的状态应为 `online`，且 `Present` 和 `Possible` 法定选票均显示相应的选票计数。

例 83 使群集节点脱离维护状态并重置法定选票计数

以下示例将群集节点及其法定设备的法定计数重置为默认值并检验结果。`cluster status` 的输出内容显示，`phys-schost-1` 的 `Node votes` 值为 1，状态为 `online`。`Quorum Summary` 还应显示选票计数的增加。

```
phys-schost-2# clquorum reset
```

- 在基于 SPARC 的系统上，运行以下命令。

```
ok boot
```

- 在基于 x86 的系统上，运行以下命令。

显示 GRUB 菜单后，选择相应的 Oracle Solaris 条目，然后按 Enter 键。

```
phys-schost-1# clquorum status
```

```
--- Quorum Votes Summary ---
```

Needed	Present	Possible
-----	-----	-----
4	6	6

```
--- Quorum Votes by Node ---
```

```
Node Name      Present    Possible    Status
-----
phys-schost-2  1          1           Online
phys-schost-3  1          1           Online

--- Quorum Votes by Device ---

Device Name      Present    Possible    Status
-----
/dev/did/rdisk/d3s2  1          1           Online
/dev/did/rdisk/d17s2 0          1           Online
/dev/did/rdisk/d31s2 1          1           Online
\
```

▼ 如何从群集节点卸载 Oracle Solaris Cluster 软件

在将某个全局群集节点从完全建立的群集配置断开之前，执行此过程从该节点取消配置 Oracle Solaris Cluster 软件。您可以使用此过程从群集中剩余的最后一个节点中卸载软件。

注 - 如果要从尚未加入群集的节点或仍处于安装模式的节点中卸载 Oracle Solaris Cluster 软件，请不要执行此过程。而应转至 [《Oracle Solaris Cluster 4.3 软件安装指南》](#) 中的“如何取消 Oracle Solaris Cluster 软件的配置以更正安装问题”。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 要删除某个群集节点，请确保您已经正确地完成了任务列表中的所有先决任务。

请参见表 16 “任务列表：删除节点”。

继续执行此过程之前，请确保您已使用 `clnode remove` 命令从群集配置中删除该节点。其他步骤可能包括将您计划卸载的节点添加到群集的节点验证列表中、卸载区域群集，等等。

注 - 要取消配置节点但保留节点上安装的 Oracle Solaris Cluster 软件，则在运行 `clnode remove` 命令后不要再继续执行进一步操作。

2. 在要卸载的节点上承担 **root** 角色。
3. 如果节点具有用于全局设备名称空间的专用分区，请将全局群集节点重新引导至非群集模式。

- 在基于 SPARC 的系统上，运行以下命令。

```
# shutdown -g0 -y -i0 ok boot -x
```

- 在基于 x86 的系统上，运行以下命令。

```
# shutdown -g0 -y -i0
...
<<< Current Boot Parameters >>>
Boot path: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/pci8086,341a@7,1/
sd@0,0:a
Boot args:

Type      b [file-name] [boot-flags] <ENTER>  to boot with options
or        i <ENTER>                          to enter boot interpreter
or        <ENTER>                            to boot with defaults

<<< timeout in 5 seconds >>>
Select (b)oot or (i)nterpreter: b -x
```

4. 在 `/etc/vfstab` 文件中，删除除 `/global/.devices` 全局挂载点以外的所有以全局方式挂载的文件系统条目。
5. 将该节点重新引导至非群集模式。

- 在基于 SPARC 的系统上，执行以下命令：

```
ok boot -x
```

- 在基于 x86 的系统上，执行以下命令：

- a. 在 GRUB 菜单中，使用方向键选择适当的 Oracle Solaris 条目，然后键入 `e` 编辑其命令。

有关基于 GRUB 的引导的更多信息，请参见 [《引导和关闭 Oracle Solaris 11.3 系统》](#) 中的“引导系统”。

- b. 在引导参数屏幕中，使用方向键选择 `kernel` 项，然后键入 `e` 编辑该项。
- c. 在命令中添加 `-x` 以指定将系统引导至非群集模式。
- d. 按 `Enter` 键接受更改，并返回到引导参数屏幕。
屏幕将显示编辑后的命令。
- e. 键入 `b` 将节点引导至非群集模式。

注 - 对内核引导参数命令所做的这一更改在系统引导后将失效。下次重新引导节点时，系统将引导至群集模式。如果希望引导至非群集模式，请执行上述步骤以再次将 `-x` 选项添加到内核引导参数命令中。

6. 转到不包含 Oracle Solaris Cluster 软件包提供的任何文件的目录，如根 (`/`) 目录。

```
phys-schost# cd /
```

7. 要取消配置节点并删除 Oracle Solaris Cluster 软件，请运行以下命令。

```
phys-schost# scinstall -r [-b bename]
```

`-r`

删除群集配置信息并从群集节点卸载 Oracle Solaris Cluster 框架和数据服务软件。然后，您可以重新安装节点或从群集删除节点。

`-b bootenvironmentname`

指定新引导环境的名称，在卸载过程完成后将引导到该引导环境。指定名称这一操作是可选的。如果您没有为引导环境指定名称，将会自动生成一个名称。

有关更多信息，请参见 [scinstall\(1M\)](#) 手册页。

8. 如果卸载完成后要在此节点上重新安装 Oracle Solaris Cluster 软件，请重新引导节点以引导至新引导环境。
9. 如果不打算在此群集上重新安装 Oracle Solaris Cluster 软件，请断开与其他群集设备之间的传输电缆和传输交换机（如果有）。
 - a. 如果卸载的节点与使用并行 SCSI 接口的存储设备相连接，请在断开传输电缆的连接后将 SCSI 端接器安装到存储设备的开路 SCSI 连接器。
如果卸载的节点与使用光纤通道接口的存储设备连接，则不需要端接器。
 - b. 有关断开连接过程，请根据您的主机适配器和服务器附带的文档进行操作。

提示 - 有关将全局设备名称空间迁移到 `lofi` 设备的更多信息，请参见[“迁移全局设备名称空间” \[115\]](#)。

对节点卸载进行故障排除

本节介绍运行 `clnode remove` 命令时可能收到的错误消息以及相应的纠正措施。

未删除的群集文件系统条目

以下错误消息表示已删除的全局群集节点在其 `vfstab` 文件中仍引用了群集文件系统。

```
Verifying that no unexpected global mounts remain in /etc/vfstab ... failed
clnode: global-mount1 is still configured as a global mount.
clnode: global-mount1 is still configured as a global mount.
clnode: /global/dg1 is still configured as a global mount.

clnode: It is not safe to uninstall with these outstanding errors.
clnode: Refer to the documentation for complete uninstall instructions.
clnode: Uninstall failed.
```

要纠正此错误，请返回到[如何从群集节点卸载 Oracle Solaris Cluster 软件 \[220\]](#)，并重复执行其中的过程。在重新运行 `clnode remove` 命令前，请确保已成功完成此过程中的[步骤 4](#)。

设备组中列出的未删除项

以下错误消息表明已删除的节点仍列在某个设备组中。

```
Verifying that no device services still reference this node ... failed
clnode: This node is still configured to host device service "service".
clnode: This node is still configured to host device service "service2".
clnode: This node is still configured to host device service "service3".
clnode: This node is still configured to host device service "dg1".

clnode: It is not safe to uninstall with these outstanding errors.
clnode: Refer to the documentation for complete uninstall instructions.
clnode: Uninstall failed.
```

创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB

本节介绍了如何创建、设置和管理简单网络管理协议 (Simple Network Management Protocol, SNMP) 事件管理信息库 (Management Information Base, MIB)。此外，本节还介绍了如何启用、禁用和更改 Oracle Solaris Cluster SNMP 事件 MIB。

Oracle Solaris Cluster 软件目前支持一个 MIB，即事件 MIB。SNMP 管理器软件可实时捕获群集事件。一经启用，SNMP 管理器就会自动向 `clsnmphost` 命令所定义的所有主机发送陷阱通知。由于群集会生成大量通知，因此只有严重程度为 `min_severity` 或更高的事件才作为陷阱通知发送。默认情况下，`min_severity` 值设置为 `NOTICE`。`log_number` 值指定在停止使用较旧的条目之前要在 MIB 表中记录的事件的数量。MIB 可维护一个只读表，其中包含最近发生的已作为陷阱发送的事件。事件数受 `log_number` 值限制。系统重新引导后此信息将不再存在。

SNMP 事件 MIB 是在 `sun-cluster-event-mib.mib` 文件中定义的，并位于 `/usr/cluster/lib/mib` 目录中。您可以使用此定义来解释 SNMP 陷阱信息。

群集事件 SNMP 接口将通用代理容器 (cacao) SNMP 适配器用作其 SNMP 代理基础结构。默认情况下，SNMP 的端口号是 11161，SNMP 陷阱的默认端口是 11162。可使用 `cacaoadm` 命令更改这些端口号。有关更多信息，请参见 `cacaoadm(1M)` 手册页。

创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB 可能涉及以下任务。

表 18 任务列表：创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB

任务	说明
启用 SNMP 事件 MIB	如何启用 SNMP 事件 MIB [224]
禁用 SNMP 事件 MIB	如何禁用 SNMP 事件 MIB [224]
更改 SNMP 事件 MIB	如何更改 SNMP 事件 MIB [225]
将 SNMP 主机添加到要接收 MIB 陷阱通知的主机列表中	如何在节点上使 SNMP 主机能够接收 SNMP 陷阱 [226]
删除 SNMP 主机	如何在节点上禁止 SNMP 主机接收 SNMP 陷阱 [227]
添加 SNMP 用户	如何在节点上添加 SNMP 用户 [228]
删除 SNMP 用户	如何从节点中删除 SNMP 用户 [228]

▼ 如何启用 SNMP 事件 MIB

此过程说明了如何启用 SNMP 事件 MIB。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

- 1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
- 2. 启用 SNMP 事件 MIB。

```
phys-schost-1# clsnmpmib enable [-n node] MIB
```

`[-n node]` 指定要启用的事件 MIB 所在的 *node*。您可以指定一个节点 ID 或节点名称。如果不指定此选项，默认情况下将使用当前节点。

MIB 指定要启用的 MIB 的名称。在本例中，MIB 的名称必须是 `event`。

▼ 如何禁用 SNMP 事件 MIB

此过程说明了如何禁用 SNMP 事件 MIB。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 禁用 SNMP 事件 MIB。

```
phys-schost-1# clsnmpmib disable -n node MIB
```

`-n node` 指定要禁用的事件 MIB 所在的 *node*。您可以指定一个节点 ID 或节点名称。如果不指定此选项，默认情况下将使用当前节点。

MIB 指定要禁用的 MIB 的类型。在本例中，必须指定 `event`。

▼ 如何更改 SNMP 事件 MIB

此过程说明如何更改 SNMP 事件 MIB 的协议、最低严重性值以及事件日志记录。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 更改 SNMP 事件 MIB 的协议、最低严重性级别值和事件日志记录。

```
phys-schost-1# clsnmpmib set -n node
-p version=SNMPv3 \
-p min_severity=WARNING \
-p log_number=100 MIB
```

`-n node`

指定要更改的事件 MIB 所在的 *node*。您可以指定一个节点 ID 或节点名称。如果不指定此选项，默认情况下将使用当前节点。

`-p version=value`

指定要用于 MIB 的 SNMP 协议版本。您可以按如下方式指定 *value*：

- `version=SNMPv2`
- `version=snmpv2`
- `version=2`
- `version=SNMPv3`

- `version=snmpv3`
- `version=3`

`-p min_severity=value`

指定要用于 MIB 的最低严重性级别值。您可以按如下方式指定 *value* :

- `min_severity=NOTICE`
- `min_severity=WARNING`
- `min_severity=ERROR`
- `min_severity=CRITICAL`
- `min_severity=FATAL`

`-p log_number=number`

指定在停止使用较旧的条目之前要在 MIB 表中记录的事件的数量。默认值为 100。取值范围为 100 到 500。您可以按如下方式指定 *value* : `log_number=100`。

MIB

指定要应用子命令的一个或多个 MIB 的名称。在本例中，必须指定 `event`。如果不指定该操作数，子命令将使用表示所有 MIB 的默认加号 (+)。如果使用 *MIB* 操作数，请在所有其他命令行选项之后在空格分隔的列表中指定 MIB。

有关更多信息，请参见 [clsnmpmib\(1CL\)](#) 手册页。

▼ 如何在节点上使 SNMP 主机能够接收 SNMP 陷阱

此过程说明如何将某个节点上的 SNMP 主机添加到将接收 MIB 陷阱通知的主机列表中。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 将主机添加到另一节点上某个团体的 SNMP 主机列表中。

```
phys-schost-1# clsnmphost add -c SNMPcommunity [-n node] host
```

`-c SNMPcommunity`

指定与主机名配合使用的 SNMP 团体名称。主机是网络中可以配置为接收陷阱的系统

如果将主机添加到除 `public` 以外的团体，必须指定 SNMP 团体名称 `SNMPcommunity`。如果使用不带 `c` 选项的 `-add` 子命令，该子命令会使用 `public` 作为默认团体名称。

如果指定的团体名称不存在，此命令将创建该团体。

`-n node`

指定用于访问群集中的 SNMP MIB 的 SNMP 主机所在群集节点的名称。您可以指定一个节点名称或节点 ID。如果不指定此选项，则默认为运行命令的节点。

`host`

指定供访问群集中的 SNMP MIB 的主机的名称、IP 地址或 IPv6 地址。这可以是群集外部的主机，也可以是尝试获取 SNMP 陷阱的群集节点本身。

▼ 如何在节点上禁止 SNMP 主机接收 SNMP 陷阱

此过程说明如何在节点上将一个 SNMP 主机从要接收 MIB 陷阱通知的主机列表中删除。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 在指定节点上将主机从某个团体的 SNMP 主机列表中删除。

```
phys-schost-1# clnmphost remove -c SNMPcommunity -n node host
```

`remove`

将指定 SNMP 主机从指定节点上删除。

`-c SNMPcommunity`

指定要从中删除 SNMP 主机的 SNMP 团体的名称。

`-n node`

指定要从配置中删除的 SNMP 主机所在群集节点的名称。您可以指定一个节点名称或节点 ID。如果不指定此选项，则默认为运行命令的节点。

`host`

指定要从配置中删除的主机的名称、IP 地址或 IPv6 地址。这可以是群集外部的主机，也可以是尝试获取 SNMP 陷阱的群集节点本身。

要删除指定 SNMP 团体中的所有主机，请使用加号 (+) 代替 `host`，并使用 `-c` 选项。

要删除所有主机，请使用加号 (+) 代替 `host`。

▼ 如何在节点上添加 SNMP 用户

此过程说明如何向节点上的 SNMP 用户配置中添加 SNMP 用户。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 添加 SNMP 用户。

```
phys-schost-1# clsnmpuser create -n node -a authentication -f password user
```

`-n node`

指定要在其中添加 SNMP 用户的节点。您可以指定一个节点 ID 或节点名称。如果不指定此选项，默认情况下将使用当前节点。

`-a authentication`

指定用于对用户进行授权的验证协议。验证协议的值可以是 SHA 或 MD5。

`-f password`

指定包含 SNMP 用户密码的文件。如果在创建新用户时未指定该选项，则此命令会提示您输入一个密码。此选项仅对 `add` 子命令有效。

必须按以下格式指定用户密码（每个密码占一行）：

```
user:password
```

密码不能包含以下字符，也不能包含空格：

- ；（分号）
- ：（冒号）
- \（反斜杠）
- \n（新行）

`user`

指定要添加的 SNMP 用户的名称。

▼ 如何从节点中删除 SNMP 用户

此过程说明如何从节点上的 SNMP 用户配置中删除 SNMP 用户。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
2. 删除 SNMP 用户。

```
phys-schost-1# clsnmpuser delete -n node user
```

`-n node` 指定要从中删除 SNMP 用户的节点。您可以指定一个节点 ID 或节点名称。如果不指定此选项，默认情况下将使用当前节点。

`user` 指定要删除的 SNMP 用户的名称。

配置负载限制

可以通过设置负载限制来启用资源组负载在各节点间的自动分配。可为每个群集节点配置一组负载限制。将负载因子分配给资源组，负载因子与已定义的节点负载限制相对应。默认行为是跨资源组节点列表中的所有可用节点均匀分配资源组负载。

资源组由 RGM 从资源组节点列表中的某个节点启动，以便不会超出该节点的负载限制。在资源组由 RGM 分配给节点后，每个节点上资源组的负载因子将会汇总来提供总负载。然后总负载会与该节点的负载限制相比较。

负载限制包含以下项：

- 用户指定的名称。
- 软限制值 – 可以临时超出软负载限制。
- 硬限制值 – 绝不能超出硬负载限制，并且应严格执行此限制。

可以使用一个命令同时设置硬限制和软限制。如果没有明确设置其中某一限制，则会使用默认值。可使用 `clnode create-loadlimit`、`clnode set-loadlimit` 和 `clnode delete-loadlimit` 命令创建和修改每个节点的硬/软负载限制。有关更多信息，请参见 [clnode\(1CL\)](#) 手册页。

可以将某个资源组配置为具有较高优先级，以减小其从特定节点被替换的可能性。还可以设置 `preemption_mode` 属性，以确定某个资源组是否会由于节点过载而被优先级较高的资源组从节点中抢占。`concentrate_load` 属性还允许您将资源组负载集中分配给尽可能少的节点。默认情况下，`concentrate_load` 属性的默认值为 `FALSE`。

注 - 您可以在全局群集或区域群集中配置节点的负载限制。可以使用命令行、`clsetup` 实用程序或 Oracle Solaris Cluster Manager 浏览器界面配置负载限制。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。以下过程说明如何使用命令行配置负载限制。

▼ 如何在节点上配置负载限制

注 - 您还可以使用 Oracle Solaris Cluster Manager 浏览器界面在全局群集节点或区域群集节点上创建并配置负载限制，或者编辑或删除现有节点负载限制。单击 "Nodes"（节点）或 "Zone Clusters"（区域群集），然后单击节点名称以访问其页面。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在全局群集的任一节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。

2. 为要使用负载平衡的节点创建并设置负载限制。

在以下示例中，区域群集名称为 `zc1`。样例属性名为 `mem_load`，软限制值为 11，硬负载限制值为 20。软限制和硬限制是可选参数，如果您没有特别定义，则默认为无限制。有关更多信息，请参见 [clnode\(1CL\)](#) 手册页。

```
# clnode create-loadlimit -p limitname=mem_load -Z zc1 \
-p softlimit=11 -p hardlimit=20 node1 node2 node3
```

3. 为每个资源组指定负载因子值。

在以下示例命令中，在两个资源组 `rg1` 和 `rg2` 上设置负载因子。负载因子设置与节点的已定义负载限制相对应。

```
# clresourcegroup set -p load_factors=mem_load@50,factor2@1 rg1 rg2
```

还可以使用 `clresourceroup create` 命令在资源组创建期间执行此步骤。有关更多信息，请参见 [clresourcegroup\(1CL\)](#) 手册页。

4. 如果需要，执行一个或多个其他可选配置任务。

- 重新分配现有负载。

```
# clresourcegroup remaster rg1 rg2
```

此命令可将资源组从当前主节点移至其他节点，以实现均匀负载分配。

- 将某些资源组的优先级设置为高于其他资源组。

```
# clresourcegroup set -p priority=600 rg1
```

默认优先级为 500。在节点分配中，具有较高优先级值的资源组优先于具有较低优先级的资源组。

- 设置 `Preemption_mode` 属性。

```
# clresourcegroup set -p Preemption_mode=No_cost rg1
```

有关 HAS_COST、NO_COST 和 NEVER 选项的更多信息，请参见 [clresourcegroup\(1CL\)](#) 手册页。

- 设置 **Concentrate_Load** 标志。

```
# cluster set -p Concentrate_load=TRUE
```

- 指定资源组之间的关联。

正向或负向强关联性优先于负载分配。不得违反强关联性，也不得违反硬负载限制。如果同时设置了强关联和硬负载限制，那么假如无法满足这两项约束，可能会强制某些资源组保持脱机状态。

以下示例指定了区域群集 `zc1` 中的资源组 `rg1` 与区域群集 `zc2` 中的资源组 `rg2` 之间的正向强关联性。

```
# clresourcegroup set -p RG_affinities=++zc2:rg2 zc1:rg1
```

5. 检验群集中所有全局群集节点和区域群集节点的状态。

```
# clnode status -Z all -v
```

输出中包含在节点上定义的任何负载限制设置。

执行区域群集管理任务

您可以在区域群集中执行其他管理任务，例如，移动区域路径、准备区域群集以运行应用程序以及克隆区域群集。所有这些命令必须从全局群集的节点上执行。

您可以通过使用 `clsetup` 实用程序启动区域群集配置向导，创建新的区域群集或者将文件系统或存储设备添加到现有区域群集。区域群集中的区域是在运行 `clzonecluster install -c` 对配置文件进行配置时配置的。有关使用 `clsetup` 实用程序或 `-c config_profile` 选项的说明，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[创建和配置区域群集](#)”。

还可以使用 Oracle Solaris Cluster Manager 浏览器界面创建区域群集或者向区域群集中添加文件系统或存储设备。您还可以使用 Oracle Solaris Cluster Manager 浏览器界面编辑区域群集的 Resource Security（资源安全性）属性。单击 "Zone Clusters"（区域群集），单击区域群集的名称以转至其页面，然后单击 "Solaris Resources"（Solaris 资源）选项卡可以管理区域群集组件，或者单击 "Properties"（属性）可以管理区域群集属性。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

注 - 只能从全局群集的节点上运行的 Oracle Solaris Cluster 命令在用于区域群集时无效。有关命令在区域群集中的有效用法的信息，请参见相应的 Oracle Solaris Cluster 手册页。

表 19 其他区域群集任务

任务	说明
将区域路径移至新的区域路径。	<code>clzonecluster move -f zonepath zone-cluster-name</code>
准备区域群集以运行应用程序	<code>clzonecluster ready -n nodename zone-cluster-name</code>
从统一归档文件恢复节点	如何从统一归档文件恢复节点 [194]
从统一归档文件配置或安装区域群集	如何从统一归档文件配置区域群集 [232] 如何从统一归档文件安装区域群集 [233]
	使用命令： <code>clzonecluster clone -Z target- zone-cluster-name [-m copymethod] source-zone-cluster-name</code> 在使用 clone 子命令之前，请停止源区域群集。目标区域群集必须已经配置。
向区域群集添加网络地址	如何向区域群集添加网络地址 [234]
向区域群集添加节点	如何向现有的群集或区域群集添加节点 [192]
从区域群集中删除节点	如何从区域群集中删除节点 [198]
删除区域群集	如何删除区域群集 [235]
从区域群集中删除文件系统	如何从区域群集中删除文件系统 [236]
从区域群集中删除存储设备	如何从区域群集中删除存储设备 [239]
从统一归档文件恢复区域群集节点	如何从统一归档文件恢复节点 [194]
对节点卸载进行故障排除	“对节点卸载进行故障排除” [222]
创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB	“创建、设置和管理 Oracle Solaris Cluster SNMP 事件 MIB” [223]

▼ 如何从统一归档文件配置区域群集

可以使用 `clzonecluster` 命令启动交互式实用程序，从统一归档文件配置 `solaris10` 或 `labeled` 标记区域群集。`clzonecluster configure` 实用程序允许您指定恢复归档文件或克隆归档文件。

在从归档文件配置区域群集时，如果您喜欢使用命令行而非交互式实用程序，请使用 `clzonecluster configure -f command-file` 命令。有关更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

注 - 如果您要安装的区域群集已通过其他受支持的方法配置，则不必从统一归档文件配置该区域群集。

- 1. 创建恢复或克隆归档文件。

```
phys-schost# archiveadm create -r archive-location
```


使用 `create` 命令创建一个克隆归档文件，或者使用 `-r` 选项创建一个恢复归档文件。有关使用 `archiveadm` 命令的更多信息，请参见 [archiveadm\(1M\)](#) 手册页。

2. 在承载着区域群集的全局群集的某个节点上承担 `root` 角色。
3. 从统一归档文件中的恢复归档文件或克隆归档文件配置区域群集。

```
phys-schost-1# clzonecluster configure zone-cluster-name
```

`clzonecluster configure zone-cluster-name` 命令启动交互式实用程序，在其中您可以指定 `create -a archive [other-options-such-as "-x"]`。归档文件可以是克隆归档文件，也可以是恢复归档文件。

注 - 在创建区域群集之前，必须向配置添加区域群集成员。

`configure` 子命令使用 `zonecfg` 命令在指定的每台计算机上配置区域。`configure` 子命令允许您指定适用于区域群集的每个节点的属性。这些属性与 `zonecfg` 命令为各个区域建立的属性具有相同的意义。`configure` 子命令支持对 `zonecfg` 命令未知的属性进行配置。如果未指定 `-f` 选项，`configure` 子命令将启动交互式 `shell`。`-f` 选项将命令文件作为其参数。`configure` 子命令使用该文件以非交互方式创建或修改区域群集。

▼ 如何从统一归档文件安装区域群集

您可以从统一归档文件安装区域群集。`clzonecluster install` 实用程序允许您指定要用于安装的归档文件或 Oracle Solaris 10 映像归档文件的绝对路径。有关支持的归档类型的详细信息，请参见 [solaris10\(5\)](#) 手册页。该归档绝对路径应能在要安装区域群集的群集的所有物理节点上访问。统一归档文件安装可以使用恢复归档文件，也可以使用克隆归档文件。

在从归档文件安装区域群集时，如果您喜欢使用命令行而非交互式实用程序，请使用 `clzonecluster create -a archive -z archived-zone` 命令。有关更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

1. 创建恢复或克隆归档文件。

```
phys-schost# archiveadm create -r archive-location
```

使用 `create` 命令创建一个克隆归档文件，或者使用 `-r` 选项创建一个恢复归档文件。有关使用 `archiveadm` 命令的更多信息，请参见 [archiveadm\(1M\)](#) 手册页。

2. 在承载着区域群集的全局群集的某个节点上承担 `root` 角色。
3. 从统一归档文件中的恢复归档文件或克隆归档文件安装区域群集。

```
phys-schost-1# clzonecluster install -a absolute_path_to_archive zone-cluster-name
```

该归档绝对路径应能在要安装区域群集的群集的所有物理节点上访问。如果具有 HTTPS 统一归档文件位置，请使用 `-x cert|ca-cert|key=file` 指定 SSL 证书、证书颁发机构 (Certificate Authority, CA) 证书和密钥文件。

统一归档文件不包含区域群集节点资源。节点资源在配置群集时指定。当您使用统一归档文件从全局区域配置区域群集时，必须设置区域路径。

如果统一归档文件包含多个区域，请使用 `zone-cluster-name` 指定安装源的区域名称。有关更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

注 - 如果用于创建统一归档文件的源不包含 Oracle Solaris Cluster 软件包，则必须运行 `pkg install ha-cluster-packages`（替代为具体的软件包名称，例如 `ha-cluster-minimal` 或 `ha-cluster-framework-full`）。您需要引导区域，运行 `zlogin` 命令，然后运行 `pkg install` 命令。此操作将在目标区域群集上安装与全局群集相同的软件包。

4. 引导新的区域群集。

```
phys-schost-1# clzonecluster boot zone-cluster-name
```

▼ 如何向区域群集添加网络地址

执行此过程可添加网络地址以供现有区域群集使用。网络地址用于在区域群集中配置逻辑主机或共享 IP 地址资源。您可以多次运行 `clsetup` 实用程序，根据需要添加尽可能多的网络地址。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面向区域群集添加网络地址。单击 "Zone Clusters" (区域群集)，单击区域群集的名称以转至其页面，然后单击 "Solaris Resources" (Solaris 资源) 选项卡可以管理区域群集组件。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在托管区域群集的全局群集的某个节点上承担 `root` 角色。
2. 在全局群集上，配置要在区域群集中使用的群集文件系统。
启动 `clsetup` 实用程序。

```
phys-schost# clsetup
```

这时将显示主菜单。

3. 选择 "Zone Cluster" (区域群集) 菜单项。

4. 选择 "Add Network Address to a Zone Cluster" (向区域群集添加网络地址) 菜单项。
5. 选择要添加网络地址的区域群集。
6. 选择相应属性以指定要添加的网络地址。

`address=value`

指定用于在区域群集中配置逻辑主机或共享 IP 地址资源的网络地址。例如，192.168.100.101。

支持以下类型的网络地址：

- 一个有效的 IPv4 地址，可选择性地后跟 / 和一个前缀长度。
- 一个有效的 IPv6 地址，必须后跟 / 和一个前缀长度。
- 解析为一个 IPv4 地址的主机名。不支持解析为 IPv6 地址的主机名。

有关网络地址的更多信息，请参见 [zonecfg\(1M\)](#) 手册页。

7. 要添加其他网络地址，请键入 **a**。
8. 键入 **c** 以保存配置更改。

将会显示配置更改结果。例如：

```
>>> Result of Configuration Change to the Zone Cluster(sczone) <<<
```

```
Adding network address to the zone cluster...
```

```
The zone cluster is being created with the following configuration
```

```
/usr/cluster/bin/clzonecluster configure sczone
add net
set address=phys-schost-1
end
```

```
All network address added successfully to sczone.
```

9. 完成后，退出 **clsetup** 实用程序。

▼ 如何删除区域群集

您可以删除特定区域群集，或者使用通配符删除在全局群集中配置的所有区域群集。在删除区域群集之前，必须对其进行配置。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面删除区域群集。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在全局群集的节点上承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
从全局群集的节点执行此过程中的所有步骤。
2. 从区域群集中删除所有资源组及其资源。

```
phys-schost# clresourcegroup delete -F -Z zone-cluster-name +
```

注 - 此步骤是从全局群集节点中执行的。要改为从区域群集的某个节点执行此步骤，请登录到该区域群集节点，并在命令中省略 `-Z zonecluster`。

3. 停止区域群集。

```
phys-schost# clzonecluster halt zone-cluster-name
```

4. 卸载区域群集。

```
phys-schost# clzonecluster uninstall zone-cluster-name
```

5. 取消区域群集的配置。

```
phys-schost# clzonecluster delete zone-cluster-name
```

例 84 从全局群集中删除区域群集

```
phys-schost# clresourcegroup delete -F -Z SCZone +  
phys-schost# clzonecluster halt SCZone  
phys-schost# clzonecluster uninstall SCZone  
phys-schost# clzonecluster delete SCZone
```

▼ 如何从区域群集中删除文件系统

可以通过直接挂载或回送挂载将文件系统导出到区域群集。

区域群集支持以下文件系统的直接挂载：

- UFS 本地文件系统
- StorageTek QFS 独立文件系统
- StorageTek QFS 共享文件系统（当用于支持 Oracle RAC 时）
- Oracle Solaris ZFS（作为数据集导出的）
- 受支持 NAS 设备中的 NFS

区域群集可以管理以下文件系统的回送挂载：

- UFS 本地文件系统

- StorageTek QFS 独立文件系统
- StorageTek QFS 共享文件系统（仅当用于支持 Oracle RAC 时）
- UFS 群集文件系统

您可以配置 `HASStoragePlus` 或 `ScalMountPoint` 资源以管理文件系统的挂载。有关向区域群集中添加文件系统的说明，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[向区域群集中添加文件系统](#)”。

如果 ZFS 文件系统将其 `mountpoint` 属性设置为 `none` 或 `legacy`，或者将其 `canmount` 属性设置为 `off`，`HASStoragePlus` 资源将不监视 ZFS 文件系统。对于所有其他 ZFS 文件系统，`HASStoragePlus` 资源故障监视器将检查文件系统是否已挂载。如果文件系统已挂载，则 `HASStoragePlus` 资源将通过对文件系统执行读写操作来探测其可访问性，具体取决于 `I00option` 属性的值 `ReadOnly/ReadWrite`。

如果 ZFS 文件系统未挂载或者文件系统探测失败，则资源故障监视器操作失败，资源被设置为 `Faulted`。RGM 将尝试重新启动资源，具体取决于资源的 `retry_count` 和 `retry_interval` 属性。如果上述特定的 `mountpoint` 和 `canmount` 属性设置未起作用，该操作会导致重新挂载文件系统。如果故障监视器操作继续失败，并且超过了 `retry_interval` 中的 `retry_count`，则 RGM 将资源故障转移到其他节点。

`phys-schost#` 提示符表示全局群集提示符。此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面从区域群集中删除文件系统。单击“Zone Clusters”（区域群集），单击区域群集的名称以转至其页面，然后单击“Solaris Resources”（Solaris 资源）选项卡可以管理区域群集组件。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在托管区域群集的全局群集的某个节点上承担 **root** 角色。
此过程的某些步骤是从全局群集的一个节点上执行的。其他步骤是从区域群集中的一个节点上执行的。
2. 删除与要删除的文件系统相关的资源。
 - a. 识别并删除为要删除的区域群集文件系统配置的 Oracle Solaris Cluster 资源类型，例如 `HASStoragePlus` 和 `SUNW.ScalMountPoint`。

```
phys-schost# clresource delete -F -Z zone-cluster-name fs_zone_resources
```

- a. 如果适用，请识别并删除在全局群集中为要删除的文件系统配置的类型为 `SUNW.qfs` 的 Oracle Solaris Cluster 资源。

```
phys-schost# clresource delete -F fs_global_resources
```

请小心使用 -F 选项，因为它会强制删除您所指定的所有资源，即使您没有首先禁用这些资源也是如此。您指定的所有资源都将从其他资源的资源依赖性设置中删除，而这可能导致群集丢失服务。未删除的相关资源可能被置于无效状态或错误状态。有关更多信息，请参见 [clresource\(1CL\)](#) 手册页。

提示 - 如果删除的资源所属的资源组稍后变为空组，则您可以放心地删除该资源组。

3. 确定文件系统挂载点目录的路径。

例如：

```
phys-schost# clzonecluster configure zone-cluster-name
```

4. 从区域群集配置中删除文件系统。

```
phys-schost# clzonecluster configure zone-cluster-name
```

```
clzc:zone-cluster-name> remove fs dir=filesystemdirectory
```

```
clzc:zone-cluster-name> commit
```

文件系统挂载点由 **dir=** 指定。

5. 检验是否删除了该文件系统。

```
phys-schost# clzonecluster show -v zone-cluster-name
```

例 85 删除区域群集中的高可用性本地文件系统

此示例说明如何删除一个具有挂载点目录 (/local/ufs-1) 的文件系统，该文件系统是在名为 **sczone** 的区域群集中配置的。资源为 **hasp-rs**，其类型为 **HAStoragePlus**。

```
phys-schost# clzonecluster show -v sczone
...
Resource Name:                fs
dir:                          /local/ufs-1
special:                      /dev/md/ds1/dsk/d0
raw:                          /dev/md/ds1/rdisk/d0
type:                         ufs
options:                      [logging]
...
phys-schost# clresource delete -F -Z sczone hasp-rs
phys-schost# clzonecluster configure sczone
clzc:sczone> remove fs dir=/local/ufs-1
clzc:sczone> commit
phys-schost# clzonecluster show -v sczone
```

例 86 删除区域群集中的高可用性 ZFS 文件系统

此示例说明了如何删除名为 HAZpool 的 ZFS 池中的一个 ZFS 文件系统，该文件系统是在 sczone 区域群集中类型为 SUNW.HAStoragePlus 的资源 hasp-rs 中配置的。

```
phys-schost# clzonecluster show -v sczone
...
Resource Name:                dataset
name:                        HAZpool
...
phys-schost# clresource delete -F -Z sczone hasp-rs
phys-schost# clzonecluster configure sczone
clzc:sczone> remove dataset name=HAZpool
clzc:sczone> commit
phys-schost# clzonecluster show -v sczone
```

▼ 如何从区域群集中删除存储设备

您可以从区域群集中删除存储设备，例如 Solaris Volume Manager 磁盘集和 DID 设备。执行此过程可从区域群集中删除存储设备。

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面从区域群集中删除存储设备。单击 "Zone Clusters" (区域群集)，单击区域群集的名称以转至其页面，然后单击 "Solaris Resources" (Solaris 资源) 选项卡可以管理区域群集组件。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager \[266\]](#)。

1. 在托管区域群集的全局群集的某个节点上承担 **root** 角色。

此过程的某些步骤是从全局群集的一个节点上执行的。其他步骤可以从区域群集的一个节点上执行。

2. 删除与要删除的设备相关的资源。

找出并删除为要删除的区域群集设备配置的 Oracle Solaris Cluster 资源类型，例如 SUNW.HAStoragePlus 和 SUNW.ScalDeviceGroup。

```
phys-schost# clresource delete -F -Z zone-cluster dev_zone_resources
```

3. 确定要删除的设备的匹配项。

```
phys-schost# clzonecluster show -v zone-cluster
...
Resource Name:    device
match:            <device_match>
...
```

4. 从区域群集配置中删除设备。

```
phys-schost# clzonecluster configure zone-cluster
clzc:zone-cluster-name> remove device match=devices-match
clzc:zone-cluster-name> commit
clzc:zone-cluster-name> end
```

5. 重新引导区域群集。

```
phys-schost# clzonecluster reboot zone-cluster
```

6. 检验是否删除了设备。

```
phys-schost# clzonecluster show -v zone-cluster
```

例 87 从区域群集中删除 Solaris Volume Manager 磁盘集

此示例说明了如何删除在名为 `sczone` 的区域群集中配置的 Solaris Volume Manager 磁盘集 `apachedg`。`apachedg` 磁盘集的编号为 3。这些设备由群集中配置的 `zc_rs` 资源使用。

```
phys-schost# clzonecluster show -v sczone
...
Resource Name:      device
match:              /dev/md/apachedg/*dsk/*
Resource Name:      device
match:              /dev/md/shared/3/*dsk/*
...
phys-schost# clresource delete -F -Z sczone zc_rs

phys-schost# ls -l /dev/md/apachedg
lrwxrwxrwx 1 root root 8 Jul 22 23:11 /dev/md/apachedg -> shared/3
phys-schost# clzonecluster configure sczone
clzc:sczone> remove device match=/dev/md/apachedg/*dsk/*
clzc:sczone> remove device match=/dev/md/shared/3/*dsk/*
clzc:sczone> commit
clzc:sczone> end
phys-schost# clzonecluster reboot sczone
phys-schost# clzonecluster show -v sczone
```

例 88 从区域群集中删除 DID 设备

此示例说明如何删除在名为 `sczone` 的区域群集中配置的 DID 设备 `d10` 和 `d11`。这些设备由群集中配置的 `zc_rs` 资源使用。

```
phys-schost# clzonecluster show -v sczone
...
Resource Name:      device
match:              /dev/did/*dsk/d10*
Resource Name:      device
match:              /dev/did/*dsk/d11*
...
```



```

phys-schost# clresource delete -F -Z sczone zc_rs
phys-schost# clzonecluster configure sczone
clzc:sczone> remove device match=/dev/did/*dsk/d10*
clzc:sczone> remove device match=/dev/did/*dsk/d11*
clzc:sczone> commit
clzc:sczone> end
phys-schost#
phys-schost# clzonecluster show -v sczone

```

用于测试的故障排除过程

本节包含故障排除过程，您可以使用该过程进行测试。

在全局群集外部运行应用程序

▼ 如何从在非群集模式下引导的节点中获取 Solaris Volume Manager 元集

使用此过程可出于测试目的在全局群集外部运行一个应用程序。

1. 确定 Solaris Volume Manager 元集中是否使用了法定设备，并确定该法定设备是否使用了 SCSI2 或 SCSI3 预留空间。

```
phys-schost# clquorum show
```

- a. 如果法定设备位于 Solaris Volume Manager 元集中，请添加一个新的法定设备，此新的法定设备不属于稍后要在非群集模式下获取的元集的一部分。

```
phys-schost# clquorum add did
```

- b. 删除旧的法定设备。

```
phys-schost# clquorum remove did
```

- c. 如果法定设备使用 SCSI2 预留空间，请从旧的法定设备中清理 SCSI2 预留空间，并确认没有留下任何 SCSI2 预留空间。

以下命令可查找持久组保留仿真 (Persistent Group Reservation Emulation, PGRE) 密钥。如果磁盘上没有密钥，将显示 `errno=22` 消息。

```
# /usr/cluster/lib/sc/pgre -c pgre_inkeys -d /dev/did/rdisk/dids2
```

找到密钥后，将清理 PGRE 密钥。

```
# /usr/cluster/lib/sc/pgre -c pgre_scrub -d /dev/did/rdisk/dids2
```



注意 - 如果从磁盘中清理活动法定设备密钥，群集将在下次重新配置时出现紧急情况，并显示 `Lost operational quorum` 消息。

2. 清除要在非群集模式下引导的全局群集节点。

```
phys-schost# clresourcegroup evacuate -n target-node
```

3. 使任一或多个资源组脱机，这些资源组包含 `HAStorage` 或 `HAStoragePlus` 资源并包含受您稍后要在非群集模式下获取的元集影响的设备或文件系统。

```
phys-schost# clresourcegroup offline resource-group
```

4. 禁用已使其脱机的资源组中的所有资源。

```
phys-schost# clresource disable resource
```

5. 取消管理资源组。

```
phys-schost# clresourcegroup unmanage resource-group
```

6. 使对应的设备组脱机。

```
phys-schost# cldevicegroup offline device-group
```

7. 禁用设备组。

```
phys-schost# cldevicegroup disable device-group
```

8. 将被动节点引导至非群集模式。

```
phys-schost# shutdown -g0 -i0 -y
ok> boot -x
```

9. 在继续之前，检验被动节点上的引导过程是否已完成。

```
phys-schost# svcs -x
```

10. 确定元集中的磁盘上是否存在任何 SCSI3 预留空间。
在元集中的所有磁盘上运行以下命令。

```
phys-schost# /usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/dids2
```

11. 如果磁盘上存在任何 SCSI3 预留空间，请清理预留空间。

```
phys-schost# /usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/dids2
```

12. 获取已清除的节点上的元集。

```
phys-schost# metaset -s name -C take -f
```

13. 挂载包含元集上定义的设备文件系统。

```
phys-schost# mount device mountpoint
```

14. 启动应用程序并执行所需的测试。完成测试后，停止应用程序。

15. 重新引导节点并等待，直到引导过程结束。

```
phys-schost# reboot
```

16. 使设备组联机。

```
phys-schost# cldevicegroup online -e device-group
```

17. 启动资源组。

```
phys-schost# clresourcegroup online -eM resource-group
```

恢复损坏的磁盘集

如果磁盘集已损坏或者群集中的节点无法获得该磁盘集的所有权，则使用此过程。如果清除该状态的尝试失败，可使用此过程作为修复磁盘集的最后措施。

这些过程适用于 Solaris Volume Manager 元集和多属主 Solaris Volume Manager 元集。

▼ 如何保存 Solaris Volume Manager 软件配置

从头开始恢复磁盘集可能非常耗时，而且容易出错。一种更好的替代方法是使用 `metastat` 命令定期备份副本，或使用 Oracle Explorer (SUNWexplo) 创建备份。然后，可以使用保存的配置来重新创建磁盘集。应该将当前配置保存到文件中（使用 `prtvtoc` 和 `metastat` 命令），然后重新创建磁盘集及其组件。请参见[如何重新创建 Solaris Volume Manager 软件配置 \[244\]](#)。

1. 保存磁盘集中每个磁盘的分区表。

```
# /usr/sbin/prtvtoc /dev/global/rdisk/disk-name > /etc/lvm/disk-name.vtoc
```

2. 保存 Solaris Volume Manager 软件配置。

```
# /bin/cp /etc/lvm/md.tab /etc/lvm/md.tab_ORIGINAL
# /usr/sbin/metastat -p -s set-name >> /etc/lvm/md.tab
```

注 - 其他配置文件，例如 `/etc/vfstab` 文件，可能引用 Solaris Volume Manager 软件。此过程假设重新构建了完全等同的 Solaris Volume Manager 软件配置，因此挂载信息相同。如果 Oracle Explorer (SUNWexplo) 在拥有磁盘集的节点上运行，它会检索 `prtvtoc` 和 `metaset -p` 信息。

▼ 如何清除损坏的磁盘集

从一个或所有节点清除磁盘集会删除配置。要从节点中清除磁盘集，该节点不能具有该磁盘集的所有权。

1. 在所有节点上运行清除命令。

```
# /usr/sbin/metaset -s set-name -P
```

运行此命令会从数据库副本以及 Oracle Solaris Cluster 系统信息库中删除磁盘集信息。-P 和 -c 选项允许在无需完全重建 Solaris Volume Manager 环境的情况下清除磁盘集。

注 - 如果在节点引导至非群集模式的情况下清除多所有者磁盘集，则可能需要从 DCS 配置文件中删除相关信息。

```
# /usr/cluster/lib/sc/dcs_config -c remove -s set-name
```

有关更多信息，请参见 [dcs_config\(1M\)](#) 手册页。

2. 如果只想删除数据库副本中的磁盘集信息，请使用以下命令。

```
# /usr/sbin/metaset -s set-name -C purge
```

通常应使用 -P 选项而不是 -c 选项。使用 -c 选项会导致在重新创建磁盘集时出现问题，因为 Oracle Solaris Cluster 软件仍识别该磁盘集。

- a. 如果在 `metaset` 命令中使用 -c 选项，请首先创建磁盘集，查看是否发生问题。
- b. 如果存在问题，请从 `dcs` 配置文件中删除相关信息。

```
# /usr/cluster/lib/sc/dcs_config -c remove -s setname
```

如果 `purge` 选项失败，请确认您已经安装了最新的内核和元设备更新，然后联系 [My Oracle Support](#)。

▼ 如何重新创建 Solaris Volume Manager 软件配置

只有在发生完全丢失 Solaris Volume Manager 软件配置的情况下才应使用此过程。这些步骤假设您已经保存了当前的 Solaris Volume Manager 配置及其组件，并清除了损坏的磁盘集。

注 - 仅在双节点群集中使用中介。

1. 创建新磁盘集。

```
# /usr/sbin/metaset -s set-name -a -h node1 node2
```

如果这是多属主磁盘集，则使用以下命令创建新磁盘集。

```
/usr/sbin/metaset -s set-name -aM -h node1 node2
```

2. 在创建了磁盘集的同时主机上，根据需要添加中介主机（仅限双节点）。

```
/usr/sbin/metaset -s set-name -a -m node1 node2
```

3. 从该主机将相同磁盘添加回磁盘集。

```
/usr/sbin/metaset -s set-name -a /dev/did/rdisk/disk-name /dev/did/rdisk/disk-name
```

4. 如果清除了磁盘集并要重新创建该磁盘集，卷目录 (Volume Table of Contents, VTOC) 应该保留在磁盘上，因此，您可以跳过此步骤。

但是，如果正在重新创建要恢复的磁盘集，应根据 `/etc/lvm/disk-name.vtoc` 文件中保存的配置格式化磁盘。例如：

```
# /usr/sbin/fmthard -s /etc/lvm/d4.vtoc /dev/global/rdsk/d4s2
```

```
# /usr/sbin/fmthard -s /etc/lvm/d8.vtoc /dev/global/rdsk/d8s2
```

您可以在任何节点上运行此命令。

5. 在现有 `/etc/lvm/md.tab` 文件中检查每个元设备的语法。

```
# /usr/sbin/metainit -s set-name -n -a metadvice
```

6. 根据保存的配置创建每个元设备。

```
# /usr/sbin/metainit -s set-name -a metadvice
```

7. 如果元设备上存在文件系统，请运行 `fsck` 命令。

```
# /usr/sbin/fsck -n /dev/md/set-name/rdsk/metadvice
```

如果 `fsck` 命令仅显示几个错误，例如超级块计数，则表明设备很可能已经正确重建。然后您可以运行不带 `-n` 选项的 `fsck` 命令。如果出现多个错误，请检验您是否正确重建了元设备。如果是，则查看 `fsck` 错误，确定是否能够恢复文件系统。如果不能，您应该从备份中恢复数据。

8. 将所有群集节点上的所有其他元集串联到 `/etc/lvm/md.tab` 文件，然后串联本地磁盘集。

```
# /usr/sbin/metastat -p >> /etc/lvm/md.tab
```


对 CPU 使用控制的配置

如果要控制 CPU 使用情况，请对 CPU 控制工具进行配置。有关配置 CPU 控制工具的更多信息，请参见 [rg_properties\(5\)](#) 手册页。本章介绍了以下相关主题：

- “CPU 控制介绍” [247]
- “配置 CPU 控制” [248]

CPU 控制介绍

Oracle Solaris Cluster 软件可用于控制 CPU 的使用情况。

CPU 控制工具是在 Oracle Solaris OS 所提供功能的基础上构建的。有关区域、项目、资源池、处理器集和调度类的信息，请参见《[Oracle Solaris Zones 介绍](#)》。

在 Oracle Solaris OS 上，您可执行以下操作：

- 将 CPU 份额分配给资源组
- 将处理器分配给资源组

注 - 还可以使用 Oracle Solaris Cluster Manager 浏览器界面查看区域群集的配置。有关 Oracle Solaris Cluster Manager 登录说明，请参见[如何访问 Oracle Solaris Cluster Manager](#) [266]。

选择方案

根据您所选择的配置和操作系统版本的具体情况，CPU 控制级别会各不相同。本章中介绍的各个 CPU 控制方面均依赖一个前提条件，即资源组属性 `RG_SLM_TYPE` 设置为 `automated`。

公平份额调度器

给资源组分配 CPU 份额的过程的第一步是将系统的调度程序设置为公平份额调度器 (Fair Share Scheduler, FSS)。默认情况下, Oracle Solaris OS 的调度类是分时调度 (Timesharing Schedule, TS)。请将调度程序设置为 FSS 以使份额配置生效。

无论选择怎样的调度程序类, 您均可创建一个专用处理器集。

配置 CPU 控制

本节介绍如何控制全局群集节点中的 CPU 使用情况。

▼ 如何在全局群集节点中控制 CPU 使用情况

执行此过程可为将在全局群集节点中执行的资源组分配 CPU 份额。

如果某个资源组分配有 CPU 份额, 则当 Oracle Solaris Cluster 软件在全局群集节点中启动该资源组的资源时, 将执行以下任务:

- 根据指定的 CPU 份额数, 增加分配给该节点的 CPU 份额数 (*zone.cpu-shares*) (如果尚未这样做)。
- 在该节点中创建一个名为 *SCSLM_resource-group* 的项目 (如果尚未这样做)。此项目特定于该资源组, 并分配有指定数目的 CPU 份额 (*project.cpu-shares*)。
- 启动 *SCSLM_resource-group* 项目中的资源。

有关配置 CPU 控制工具的更多信息, 请参见 [rg_properties\(5\)](#) 手册页。

1. 将系统的默认调度程序设置为公平份额调度器 (Fair Share Scheduler, FSS)。

```
# dispadmin -d FSS
```

下次重新引导时, FSS 将成为默认调度程序。要使此配置立即生效, 请使用 [priocntl\(1\)](#) 命令。

```
# priocntl -s -C FSS
```

组合使用 [priocntl](#) 命令和 [dispadmin](#) 命令可确保 FSS 立即成为默认调度器, 并在重新引导后也保持不变。有关设置调度类的更多信息, 请参见 [dispadmin\(1M\)](#) 和 [priocntl\(1\)](#) 手册页。

注 - 如果 FSS 不是默认调度程序, 您分配的 CPU 份额将不会生效。

2. 在每个要使用 CPU 控制的节点上，配置全局群集节点的份额数以及默认处理器集中可用 CPU 的最小数目。

如果没有给 `globalzoneshares` 和 `defaultpsetmin` 属性赋值，这些属性将采用各自的默认值。

```
# clnode set [-p globalzoneshares=integer] [-p defaultpsetmin=integer] node
```

```
-p defaultpsetmin=integer
```

设置默认处理器集中可用的最小 CPU 数。缺省值为 1。

```
-p globalzoneshares=integer
```

设置分配给节点的份额数。缺省值为 1。

```
node
```

指定要设置其属性的节点。

设置这些属性即是设置该节点的属性。

3. 验证是否正确设置了这些属性。

```
# clnode show node
```

对于您指定的节点，`clnode` 命令可显示属性集以及为这些属性设置的值。如果未使用 `clnode` 设置 CPU 控制属性，这些属性将采用默认值。

4. 配置 CPU 控制工具。

```
# clresourcegroup create -p RG_SLM_TYPE=automated [-p RG_SLM_CPU_SHARES=value] resource-group
```

```
-pRG_SLM_TYPE=automated
```

允许您控制 CPU 使用情况，并自动执行一些步骤以在 Oracle Solaris OS 中配置系统资源管理。

```
-p RG_SLM_CPU_SHARES=value
```

指定分配给资源组特定项目的 CPU 份额数（即 `project.cpu-shares`），并确定分配给节点的 CPU 份额数（即 `zone.cpu-shares`）。

```
resource-group
```

指定资源组的名称。

在此过程中，未设置 `RG_SLM_PSET_TYPE` 属性。在节点中，该属性采用 `default` 值。

这步操作将创建一个资源组。此外，您还可使用 `clresourcegroup set` 命令修改现有资源组。

5. 激活配置更改。

```
# clresourcegroup online -eM resource-group
```

resource-group 指定资源组的名称。

注 - 请不要删除或修改 *SCSLM_resource-group* 项目。您可以手动将更多资源控制添加到项目中，例如，通过配置 *project.max-lwps* 属性。有关更多信息，请参见 [projmod\(1M\)](#) 手册页。

更新您的软件

本章在以下各节中提供了用于更新 Oracle Solaris Cluster 软件的信息和说明。

- [“Oracle Solaris Cluster 软件更新概述” \[251\]](#)
- [“更新 Oracle Solaris Cluster 软件” \[252\]](#)
- [“卸载软件包” \[257\]](#)

Oracle Solaris Cluster 软件更新概述

必须对所有群集成员节点应用相同的更新，群集才能正常运转。更新节点时，有时候可能需要先暂时取消该节点的群集成员身份或停止整个群集，然后才能执行更新。

可以采用两种方式更新 Oracle Solaris Cluster 软件。

- **升级** – 将群集升级到最新的主要或次要 Oracle Solaris Cluster 发行版，并通过更新所有软件包来更新 Oracle Solaris OS。主要发行版的示例：从 Oracle Solaris Cluster 4.0 升级到 5.0。次要发行版的示例：从 Oracle Solaris Cluster 4.1 升级到 4.2。

运行 `scinstall` 实用程序或 `scinstall -u update` 命令以创建新的引导环境（映像的可引导实例），将引导环境挂载到某个未使用的挂载点上，对位进行更新，然后激活新的引导环境。创建克隆环境起初不会占用额外的空间，因此是即刻完成的。在执行此更新后，必须重新引导群集。升级时，还会将 Oracle Solaris OS 升级到最新的兼容版本。有关详细说明，请参见 [《Oracle Solaris Cluster 4.3 Upgrade Guide》](#)。

如果您具有标记类型为 `solaris` 的故障转移区域，请按 [《Oracle Solaris Cluster 4.3 Upgrade Guide》](#) 中的 [“How to Upgrade a solaris Branded Failover Zone”](#) 中的说明进行操作。

如果您在区域群集中具有 `solaris10` 标记区域，请按照 [《Oracle Solaris Cluster 4.3 Upgrade Guide》](#) 中的 [“Upgrading a solaris10 Branded Zone in a Zone Cluster”](#) 中的升级说明进行操作。

要更新未标记的区域群集，请遵循[“更新特定软件包” \[253\]](#)中的过程更新底层全局群集。更新全局群集后，也会自动更新其区域群集。

注 - 应用 Oracle Solaris Cluster Core SRU 不会产生与将软件升级到其他 Oracle Solaris Cluster 发行版相同的结果。

- 更新 – 将特定 Oracle Solaris Cluster 软件包更新到不同的 SRU 级别。您可以使用某个 pkg 命令来更新支持系统信息库更新 (Support Repository Update, SRU) 中的映像包管理系统 (Image Packaging System, IPS) 软件包。
SRU 通常定期发布，包含更新的软件包和缺陷修复程序。系统信息库包含所有 IPS 软件包和更新的软件包。运行 pkg update 命令可同时将 Oracle Solaris 操作系统和 Oracle Solaris Cluster 软件更新到兼容的版本。执行此更新后，可能需要重新引导群集。有关说明，请参见[如何更新特定软件包 \[253\]](#)。
您必须是 My Oracle Support 的注册用户才能查看和下载 Oracle Solaris Cluster 产品所需的软件更新。如果您没有 My Oracle Support 帐户，请与您的 Oracle 服务代表或销售工程师联系，或者在 <http://support.oracle.com> 上进行联机注册。有关固件更新的信息，请参见您的硬件文档。

注 - 在应用或删除任何更新之前，请阅读软件更新自述文件。

有关使用 Oracle Solaris 软件包管理实用程序 pkg 的信息，请参见《[在 Oracle Solaris 11.3 中添加和更新软件](#)》中的第 3 章, “安装和更新软件包”。

更新 Oracle Solaris Cluster 软件

请参阅下表来确定如何升级或更新 Oracle Solaris Cluster 软件中的 Oracle Solaris Cluster 发行版或软件包。

表 20 更新 Oracle Solaris Cluster 软件

任务	说明
将整个群集升级到新的主要或次要发行版	《Oracle Solaris Cluster 4.3 Upgrade Guide》中的 “How to Upgrade the Software (Standard Upgrade)”
查看成功更新的提示	“更新技巧” [253]
更新特定软件包	如何更新特定软件包 [253]
更新法定服务器或 AI 安装服务器	如何更新法定服务器或 AI 安装服务器 [257]
更新 solaris 标记区域群集	如何更新 solaris 标记区域群集 [255]
修补 solaris10 标记区域群集	如何修补 solaris10 标记区域群集 (clzonecluster) [255] 如何修补区域群集中的 solaris10 标记区域 (patchadd) [256] 如何为 solaris10 标记区域群集中的区域修补缺备用引导环境 [256]

任务	说明
删除 Oracle Solaris Cluster 软件包	如何卸载软件包 [257] 如何卸载法定服务器或 AI 安装服务器软件包 [257]

更新技巧

使用以下技巧来更有效地管理 Oracle Solaris Cluster 更新：

- 在执行更新之前，阅读 SRU 的 README 文件。
- 检查您的存储设备的更新要求。
- 在生产环境中运行群集前应用所有的更新。
- 检查硬件固件级别，并安装所有可能需要的固件更新。有关固件更新的信息，请参阅您的硬件文档。
- 充当群集成员的所有节点必须具有相同的更新。
- 使群集子系统的更新保持最新。例如，这些更新包括：卷管理、存储设备固件和群集传输。
- 完成主要的更新后，测试故障转移。如果群集运转性能降低或弱化，则应做好取消该更新的准备。

更新特定软件包

Oracle Solaris 11 操作系统引入了 IPS 软件包。每个 IPS 软件包由一个故障管理资源标识符 (Fault Managed Resource Indicator, FMRI) 进行描述，您可以使用 `pkg` 命令来执行 SRU 更新。此外，您还可以使用 `scinstall -u` 命令来执行 SRU 更新。

您可能希望对特定软件包进行更新，以便使用更新的 Oracle Solaris Cluster 数据服务代理。

▼ 如何更新特定软件包

请在要更新的每个全局群集节点上执行此过程。群集节点上任何未标记的区域群集也将自动收到此更新。

1. 承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 更新软件包。

例如，要更新来自特定发布者的软件包，请在 `pkg-fmri` 中指定发布者名称。

```
# pkg update pkg-fmri
```



注意 - 如果您在使用 `pkg update` 命令时没有指定 *pkg-fmri*，则会更新有更新可用的所有已安装软件包。

如果已安装软件包有较新的版本可用且该版本与映像的其余部分兼容，则会将软件包更新到该版本。如果软件包包含 `reboot-needed` 标志设置为 `true` 的二进制代码，则执行 `pkg update pkg-fmri` 时会自动创建新的引导环境，并且在更新以后，您将引导至新引导环境。如果要更新的软件包未包含任何强制重新引导的二进制代码，则 `pkg update` 命令会更新实时映像，且不必重新引导。

3. 如果要更新的是数据服务代理 (`ha-cluster/data-service/*` 或 `ha-cluster/ha-service/gds` 的通用数据服务代理)，请运行以下命令。

```
# pkg change-facet facet.version-lock.pkg-name=false
# pkg update pkg-name
```

例如：

```
# pkg change-facet facet.version-lock.ha-cluster/data-service/weblogic=false
# pkg update ha-cluster/data-service/weblogic
```

注 - 如果您要冻结某个代理，防止对其进行更新，请运行以下命令。

```
# pkg change-facet facet.version-lock.pkg-name=false
# pkg freeze pkg-name
```

有关冻结特定代理的更多信息，请参见《在 Oracle Solaris 11.3 中添加和更新软件》中的“控制可选组件的安装”。

4. 验证软件包已更新。

```
# pkg verify -v pkg-fmri
```

更新或修补区域群集

要更新 `solaris` 标记区域群集，请使用 `scinstall-u update` 命令应用 SRU。

要修补 `solaris10` 标记区域群集，请使用 `clzonecluster install-cluster -p` 命令或 `patchadd` 命令或者通过修补备用引导环境来应用修补程序。

本节提供以下过程：

- [如何更新 solaris 标记区域群集 \[255\]](#)

- [如何修补 solaris10 标记区域群集 \(clzonecluster\) \[255\]](#)
- [如何修补区域群集中的 solaris10 标记区域 \(patchadd\) \[256\]](#)
- [如何为 solaris10 标记区域群集中的区域修补备用引导环境 \[256\]](#)

▼ 如何更新 solaris 标记区域群集

执行此过程可使用 `scinstall -u update` 命令应用 SRU 来更新 solaris 标记区域群集。

1. 在全局群集的节点上承担提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 从全局群集的某个节点，更新整个节点。

```
phys-schost# scinstall -u update [-b be-name]
```

针对每个群集节点重复此步骤。

3. 重新引导群集。

```
phys-schost# clzonecluster reboot
```

▼ 如何修补 solaris10 标记区域群集 (clzonecluster)

请在全局群集的一个节点上执行此过程。

1. 在全局群集的节点上承担提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 将区域群集引导为 Offline Running (脱机运行) 模式。

- 如果区域群集未引导，则运行以下命令。

```
phys-schost# clzonecluster boot -o zone-cluster
```

- 如果区域群集已引导，则运行以下命令。

```
phys-schost# clzonecluster reboot -o zone-cluster
```

3. 从全局群集的一个节点，修补整个 solaris10 标记区域群集。

```
phys-schost# clzonecluster install-cluster -p patch-spec [options] zone-cluster
```

有关 `install-cluster` 子命令的更多信息，请参见 [clzonecluster\(1CL\)](#) 手册页。

注 - 如果 `scinstall -u update` 命令失败并显示消息，指示 `pkg` 命令过时，请遵循该消息中的说明。

4. 重新引导区域群集。

```
phys-schost# clzonecluster reboot zone-cluster
```

▼ 如何修补区域群集中的 solaris10 标记区域 (patchadd)

对每个配置的区域群集节点执行以下步骤。

1. 使区域群集变为脱机运行状态。

```
phys-schost# clzonecluster reboot -o zone-cluster
```

2. 登录到区域群集节点。

```
phys-schost# zlogin zone-cluster
```

3. 在命令行中，键入 patchadd。

```
zchost# patchadd
```

4. 安装新发行版对应的修补程序。

▼ 如何为 solaris10 标记区域群集中的区域修补备用引导环境

使用此过程为区域群集节点的区域修补备用引导环境 (boot environment, BE)。使用此方法，可以在需要时退出修补的 BE。

1. 成为要修补的区域群集节点所在区域的管理员。
2. 创建、修补和激活新引导环境，然后使用 shutdown 命令重新引导区域。

请遵循《[创建和使用 Oracle Solaris 10 区域](#)》中的“[关于引导 solaris10 标记区域](#)”一节中的“如何在 solaris10 标记区域上创建和激活多个引导环境”中包含的这些任务的说明。

3. 对需要修补的任何其他区域群集节点所在区域，重复前面的步骤。

更新法定服务器或 AI 安装服务器

使用以下过程为法定服务器或 Oracle Solaris 11 自动化安装程序 (Automated Installer, AI) 安装服务器更新软件包。有关法定服务器的更多信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[如何安装和配置 Oracle Solaris Cluster 法定服务器软件](#)”。有关如何使用 AI 的更多信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的“[如何安装和配置 Oracle Solaris 和 Oracle Solaris Cluster 软件 \(IPS 系统信息库\)](#)”。

▼ 如何更新法定服务器或 AI 安装服务器

1. 承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 更新法定服务器或 AI 安装服务器软件包。

```
# pkg update ha-cluster/*
```

如果已安装的 `ha-cluster` 软件包有较新的版本可用且该版本与映像的其余部分兼容，则会将软件包更新到该版本。



注意 - 运行 `pkg update` 命令会更新系统上安装的所有 `ha-cluster` 软件包。

卸载软件包

您可以删除单个或多个软件包。

▼ 如何卸载软件包

1. 承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 卸载一个或多个现有软件包。

```
# pkg uninstall pkg-fmri [pkg-fmri ...]
```

如果有其他已安装的软件包依赖于您正在卸载的 `pkg-fmri`，则 `pkg uninstall` 命令将失败。要卸载 `pkg-fmri`，您必须将 `pkg uninstall` 命令提供给所有 `pkg-fmri` 相关项。有关卸载软件包的其他信息，请参见《在 Oracle Solaris 11.3 中添加和更新软件》和 `pkg(1)` 手册页。

▼ 如何卸载法定服务器或 AI 安装服务器软件包

1. 承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 卸载法定服务器或 AI 安装服务器软件包。

```
# pkg uninstall ha-cluster/*
```



注意 - 此命令可卸载系统上安装的所有 `ha-cluster` 软件包。

备份和恢复群集

本章包括以下几节：

- “备份群集” [259]
- “恢复群集文件” [261]
- “恢复群集节点” [194]

备份群集

在备份群集之前，请找到要备份的文件系统的名称，计算出您需要多少磁带来存储完整的备份内容，然后再备份 ZFS 根文件系统。

表 21 任务列表：备份群集文件

任务	说明
为镜像的文件系统或网状文件系统执行联机备份	如何为镜像执行联机备份 (Solaris Volume Manager) [259]
备份群集配置	如何备份群集配置 [261]
备份存储磁盘的磁盘分区配置	参见存储磁盘的相关文档

▼ 如何为镜像执行联机备份 (Solaris Volume Manager)

无需卸载镜像的 Solaris Volume Manager 卷或使整个镜像脱机即可备份该卷。您必须暂时使其中一个子镜像脱机（因而失去镜像），但备份完成后可立即使之联机并重新同步，这样就不必停止系统，也不用拒绝用户访问数据。通过使用镜像来执行联机备份，可创建活动文件系统的“快照”备份。

如果在某个程序将数据写入卷后又立即运行了 `lockfs` 命令，则可能会出现问題。要避免此故障，请暂时停止在此节点上运行的所有服务。另外，在执行此备份过程之前，请确保群集正在无故障运行。

`phys-schost#` 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在要备份的群集节点上，承担等效角色。
2. 使用 **metaset** 命令来确定哪一个节点具有已备份卷的所有权。

```
# metaset -s setname
```

-s *setname* 指定磁盘集名称。

有关更多信息，请参见 [metaset\(1M\)](#) 手册页。

3. 使用 **lockfs** 命令以及 -w 选项锁定文件系统，防止写入。

```
# lockfs -w mountpoint
```

有关更多信息，请参见 [lockfs\(1M\)](#) 手册页。

4. 使用 **metastat** 命令来确定子镜像的名称。

```
# metastat -s setname -p
```

-p 以类似于 md.tab 文件的格式显示状态。

有关更多信息，请参见 [metastat\(1M\)](#) 手册页。

5. 使用 **metadetach** 命令使一个子镜像从镜像脱机。

```
# metadetach -s setname mirror submirror
```

有关更多信息，请参见 [metadetach\(1M\)](#) 手册页。

注 - 将继续从其他子镜像进行读取。但是，对镜像进行第一次写操作后，脱机子镜像将立即不再同步。脱机子镜像重新联机后，这种不一致现象就会得到纠正。您无需运行 **fsck** 命令。

6. 通过结合使用 **lockfs** 命令和 -u 选项，解除对文件系统的锁定，允许写入操作继续执行。

```
# lockfs -u mountpoint
```

7. 执行文件系统检查。

```
# fsck /dev/md/diskset/rdsk/submirror
```

8. 将脱机子镜像备份到磁带或另一介质上。

注 - 对子镜像使用原始设备 (/rdsk) 名称，而不使用块设备 (/dsk) 名称。

9. 使用 `metattach` 命令将元设备或卷切换回联机状态。

```
# metattach -s setname mirror submirror
```

当元设备或卷处于联机状态后，将会自动与镜像重新同步。有关更多信息，请参见 [metattach\(1M\)](#) 手册页。

10. 使用 `metastat` 命令可检验该子镜像是否正在重新同步。

```
# metastat -s setname mirror
```

有关更多信息，请参见《在 Oracle Solaris 11.3 中管理 ZFS 文件系统》。

▼ 如何备份群集配置

为确保群集配置已归档并可轻松恢复，请定期对群集配置加以备份。Oracle Solaris Cluster 可将群集配置导出到一个可扩展标记语言(eXtensible Markup Language, XML) 文件。

1. 登录到群集的任一节点，承担可提供 `solaris.cluster.read` RBAC 授权的角色。
2. 将群集配置信息导出到一个文件中。

```
# cluster export -o configfile
```

configfile 群集命令正在将群集配置信息导出到的 XML 配置文件的名称。有关 XML 配置文件的信息，请参见 [clconfiguration\(5CL\)](#) 手册页。

3. 检验群集配置信息是否已成功导出至 XML 文件。

```
# pfedit configfile
```

恢复群集文件

您可以将 ZFS 根文件系统恢复到新磁盘上。

可以从统一归档文件恢复群集或节点，也可以恢复特定的文件或文件系统。

在开始恢复文件或文件系统前，您需要了解以下信息。

- 需要哪些磁带
- 正在其上恢复文件系统的原始设备名称

- 正在使用的磁带驱动器类型
- 磁带驱动器的设备名称（本地或远程）
- 所有故障磁盘的分区方案，由于分区和文件系统必须完全复制到替换磁盘上

表 22 任务列表：恢复群集文件

任务	说明
对于 Solaris Volume Manager，恢复 ZFS 根 (/) 文件系统	如何恢复 ZFS 根 (/) 文件系统 (Solaris Volume Manager) [262]

▼ 如何恢复 ZFS 根 (/) 文件系统 (Solaris Volume Manager)

使用此过程可将 ZFS 根 (/) 文件系统恢复到新磁盘（例如，在更换损坏的根磁盘之后）。不应引导正在恢复的节点。执行恢复过程之前，请确保群集正在无故障运行。支持 UFS，但用作根文件系统时除外。UFS 可以用在共享磁盘上的 Solaris Volume Manager 元集中的元设备上。

注 - 由于新磁盘的分区格式必须与故障磁盘的分区格式相同，所以在开始此过程之前，请先确定分区方案，然后再相应地重新创建文件系统。

phys-schost# 提示符表示全局群集提示符。此操作过程适用于全局群集。

此过程提供了 Oracle Solaris Cluster 命令的长格式。此外，大多数命令还有短形式。这些命令除了名称长短的不同以外，其功能都是相同的。

1. 在对附加了待恢复节点的磁盘集具有访问权限的群集节点上，承担可提供 `solaris.cluster.modify` RBAC 授权的角色。
使用除要恢复的节点以外的其他节点。
2. 将要恢复的节点的主机名从所有元集中删除。
从元集中除正要删除的节点以外的一个节点上运行此命令。由于要恢复的节点处于脱机状态，因此系统会显示 `RPC: Rpcbind failure - RPC: Timed out` 错误。忽略此错误并继续执行下一步。

```
# metaset -s setname -f -d -h nodelist
```

- s setname 指定磁盘集名称。
- f 从磁盘集中删除最后一个主机。
- d 从磁盘集删除。

`-h nodelist` 指定要从磁盘集删除的节点的名称。

3. 恢复 ZFS 根文件系统 (/)。

有关更多信息，请参见《在 Oracle Solaris 11.3 中管理 ZFS 文件系统》中的“替换 ZFS 根池中的磁盘”。

要恢复 ZFS 根池或根池快照，请遵循《在 Oracle Solaris 11.3 中管理 ZFS 文件系统》中的“替换 ZFS 根池中的磁盘”中的过程。

注 - 请确保创建 `/global/.devices/node@nodeid` 文件系统。

如果 `/globaldevices` 备份文件存在于备份目录中，则它会与 ZFS 根一起恢复。该文件不是由 `globaldevices` SMF 服务自动创建的。

4. 在多用户模式下重新引导节点。

`# reboot`

5. 替换设备 ID。

`# cldevice repair root-disk`

6. 使用 `metadb` 命令可重新创建状态数据库副本。

`# metadb -c copies -af raw-disk-device`

`-c copies` 指定要创建的副本数。

`-f raw-disk-device` 要在其上创建副本的原始磁盘设备。

`-a` 添加副本。

有关更多信息，请参见 `metadb(1M)` 手册页。

7. 从一个群集节点（非恢复的节点）上将恢复后的节点添加到所有磁盘集。

`phys-schost-2# metaset -s setname -a -h nodelist`

`-a` 创建主机并将其添加到磁盘集中。

已将节点重新引导到群集模式下。群集已经就绪。

例 89 恢复 ZFS 根 (/) 文件系统 (Solaris Volume Manager)

下面的示例显示了将根 (/) 文件系统恢复到节点 `phys-schost-1` 的过程。从群集中的另一节点 `phys-schost-2` 运行 `metaset` 命令，以便从磁盘集 `schost-1` 中删除节点 `phys-`

schost-1，然后再将其重新添加到该磁盘集中。所有其他命令都是从 phys-schost-1 运行的。系统在 /dev/rdisk/c0t0d0s0 上创建了一个新的引导块，在 /dev/rdisk/c0t0d0s4 上重新创建了三个状态数据库副本。有关恢复数据的更多信息，请参见《在 Oracle Solaris 11.3 中管理 ZFS 文件系统》中的“解决 ZFS 存储池中的数据问题”。

从元集中删除节点

```
phys-schost-2# metaset -s schost-1 -f -d -h phys-schost-1
```

替换发生故障的磁盘并引导节点

使用 Oracle Solaris 文档中的过程恢复根 (/) 和 /usr 文件系统

重新引导该节点

```
# reboot
```

替换磁盘 ID

```
# cldevice repair /dev/dsk/c0t0d0
```

重新创建状态数据库副本

```
# metadb -c 3 -af /dev/rdisk/c0t0d0s4
```

将节点添加回元集

```
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```


使用 Oracle Solaris Cluster Manager 浏览器界面

本章提供了有关 Oracle Solaris Cluster Manager 浏览器界面的说明，您可以使用该界面管理群集的多个方面。此外，本章还介绍了如何访问和使用 Oracle Solaris Cluster Manager。

注 - Oracle Solaris Cluster Manager 使用 Oracle GlassFish Server 软件的专用版本，该版本随 Oracle Solaris Cluster 产品提供。请勿尝试安装 Oracle GlassFish Server 软件的公用版本，或更新到其任何修补程序集。否则可能会在更新 Oracle Solaris Cluster 软件或安装 Oracle Solaris Cluster SRU 时导致软件包问题。Oracle Solaris Cluster Manager 所需的 Oracle GlassFish Server 专用版本的所有错误修补都在 Oracle Solaris Cluster SRU 中提供。

这一章包括以下主题：

- [“Oracle Solaris Cluster Manager 概述” \[265\]](#)
- [“访问 Oracle Solaris Cluster Manager 软件” \[266\]](#)
- [“为 Oracle Solaris Cluster Manager 配置辅助功能支持” \[269\]](#)
- [“使用拓扑监视群集” \[269\]](#)

Oracle Solaris Cluster Manager 概述

借助 Oracle Solaris Cluster Manager 浏览器界面，可通过图形方式显示群集信息、检查群集组件状态以及监视配置更改。Oracle Solaris Cluster Manager 还可用于执行许多管理任务，包括管理以下 Oracle Solaris Cluster 组件：

- 数据服务
- 区域群集
- 节点
- 专用适配器、电缆和交换机
- 法定设备

- 设备组
- 磁盘
- NAS 设备
- 节点负载限制
- 资源组和资源
- Geographic Edition 伙伴关系

Oracle Solaris Cluster Manager 当前还不能执行所有 Oracle Solaris Cluster 管理任务。您必须使用命令行界面来执行其他操作。

访问 Oracle Solaris Cluster Manager 软件

Oracle Solaris Cluster Manager 浏览器界面提供了在 Oracle Solaris Cluster 软件中轻松管理许多任务的方式。有关更多信息，请参见 Oracle Solaris Cluster Manager 联机帮助。

引导群集时，Common Agent Container 会自动启动。如果需要验证 Common Agent Container 是否正在运行，请参见[“排除 Oracle Solaris Cluster Manager 的故障” \[267\]](#)。

提示 - 请勿单击浏览器中的 "Back"（后退）从 Oracle Solaris Cluster Manager 中退出。

▼ 如何访问 Oracle Solaris Cluster Manager

此过程说明了如何访问群集上的 Oracle Solaris Cluster Manager。

1. 在群集成员上，承担 root 角色。
2. 从管理控制台或群集外部的任何其他计算机上，启动一个浏览器。
3. 确保浏览器的磁盘和内存高速缓存大小设置为大于 0 的值。
4. 确认在浏览器中启用了 Java 和 Javascript。
5. 通过浏览器连接到群集中一个节点上的 Oracle Solaris Cluster Manager 端口。
默认端口号为 8998。

`https://node:8998/scm`

6. 接受 Web 浏览器提供的所有证书。

此时将显示 Oracle Solaris Cluster Manager 登录页。

7. 输入要管理的群集中的节点的名称，或者接受默认值 localhost 来管理当前群集。
8. 输入节点的用户名和密码。
9. 单击 "Sign In" (登录)。

此时将显示 Oracle Solaris Cluster Manager 应用程序启动页。

注 - 如果您配置了多个群集，则可以从下拉列表中选择 "Other" (其他) 并登录到另一个群集来显示该群集的信息。如果群集是一个或多个伙伴关系的成员，则在您访问 Partnerships 文件夹之后，所有伙伴名称都会自动添加到该下拉列表。在您执行身份验证之后，可以选择 "Switch Cluster" (切换群集)。

如果无法连接到 Oracle Solaris Cluster Manager，请参见[“排除 Oracle Solaris Cluster Manager 的故障” \[267\]](#)。如果在 Oracle Solaris 安装期间选择了受限的网络配置文件，则对 Oracle Solaris Cluster Manager 的外部访问会受到限制。使用 Oracle Solaris Cluster Manager 浏览器界面需要此网络。

排除 Oracle Solaris Cluster Manager 的故障

- 验证两个管理器服务是否正在运行。

```
# svcs system/cluster/manager\*
```

STATE	STIME	FMRI
online	Oct_30	svc:/system/cluster/manager-glassfish3:default
online	Oct_30	svc:/system/cluster/manager:default

使用 `svcadm` 命令禁用或启用 `system/cluster/manager-glassfish3`。此操作会分别停止并重新启动应用服务器。您应让 `system/cluster/manager` 保持联机。您不需要将其禁用或启用。

- 如果无法连接到 Oracle Solaris Cluster Manager，请通过输入 `usr/sbin/cacoadm status` 来确定 Common Agent Container 是否正在运行。如果 Common Agent Container 没有运行，则将看到登录页面，但您无法进行身份验证。

您可以通过输入 `/usr/sbin/cacoadm start` 手动启动 Common Agent Container。

▼ 如何配置 Common Agent Container 安全密钥

Oracle Solaris Cluster Manager 使用强加密技术来确保 Oracle Solaris Cluster Manager Web 服务器和各群集节点之间通信的安全。

当您在 Oracle Solaris Cluster Manager 中使用数据服务配置向导或执行其他 Oracle Solaris Cluster Manager 任务时，可能会发生 Cacao 连接错误。此过程将 Common Agent Container 的安全性文件复制到所有群集节点。这可以确保 Common Agent Container 的安全性文件在所有群集节点上完全相同，并且复制的文件保留有正确的文件权限。执行此过程会同步安全密钥。

1. 在每个节点上，停止安全性文件代理。

```
phys-schost# /usr/sbin/cacaoadm stop
```

2. 在一个节点上切换至 `/etc/cacao/instances/default/` 目录。

```
phys-schost-1# cd /etc/cacao/instances/default/
```

3. 创建 `/etc/cacao/instances/default/` 目录的 tar 文件。

```
phys-schost-1# tar cf /tmp/SECURITY.tar security
```

4. 将 `/tmp/Security.tar` 文件复制到每个群集节点。

5. 在复制 `/tmp/SECURITY.tar` 文件的每个节点上，提取安全性文件。

`/etc/cacao/instances/default/` 目录中已存在的所有安全性文件都将被覆盖。

```
phys-schost-2# cd /etc/cacao/instances/default/
```

```
phys-schost-2# tar xf /tmp/SECURITY.tar
```

6. 删除 tar 文件的所有副本以避免安全隐患。
必须删除 tar 文件的所有副本以避免安全隐患。

```
phys-schost-1# rm /tmp/SECURITY.tar
```

```
phys-schost-2# rm /tmp/SECURITY.tar
```

7. 在每个节点上，启动安全性文件代理。

```
phys-schost# /usr/sbin/cacaoadm start
```

▼ 如何检查网络绑定地址

在尝试查看有关不运行 Oracle Solaris Cluster Manager 的节点的信息时，如果收到系统错误消息，请检查 Common Agent Container 的 `network-bind-address` 参数是否已设置为正确值 `0.0.0.0`。

在群集的每个节点上执行以下步骤。

1. 确定网络绑定地址。

```
phys-schost# cacaoadm list-params | grep network
```

```
network-bind-address=0.0.0.0
```

如果网络绑定地址设置为 0.0.0.0 以外的任何其他值，则您需要将其更改为所需的地址。

2. 更改前请停止 cacao，更改后再将其启动。

```
phys-schost# cacaoadm stop
phys-schost# cacaoadm set-param network-bind-address=0.0.0.0
phys-schost# cacaoadm start
```

为 Oracle Solaris Cluster Manager 配置辅助功能支持

Oracle Solaris Cluster Manager 提供了以下辅助功能设置：

- 屏幕阅读器
- 高对比度
- 大字体

要设置这些控件中的一个或多个，请单击 Oracle Solaris Cluster Manager 顶部菜单栏中的 "Accessibility"（辅助功能）。然后单击要启用或禁用的辅助功能控制的复选框。

对辅助功能控制设置的更改在登录会话之外不会持续存在。如果您在同一登录会话期间对其他群集节点进行了验证，该设置会持续存在。

使用拓扑监视群集

"Topology"（拓扑）视图可帮助您监视群集和发现问题。您可以快速查看对象之间的关系，了解哪些资源组和资源属于每个节点。

▼ 如何使用拓扑监视和更新群集

要访问 "Topology"（拓扑）页面，请登录到 Oracle Solaris Cluster Manager，单击 "Resource Groups"（资源组），然后单击 "Topology"（拓扑）选项卡。各行表示依赖性和同位关系。

联机帮助提供了有关各视图元素的详细说明，以及如何选择对象来筛选视图，并且单击右键可查看选定对象的上下文操作菜单。通过单击联机帮助旁边的箭头，可以折叠或恢复联机帮助。您也可以折叠或恢复筛选器。

下表提供了 "Resource Topology" (资源拓扑) 页面上的控件列表。

控件	功能	描述
	缩放	放大或缩小部分页面。
	概览	将取景框放到图上可平移视图。
	隔离	单击某个资源组或资源可从显示中删除其他所有对象。
	钻取	单击某个资源组可向下钻取其资源。
	重置	在隔离或钻取后恢复为完全视图。
	过滤	通过按类型、实例或状态选择对象来缩减显示内容。

以下过程说明了如何监视群集节点中的严重错误：

1. 在 "Topology" (拓扑) 选项卡中，找到 "Potential Masters" (潜在主节点) 区域。
2. 放大以查看群集中每个节点的状态。
3. 找到有红色的 "Critical" (严重) 状态图标的节点，右键单击该节点并选择 "Show Details" (显示详细信息)。
4. 在该节点的 "Status" (状态) 页中，单击 *System Log* (系统日志) 以查看和过滤日志消息。

部署示例：使用 Availability Suite 软件在群集间配置基于主机的数据复制

本附录提供了一种方法，可以在不使用 Oracle Solaris Cluster Geographic Edition 的情况下使用基于主机的复制。不过，使用 Oracle Solaris Cluster Geographic Edition 进行基于主机的复制，可简化群集之间基于主机的复制的配置和操作。请参见[“了解数据复制” \[93\]](#)。

本附录中的示例说明了如何使用 Oracle Solaris 的 Availability Suite 功能软件配置群集间基于主机的数据复制。该示例描绘了一个 NFS 应用程序（详细介绍如何执行各项具体任务）的完整群集配置。所有任务都应该在全局群集中执行。该示例不包括其他应用程序或群集配置所需的所有步骤。

如果使用基于角色的访问控制 (role-based access control, RBAC) 来访问群集节点，请确保承担可对所有 Oracle Solaris Cluster 命令提供授权的 RBAC 角色。需要具备以下 Oracle Solaris Cluster RBAC 授权才能执行这一系列的数据复制过程：

- `solaris.cluster.modify`
- `solaris.cluster.admin`
- `solaris.cluster.read`

有关使用 RBAC 角色的更多信息，请参见[《在 Oracle Solaris 11.3 中确保用户和进程的安全》](#)。有关每个 Oracle Solaris Cluster 子命令所需 RBAC 授权的信息，请参见 Oracle Solaris Cluster 手册页。

理解 Availability Suite 软件在群集中的应用

本节介绍了容灾性，并且描述了 Availability Suite 软件使用的数据复制方法。

容灾性是指当主群集发生故障时，在备用群集上恢复应用程序的能力。容灾性基于数据复制和接管。接管是指通过让一个或多个资源组和设备组联机，将某项应用程序服务重定位到辅助群集。

如果数据是在主群集和辅助群集之间同步复制的，则当主站点发生故障时，不会丢失已提交的数据。但是，如果数据是以异步方式复制的，则在主站点发生故障之前，某些数据可能尚未复制到辅助群集中，因此会丢失。

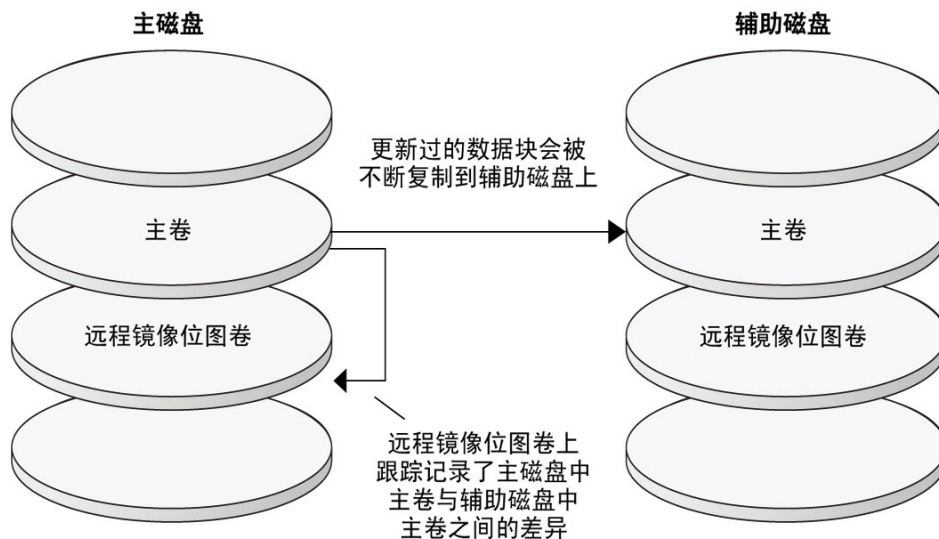
Availability Suite 软件使用的数据复制方法

本节描述了 Availability Suite 软件使用的远程镜像复制方法和时间点快照方法。该软件使用 `sndradm` 和 `iiadm` 命令复制数据。有关更多信息，请参见 [sndradm\(1M\)](#) 和 [iiadm\(1M\)](#) 手册页。

远程镜像复制

图 2 显示了远程镜像复制。通过 TCP/IP 连接可以将主磁盘主卷中的数据复制到辅助磁盘的主卷中。该软件使用远程镜像位图来跟踪主磁盘上的主卷与辅助磁盘上的主卷之间的差别。

图 2 远程镜像复制



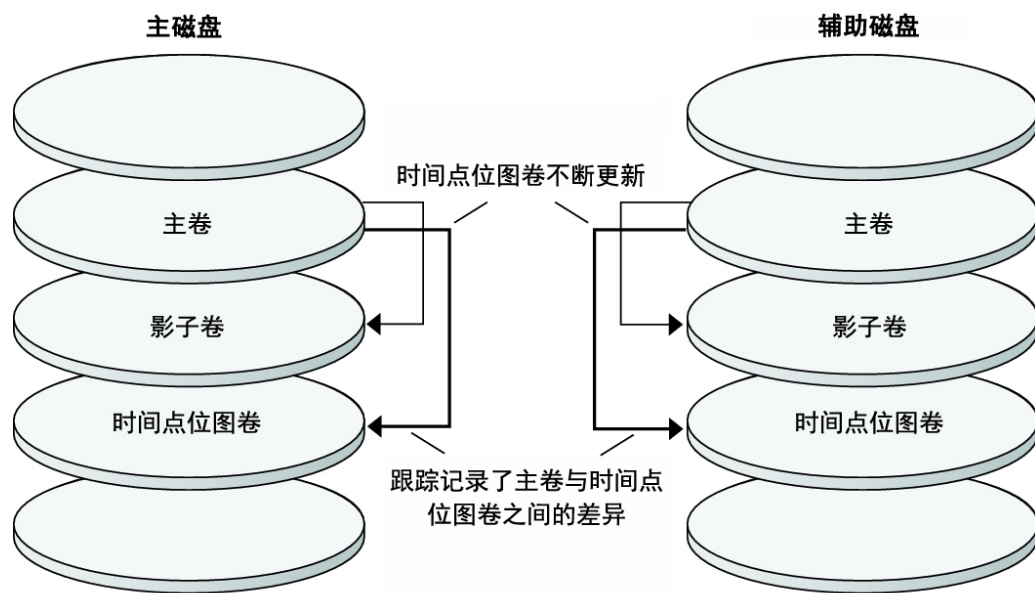
可以实时同步或异步执行远程镜像复制。可以为同步复制或异步复制单独配置每个群集中的每个卷。

- 在同步数据复制中，只有在更新了远程卷之后才能确认写入操作是否完成。
- 在异步数据复制中，在更新远程卷之前确认写入操作是否完成。异步数据复制以其长距离、低带宽而提供了更大的灵活性。

时间点快照

图 3 显示了时间点快照。每个磁盘的主卷中的数据都被复制到同一磁盘上的影子卷中。时间点位图跟踪主卷和影子卷之间的差异。数据被复制到影子卷之后，时间点位图将被复位。

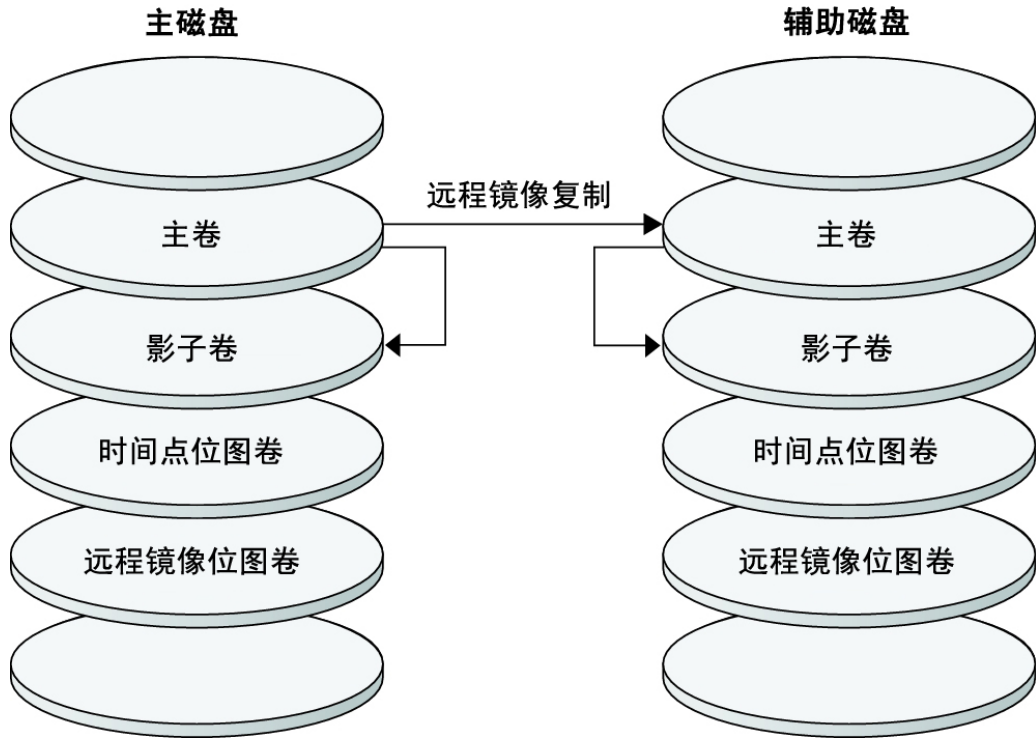
图 3 时间点快照



复制示例配置

图 4 展示了此配置示例是如何使用远程镜像复制和时间点快照的。

图 4 复制示例配置



在群集间配置基于主机的数据复制的准则

本节提供了有关配置群集间数据复制的准则。本节还包含了配置复制资源组和应用程序资源组的提示。为群集配置数据复制时，请使用这些指导信息。

本节包括以下主题：

- “配置复制资源组” [275]
- “配置应用程序资源组” [275]
 - “为故障转移应用程序配置资源组” [276]
 - “为可伸缩应用程序配置资源组” [277]
- “管理接管的准则” [278]

配置复制资源组

复制资源组可将设备组与逻辑主机名资源并置在 Availability Suite 软件控制下。逻辑主机名必须存在于每个数据复制流的末尾，且必须位于可充当设备主 I/O 路径的同一群集节点上。复制资源组必须具有以下特征：

- 是一个故障转移资源组
 - 每次只能在一个节点上运行故障转移资源。发生故障转移时，故障转移资源参与故障转移。
- 具有逻辑主机名资源
 - 逻辑主机名承载在每个群集（主群集和辅助群集）的一个节点上，用于提供 Availability Suite 软件数据复制流的源地址和目标地址。
- 具有 HASToragePlus 资源
 - 复制资源组被切换或故障转移之后，HASToragePlus 资源将强制执行设备组的故障转移。设备组被切换之后，Oracle Solaris Cluster 软件还将强制执行复制资源组的故障转移。这样，复制资源组和设备组将始终由同一节点配置或控制。
 - 必须在 HASToragePlus 资源中定义以下扩展属性：
 - GlobalDevicePaths。此扩展属性定义了卷所属的设备组。
 - AffinityOn 属性 = True。在复制资源组切换或故障转移时，此扩展特性使设备组进行切换或故障转移。该特性称作关联性切换。
 - 有关 HASToragePlus 的更多信息，请参见 [SUNW.HASToragePlus\(5\)](#) 手册页。
- 根据与其协同定位的设备组命名，后面跟 -stor-rg
 - 例如，devgrp-stor-rg。
- 同时在主群集和辅助群集上联机

配置应用程序资源组

要具有较高的可用性，应用程序必须以应用程序资源组中的资源形式接受管理。可以将应用程序资源组配置为故障转移应用程序或可伸缩应用程序。

必须在 HASToragePlus 资源中定义 ZPoolsSearchDir 扩展属性。此扩展属性是使用 ZFS 文件系统时所必需的。

主群集上配置的应用程序资源和应用程序资源组也必须在辅助群集上配置。而且，应用程序资源访问的数据也必须被复制到辅助群集上。

本节提供了配置以下应用程序资源组的准则：

- “为故障转移应用程序配置资源组” [276]
- “为可伸缩应用程序配置资源组” [277]

为故障转移应用程序配置资源组

在故障转移应用程序中，一个应用程序一次在一个节点上运行。如果此节点发生故障，应用程序将故障转移到同一群集中的另一个节点。用于故障转移应用程序的资源组必须具有以下特征：

- 当应用程序资源组发生切换或故障转移时，具有 `HAStoragePlus` 资源可以强制文件系统或 `zpool` 进行故障转移。

设备组与复制资源组和应用程序资源组位于相同的位置。因此，应用程序资源组的故障转移将强制执行设备组和复制资源组的故障转移。应用程序资源组、复制资源组和设备组由同一节点控制。

但是请注意，设备组或复制资源组的故障转移不会引起应用程序资源组的故障转移。

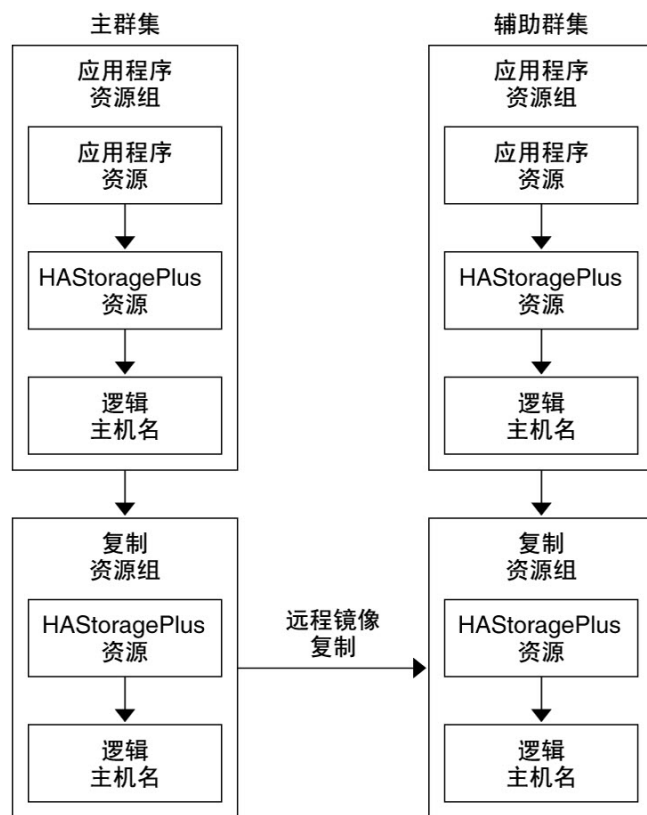
- 如果已全局安装了应用程序数据，则不需要使 `HAStoragePlus` 资源存在于应用程序资源组中，但建议使其存在。
- 如果局部安装了应用程序数据，则需要使 `HAStoragePlus` 资源存在于应用程序资源组中。

有关 `HAStoragePlus` 的更多信息，请参见 [SUNW.HAStoragePlus\(5\)](#) 手册页。

- 必须在主群集上联机而在辅助群集上脱机。
辅助群集成为主群集时，必须使应用程序资源组在辅助群集上联机。

[图 5](#) 说明了故障转移应用程序中应用程序资源组和复制资源组的配置。

图 5 配置故障转移应用程序中的资源组



为可伸缩应用程序配置资源组

在可伸缩应用程序中，一个应用程序可以在多个节点上运行以创建单一逻辑服务。如果运行可伸缩应用程序的节点发生故障，将不会发生故障转移。该应用程序将在其他节点上继续运行。

如果将可伸缩应用程序作为应用程序资源组中的资源管理，则无需将设备组配置给应用程序资源组。因此，也无需为应用程序资源组创建 HASToragePlus 资源。

用于可伸缩应用程序的资源组必须具有以下特征：

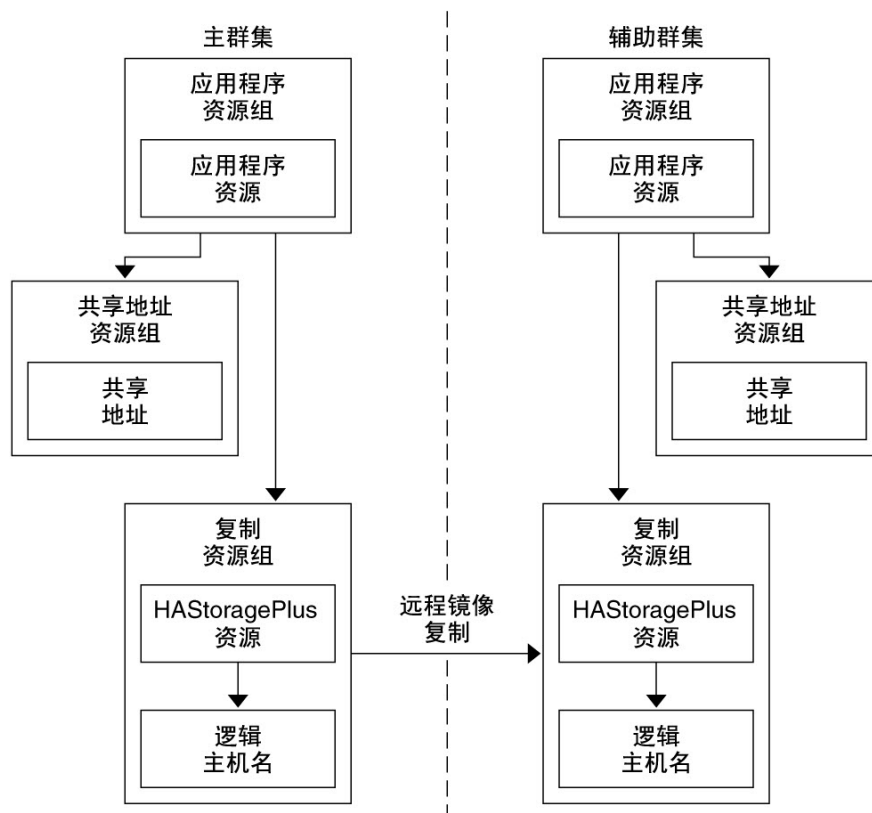
- 对共享地址资源组具有依赖性

运行可伸缩应用程序的节点要使用共享地址来分配传入的数据。

- 在主群集上联机而在辅助群集上脱机

图 6 说明了可伸缩应用程序中的资源组配置。

图 6 配置可伸缩应用程序中的资源组

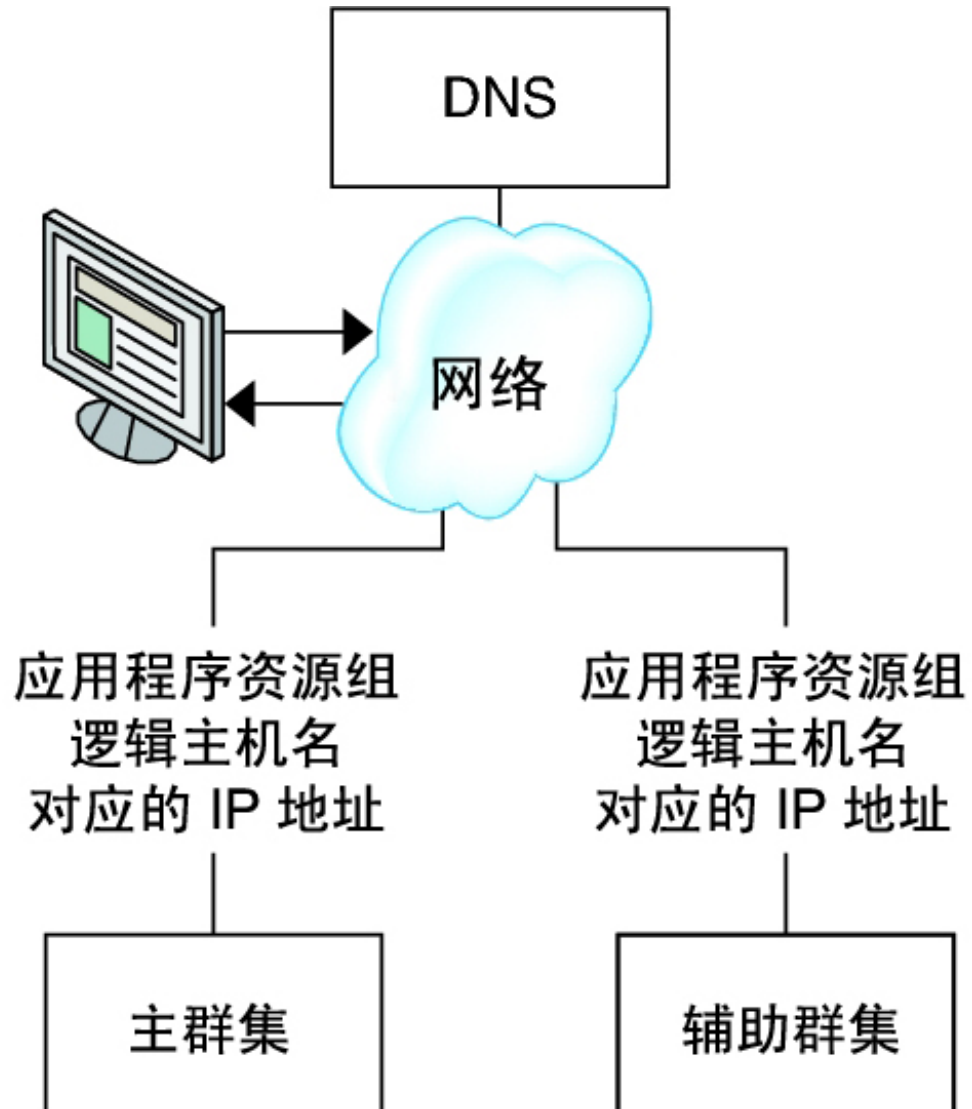


管理接管的准则

如果主群集发生故障，必须尽快将应用程序切换到辅助群集。要使辅助群集接任主群集，必须更新 DNS。

客户机使用 DNS 将应用程序的逻辑主机名映射到某个 IP 地址。完成切换（即将应用程序转移到辅助群集）后，必须更新 DNS 信息，使之反映应用程序的逻辑主机名与新 IP 地址之间的映射关系。

图 7 客户机到群集的 DNS 映射



要更新 DNS，请使用 `nsupdate` 命令。有关信息，请参见 [nsupdate\(1M\)](#) 手册页。有关如何管理接管示例，请参见“[用以管理接管的方法示例](#)” [304]。

修复后，可以使主群集重新联机。要切回到原始主群集，请执行以下任务：

- 1. 使主群集与辅助群集同步以确保主卷最新。要实现此目的，您可以停止辅助节点上的资源组，以便复制数据流可以排空。
- 2. 逆转数据复制的方向，以便原始的主节点现在可以再次将数据复制到原始的辅助节点。
- 3. 启动主群集上的资源组。
- 4. 更新 DNS 以使客户机能够访问主群集上的应用程序。

任务列表：数据复制配置示例

表 23 “任务列表：数据复制配置示例”列出了本示例中的任务，本示例描述了如何使用 Availability Suite 软件为 NFS 应用程序配置数据复制。

表 23 任务列表：数据复制配置示例

任务	说明
1. 连接和安装群集	“连接和安装群集” [280]
2. 在主群集和辅助群集上配置设备组、NFS 应用程序文件系统以及资源组	“如何配置设备组和资源组的示例” [282]
3. 在主群集和辅助群集上启用数据复制	如何在主群集上启用复制 [296] 如何在辅助群集上启用复制 [297]
4. 执行数据复制	如何执行远程镜像复制 [299] 如何执行时间点快照 [300]
5. 检验数据复制配置	如何检验复制是否已正确配置 [301]

连接和安装群集

图 8 说明了配置示例所使用的群集配置。配置示例中的辅助群集包含一个节点，但是可以使用其他群集配置。

图 8 群集配置示例

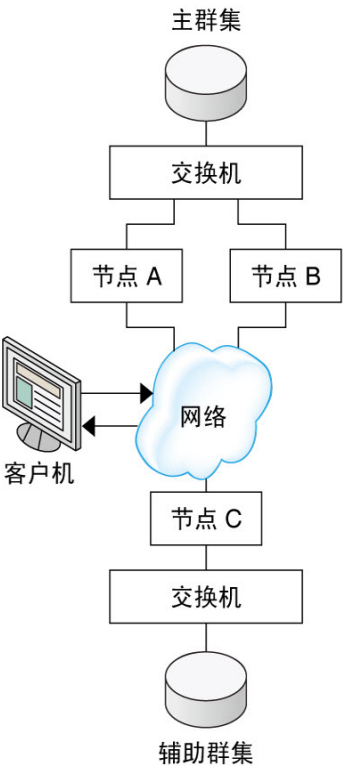


表 24 “需要的硬件和软件”概述了配置实例必需的硬件和软件。在安装 Availability Suite 软件和软件更新之前，必须先 在群集节点上安装 Oracle Solaris OS、Oracle Solaris Cluster 软件和卷管理器软件。

表 24 需要的硬件和软件

硬件或软件	要求
节点硬件	所有使用 Oracle Solaris OS 的服务器均支持 Availability Suite 软件。 有关要使用哪个硬件的信息，请参见 《Oracle Solaris Cluster Hardware Administration Manual》 。
磁盘空间	大约 15 MB。
Oracle Solaris OS	Oracle Solaris Cluster 软件支持的 Oracle Solaris OS 发行版。所有节点必须使用同一版本的 Oracle Solaris OS。 有关安装信息，请参见 《Oracle Solaris Cluster 4.3 软件安装指南》
Oracle Solaris Cluster 软件	Oracle Solaris Cluster 4.1 或更高版本软件。

硬件或软件	要求
	有关安装信息，请参见《 Oracle Solaris Cluster 4.3 软件安装指南 》。
卷管理器软件	Solaris Volume Manager 软件。所有节点必须使用相同版本的卷管理器软件。 有关安装信息，请参见《 Oracle Solaris Cluster 4.3 软件安装指南 》中的第 4 章, “ 配置 Solaris Volume Manager 软件 ”。
Availability Suite 软件	不同的群集可以使用不同版本的 Oracle Solaris OS 和 Oracle Solaris Cluster 软件，但在群集之间必须使用相同版本的 Availability Suite 软件。 有关如何安装该软件的信息，请参见 Availability Suite 软件发行版的安装手册。
Availability Suite 软件更新	有关最新软件更新的信息，请登录到 My Oracle Support 。

如何配置设备组和资源组的示例

本节介绍如何为 NFS 应用程序配置设备组和资源组。有关其他信息，请参见“[配置复制资源组](#)” [275]和“[配置应用程序资源组](#)” [275]。

本节包含以下过程：

- [如何在主群集上配置设备组](#) [283]
- [如何在辅助群集上配置设备组](#) [284]
- [如何在主群集上为 NFS 应用程序配置文件系统](#) [285]
- [如何在辅助群集上为 NFS 应用程序配置文件系统](#) [286]
- [如何在主群集上创建复制资源组](#) [287]
- [如何在辅助群集上创建复制资源组](#) [289]
- [如何在主群集上创建 NFS 应用程序资源组](#) [291]
- [如何在辅助群集上创建 NFS 应用程序资源组](#) [293]
- [如何检验复制是否已正确配置](#) [301]

下表列出了为示例配置创建的组和资源名称。

表 25 配置示例中的组 and 资源的摘要

组或资源	名称	描述
设备组	devgrp	设备组
复制资源组和资源	devgrp-stor-rg	复制资源组
	lhost-reprg-prim、lhost-reprg-sec	主群集和辅助群集上复制资源组的逻辑主机名
	devgrp-stor	复制资源组的 HAStoragePlus 资源
应用程序资源组和资源	nfs-rg	应用程序资源组
	lhost-nfsrg-prim、lhost-nfsrg-sec	主群集和辅助群集上应用程序资源组的逻辑主机名

组或资源	名称	描述
	nfs-dg-rs	应用程序的 HAStoragePlus 资源
	nfs-rs	NFS 资源

除 devgrp-stor-rg 以外，组和资源的名称均为示例名称，可根据需要更改。复制资源组的名称必须使用以下格式：*devicegroupname -stor-rg*。

有关 Solaris Volume Manager 软件的信息，请参见《[Oracle Solaris Cluster 4.3 软件安装指南](#)》中的第 4 章，“配置 Solaris Volume Manager 软件”。

▼ 如何在主群集上配置设备组

开始之前 确保您已完成以下任务：

- 阅读以下各节中的指导信息和要求：
 - “[理解 Availability Suite 软件在群集中的应用](#)” [271]
 - “[在群集间配置基于主机的数据复制的准则](#)” [274]
- 按照“[连接和安装群集](#)” [280]中的描述，设置主群集和辅助群集。

1. 通过承担可提供 `solaris.cluster.modify` RBAC 授权的角色访问 nodeA。
节点 nodeA 是主群集中的第一个节点。有关哪个节点是 nodeA 的提示，请参见图 8。

2. 创建一个元集，用以包含 NFS 数据和关联的复制。

```
nodeA# metaset -s nfsset a -h nodeA nodeB
```

3. 向元集添加磁盘。

```
nodeA# metaset -s nfsset -a /dev/did/dsk/d6 /dev/did/dsk/d7
```

4. 向元集添加中介。

```
nodeA# metaset -s nfsset -a -m nodeA nodeB
```

5. 创建所需的卷（或元设备）。

创建镜像的两个组件：

```
nodeA# metainit -s nfsset d101 1 1 /dev/did/dsk/d6s2
nodeA# metainit -s nfsset d102 1 1 /dev/did/dsk/d7s2
```

使用下述某个组件创建镜像：

```
nodeA# metainit -s nfsset d100 -m d101
```

将另一个组件附加到镜像并允许其同步：

```
nodeA# metattach -s nfsset d100 d102
```

根据以下示例从镜像创建软分区：

- **d200** – NFS 数据（主卷）：

```
nodeA# metainit -s nfsset d200 -p d100 50G
```

- **d201** – NFS 数据的时间点副本卷：

```
nodeA# metainit -s nfsset d201 -p d100 50G
```

- **d202** – 时间点位图卷：

```
nodeA# metainit -s nfsset d202 -p d100 10M
```

- **d203** – 远程影子位图卷：

```
nodeA# metainit -s nfsset d203 -p d100 10M
```

- **d204** – Oracle Solaris Cluster SUNW.NFS 配置信息卷：

```
nodeA# metainit -s nfsset d204 -p d100 100M
```

6. 创建 NFS 数据和配置卷的文件系统。

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d200
```

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d204
```

接下来的步骤 请转到[如何在辅助群集上配置设备组 \[284\]](#)。

▼ 如何在辅助群集上配置设备组

开始之前 请完成[如何在主群集上配置设备组 \[283\]](#)中所述的过程。

1. 通过承担可提供 `solaris.cluster.modify` RBAC 授权的角色访问 `nodeC`。
2. 创建一个元集，用以包含 NFS 数据和关联的复制。

```
nodeC# metaset -s nfsset a -h nodeC
```

3. 向元集添加磁盘。

在下面的示例中，假定磁盘 DID 编号各不相同。

```
nodeC# metaset -s nfsset -a /dev/did/dsk/d3 /dev/did/dsk/d4
```

注 - 在单节点群集上，中介不是必需的。

4. 创建所需的卷（或元设备）。

创建镜像的两个组件：

```
nodeC# metainit -s nfsset d101 1 1 /dev/did/dsk/d3s2
nodeC# metainit -s nfsset d102 1 1 /dev/did/dsk/d4s2
```

使用下述某个组件创建镜像：

```
nodeC# metainit -s nfsset d100 -m d101
```

将另一个组件附加到镜像并允许其同步：

```
metattach -s nfsset d100 d102
```

根据以下示例从镜像创建软分区：

- *d200* – NFS 数据（主卷）：

```
nodeC# metainit -s nfsset d200 -p d100 50G
```

- *d201* – NFS 数据的时间点副本卷：

```
nodeC# metainit -s nfsset d201 -p d100 50G
```

- *d202* – 时间点位图卷：

```
nodeC# metainit -s nfsset d202 -p d100 10M
```

- *d203* – 远程影子位图卷：

```
nodeC# metainit -s nfsset d203 -p d100 10M
```

- *d204* – Oracle Solaris Cluster SUNW.NFS 配置信息卷：

```
nodeC# metainit -s nfsset d204 -p d100 100M
```

5. 创建 NFS 数据和配置卷的文件系统。

```
nodeC# yes | newfs /dev/md/nfsset/rdisk/d200
nodeC# yes | newfs /dev/md/nfsset/rdisk/d204
```

接下来的步骤 请转到[如何在主群集上为 NFS 应用程序配置文件系统 \[285\]](#)。

▼ 如何在主群集上为 NFS 应用程序配置文件系统

开始之前 请完成[如何在辅助群集上配置设备组 \[284\]](#)中所述的过程。

1. 在 *nodeA* 和 *nodeB* 上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 在 *nodeA* 和 *nodeB* 上，为 NFS 文件系统创建一个挂载点目录。

例如：

```
nodeA# mkdir /global/mountpoint
```

3. 在 **nodeA** 和 **nodeB** 上，将主卷配置为不自动挂载到挂载点上。
在 **nodeA** 和 **nodeB** 上的 `/etc/vfstab` 文件中添加或替换以下文本。该文本必须在一行。

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \  
/global/mountpoint ufs 3 no global,logging
```

4. 在 **nodeA** 和 **nodeB** 上，为元设备 **d204** 创建一个挂载点。
以下示例创建挂载点 `/global/etc`。

```
nodeA# mkdir /global/etc
```

5. 在 **nodeA** 和 **nodeB** 上，将元设备 **d204** 配置为自动挂载到挂载点上。
在 **nodeA** 和 **nodeB** 上的 `/etc/vfstab` 文件中添加或替换以下文本。该文本必须在一行。

```
/dev/md/nfsset/dsk/d204 /dev/md/nfsset/rdisk/d204 \  
/global/etc ufs 3 yes global,logging
```

6. 将元设备 **d204** 挂载到 **nodeA** 上。

```
nodeA# mount /global/etc
```

7. 为 Oracle Solaris Cluster HA for NFS 数据服务创建配置文件和信息。

- a. 在 **nodeA** 上创建一个名为 `/global/etc/SUNW.nfs` 的目录。

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

- b. 在 **nodeA** 上创建文件 `/global/etc/SUNW.nfs/dfstab.nfs-rs`。

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. 将下面一行添加到 **nodeA** 上的 `/global/etc/SUNW.nfs/dfstab.nfs-rs` 文件中。

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

接下来的步骤 请转到[如何在辅助群集上为 NFS 应用程序配置文件系统 \[286\]](#)。

▼ 如何在辅助群集上为 NFS 应用程序配置文件系统

开始之前 完成[如何在主群集上为 NFS 应用程序配置文件系统 \[285\]](#)中所述的过程。

1. 在 **nodeC** 上，承担可提供 `solaris.cluster.admin` RBAC 授权的角色。
2. 在 **nodeC** 上，为 NFS 文件系统创建一个挂载点目录。

例如：

```
nodeC# mkdir /global/mountpoint
```

3. 在 **nodeC** 上，将主卷配置为自动在挂载点上挂载。
在 **nodeC** 上的 `/etc/vfstab` 文件中添加或替换以下文本。该文本必须在一行。

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \  
/global/mountpoint ufs 3 yes global,logging
```

4. 将元设备 **d204** 挂载到 **nodeA** 上。

```
nodeC# mount /global/etc
```

5. 为 Oracle Solaris Cluster HA for NFS 数据服务创建配置文件和信息。

- a. 在 **nodeA** 上创建一个名为 `/global/etc/SUNW.nfs` 的目录。

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

- b. 在 **nodeA** 上创建文件 `/global/etc/SUNW.nfs/dfstab.nfs-rs`。

```
nodeC# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. 将下面一行添加到 **nodeA** 上的 `/global/etc/SUNW.nfs/dfstab.nfs-rs` 文件中。

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

接下来的步骤 请转到[如何在主群集上创建复制资源组 \[287\]](#)。

▼ 如何在主群集上创建复制资源组

- 开始之前
- 完成[如何在辅助群集上为 NFS 应用程序配置文件系统 \[286\]](#)中所述的过程。
 - 请确保 `/etc/netmasks` 文件具有所有逻辑主机名对应的 IP 地址子网和网络掩码条目。如有必要，编辑 `/etc/netmasks` 文件以添加缺少的任何条目。

1. 作为可提供 `solaris.cluster.modify`、`solaris.cluster.admin` 和 `solaris.cluster.read` RBAC 授权的角色访问 **nodeA**。

2. 注册 **SUNW.HAStoragePlus** 资源类型。

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

3. 为设备组创建一个复制资源组。

```
nodeA# clresourcegroup create -n nodeA,nodeB devgrp-stor-rg
```

-n nodeA,nodeB

指定群集节点 nodeA 和 nodeB 可控制该复制资源组。

devgrp-stor-rg

复制资源组的名称。在此名称中，devgrp 指定了设备组的名称。

4. 向复制资源组添加 **SUNW.HAStoragePlus** 资源。

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HAStoragePlus \  
-p GlobalDevicePaths=nfsset \  
-p AffinityOn=True \  
devgrp-stor
```

-g

指定资源将被添加到哪一个资源组。

-p GlobalDevicePaths=

指定 Availability Suite 软件所依赖的设备组。

-p AffinityOn=True

指定 SUNW.HAStoragePlus 资源必须对 -p GlobalDevicePaths= 所定义的全局设备和群集文件系统执行关联性切换。因此，复制资源组发生故障转移或被切换后，相关的设备组也将被切换。

有关这些扩展属性的更多信息，请参见 [SUNW.HAStoragePlus\(5\)](#) 手册页。

5. 为复制资源组添加逻辑主机名资源。

```
nodeA# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-prim
```

主群集上复制资源组的逻辑主机名为 lhost-reprg-prim。

6. 启用资源、管理资源组并使资源组联机。

```
nodeA# clresourcegroup online -emM -n nodeA devgrp-stor-rg
```

-e

启用相关联的资源。

-M

管理资源组。

-n

指定在哪个节点上使资源组联机。

7. 检验资源组是否处于联机状态。


```
nodeA# clresourcegroup status devgrp-stor-rg
```

检查资源组状态字段以确认该复制资源组在 nodeA 上处于联机状态。

接下来的步骤 请转到[如何在辅助群集上创建复制资源组 \[289\]](#)。

▼ 如何在辅助群集上创建复制资源组

- 开始之前
- 完成[如何在主群集上创建复制资源组 \[287\]](#)中所述的过程。
 - 请确保 `/etc/netmasks` 文件具有所有逻辑主机名对应的 IP 地址子网和网络掩码条目。如有必要，编辑 `/etc/netmasks` 文件以添加缺少的任何条目。

1. 作为可提供 `solaris.cluster.modify`、`solaris.cluster.admin` 和 `solaris.cluster.read` RBAC 授权的角色访问 nodeC。

2. 将 `SUNW.HAStoragePlus` 注册为资源类型。

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

3. 为设备组创建一个复制资源组。

```
nodeC# clresourcegroup create -n nodeC devgrp-stor-rg
```

```
create
```

创建资源组。

```
-n
```

指定资源组的节点列表。

```
devgrp
```

设备组的名称。

```
devgrp-stor-rg
```

复制资源组的名称。

4. 向复制资源组中添加一个 `SUNW.HAStoragePlus` 资源。

```
nodeC# clresource create \
-t SUNW.HAStoragePlus \
-p GlobalDevicePaths=nfsset \
-p AffinityOn=True \
devgrp-stor
```

```
create
```

创建资源。

-t

指定资源类型。

-p GlobalDevicePaths=

指定 Availability Suite 软件所依赖的设备组。

-p AffinityOn=True

指定 SUNW.HAStoragePlus 资源必须对 -p GlobalDevicePaths= 所定义的全局设备和群集文件系统执行关联性切换。因此，复制资源组发生故障转移或被切换后，相关的设备组也将被切换。

devgrp-stor

复制资源组的 HAStoragePlus 资源。

有关这些扩展属性的更多信息，请参见 [SUNW.HAStoragePlus\(5\)](#) 手册页。

5. 为复制资源组添加逻辑主机名资源。

```
nodeC# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-sec
```

辅助群集上复制资源组的逻辑主机名为 lhost-reprg-sec。

6. 启用资源、管理资源组并使资源组联机。

```
nodeC# clresourcegroup online -eM -n nodeC devgrp-stor-rg
```

online

联机。

-e

启用相关联的资源。

-M

管理资源组。

-n

指定在哪个节点上使资源组联机。

7. 检验资源组是否处于联机状态。

```
nodeC# clresourcegroup status devgrp-stor-rg
```

检查资源组状态字段以确认该复制资源组在 nodeC 上处于联机状态。

接下来的步骤 请转到[如何在主群集上创建 NFS 应用程序资源组 \[291\]](#)。

▼ 如何在主群集上创建 NFS 应用程序资源组

此过程描述了如何为 NFS 创建应用程序资源组。此过程是特定于该应用程序的，且不能用于其他类型的应用程序。

- 开始之前
- 完成[如何在辅助群集上创建复制资源组 \[289\]](#)中所述的过程。
 - 请确保 `/etc/netmasks` 文件具有所有逻辑主机名对应的 IP 地址子网和网络掩码条目。如有必要，编辑 `/etc/netmasks` 文件以添加缺少的任何条目。

1. 作为可提供 `solaris.cluster.modify`、`solaris.cluster.admin` 和 `solaris.cluster.read` RBAC 授权的角色访问 `nodeA`。

2. 将 `SUNW.nfs` 注册为资源类型。

```
nodeA# clresourcetype register SUNW.nfs
```

3. 将 `SUNW.HAStoragePlus` 注册为资源类型（如果它尚未注册）。

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

4. 为 NFS 服务创建一个应用程序资源组。

```
nodeA# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_affinities=+++devgrp-stor-rg \
nfs-rg
```

```
Pathprefix=/global/etc
```

指定组中资源可将管理文件写入哪个目录。

```
Auto_start_on_new_cluster=False
```

指定不自动启动应用程序资源组。

```
RG_affinities=+++devgrp-stor-rg
```

指定应用程序资源组必须与之并置的资源组。在本示例中，该应用程序资源组必须与复制资源组 `devgrp-stor-rg` 并置。

如果复制资源组被切换到新的主节点上，应用程序资源组也会被自动切换。不过，系统将阻止您将应用程序资源组切换到新的主节点，因为该操作会违反并置要求。

```
nfs-rg
```

应用程序资源组的名称

5. 向应用程序资源组中添加一个 `SUNW.HAStoragePlus` 资源。

```
nodeA# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
```

```
-p FileSystemMountPoints=/global/mountpoint \  
-p AffinityOn=True \  
nfs-dg-rs
```

create

创建资源。

-g

指定资源将被添加到哪个资源组。

-t SUNW.HAStoragePlus

指定资源的类型是 SUNW.HAStoragePlus。

-p FileSystemMountPoints=/global/mountpoint

指定文件系统的挂载点为全局挂载点。

-p AffinityOn=True

指定应用程序资源必须对 -p FileSystemMountPoints 所定义的全局设备和群集文件系统执行关联性切换。因此，当应用程序资源组发生故障转移或被切换后，相关的设备组也将被切换。

nfs-dg-rs

用于 NFS 应用程序的 HAStoragePlus 资源的名称。

有关这些扩展属性的更多信息，请参见 [SUNW.HAStoragePlus\(5\)](#) 手册页。

6. 为应用程序资源组添加逻辑主机名资源。

```
nodeA# clreslogicalhostname create -g nfs-rg \  
lhost-nfsrg-prim
```

主群集上应用程序资源组的逻辑主机名为 lhost-nfsrg-prim。

7. 创建 `dfstab.resource-name` 配置文件并将其放置在包含资源组的 `Pathprefix` 目录下的 `SUNW.nfs` 子目录中。

a. 在 `nodeA` 上创建一个名为 `SUNW.nfs` 的目录。

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

b. 在 `nodeA` 上创建文件 `dfstab.resource-name`。

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

c. 将下面一行添加到 `nodeA` 上的 `/global/etc/SUNW.nfs/dfstab.nfs-rs` 文件中。

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

8. 使应用程序资源组联机。

```
nodeA# clresourcegroup online -M -n nodeA nfs-rg
```

```
online
```

使资源组联机。

```
-e
```

启用相关联的资源。

```
-M
```

管理资源组。

```
-n
```

指定在哪个节点上使资源组联机。

```
nfs-rg
```

资源组的名称。

9. 检验应用程序资源组是否处于联机状态。

```
nodeA# clresourcegroup status
```

检查资源组状态字段，确定该应用程序资源组在 nodeA 和 nodeB 上是否处于联机状态。

接下来的步骤 请转到[如何在辅助群集上创建 NFS 应用程序资源组 \[293\]](#)。

▼ 如何在辅助群集上创建 NFS 应用程序资源组

开始之前

- 完成[如何在主群集上创建 NFS 应用程序资源组 \[291\]](#)中所述的过程。
- 请确保 /etc/netmasks 文件具有所有逻辑主机名对应的 IP 地址子网和网络掩码条目。如有必要，编辑 /etc/netmasks 文件以添加缺少的任何条目。

1. 作为可提供 `solaris.cluster.modify`、`solaris.cluster.admin` 和 `solaris.cluster.read` RBAC 授权的角色访问 nodeC。

2. 将 `SUNW.nfs` 注册为资源类型。

```
nodeC# clresourcetype register SUNW.nfs
```

3. 将 `SUNW.HAStoragePlus` 注册为资源类型（如果它尚未注册）。

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

4. 为设备组创建一个应用程序资源组。

```
nodeC# clresourcegroup create \  
-p Pathprefix=/global/etc \  
-p Auto_start_on_new_cluster=False \  
-p RG_affinities=+++devgrp-stor-rg \  
nfs-rg
```

create

创建资源组。

-p

指定资源组的属性。

Pathprefix=/global/etc

指定组中资源可以在哪个目录中写入管理文件。

Auto_start_on_new_cluster=False

指定不自动启动应用程序资源组。

RG_affinities=+++devgrp-stor-rg

指定应用程序资源组必须与之并置的资源组。在本示例中，该应用程序资源组必须与复制资源组 devgrp-stor-rg 并置。

如果复制资源组被切换到新的主节点上，应用程序资源组也会被自动切换。不过，系统将阻止您将应用程序资源组切换到新的主节点，因为该操作会违反并置要求。

nfs-rg

应用程序资源组的名称

5. 向应用程序资源组中添加一个 SUNW.HAStoragePlus 资源。

```
nodeC# clresource create -g nfs-rg \  
-t SUNW.HAStoragePlus \  
-p FileSystemMountPoints=/global/mountpoint \  
-p AffinityOn=True \  
nfs-dg-rs
```

create

创建资源。

-g

指定资源将被添加到哪个资源组。

-t SUNW.HAStoragePlus

指定资源的类型是 SUNW.HAStoragePlus。

-p

指定资源的属性。

```
FileSystemMountPoints=/global/mountpoint
```

指定文件系统的挂载点为全局挂载点。

```
AffinityOn=True
```

指定应用程序资源必须对 -p FileSystemMountPoints= 所定义的全局设备和群集文件系统执行关联性切换。因此，当应用程序资源组发生故障转移或被切换后，相关的设备组也将被切换。

```
nfs-dg-rs
```

用于 NFS 应用程序的 HASToragePlus 资源的名称。

6. 为应用程序资源组添加逻辑主机名资源。

```
nodeC# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-sec
```

辅助群集上应用程序资源组的逻辑主机名为 lhost-nfsrg-sec。

7. 向应用程序资源组中添加 NFS 资源。

```
nodeC# clresource create -g nfs-rg \
-t SUNW.nfs -p Resource_dependencies=nfs-dg-rs nfs-rg
```

8. 如果在主群集上挂载全局卷，应从辅助群集上卸载全局卷。

```
nodeC# umount /global/mountpoint
```

如果在辅助群集上挂载卷，同步将失败。

接下来的步骤 请转到[“如何启用数据复制的示例” \[295\]](#)。

如何启用数据复制的示例

本节描述了如何为示例配置启用数据复制。本节使用了 Availability Suite 软件的命令 `sndradm` 和 `iiadm`。有关这些命令的更多信息，请参见 Availability Suite 文档。

本节包含以下过程：

- [如何在主群集上启用复制 \[296\]](#)
- [如何在辅助群集上启用复制 \[297\]](#)

▼ 如何在主群集上启用复制

1. 作为可提供 `solaris.cluster.read` RBAC 授权的角色访问 `nodeA`。
2. 刷新所有事务。

```
nodeA# lockfs -a -f
```

3. 确认逻辑主机名 `lhost-reprg-prim` 和 `lhost-reprg-sec` 处于联机状态。

```
nodeA# clresourcegroup status
nodeC# clresourcegroup status
```

检查资源组的状态字段。

4. 启用从主群集到辅助群集的远程镜像复制。

此步骤启用从主群集到辅助群集的复制。此步骤启用从主群集上的主卷 (`d200`) 到辅助群集上的主卷 (`d200`) 的复制。此外，此步骤还启用到 `d203` 上的远程镜像位图的复制。

- 如果主群集和辅助群集不同步，请对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

- 如果主群集和辅助群集同步，请对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -E lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

5. 启用自动同步。

对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -a on lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

该步骤启用了自动同步。如果自动同步的活动状态设置为 `on`，则当系统重新引导或发生故障时，将重新同步卷集。

6. 检验群集是否处于记录模式。

对 Availability Suite 软件使用以下命令：

```
nodeA# /usr/sbin/sndradm -P
```

输出应与以下所示类似：

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

在日志记录模式下，状态为 logging，自动同步的活动状态为 off。当磁盘上的数据卷被写入时，即更新同一磁盘上的位图文件。

7. 启用时间点快照。

对 Availability Suite 软件使用以下命令：

```
nodeA# /usr/sbin/iiadm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeA# /usr/sbin/iiadm -w \
/dev/md/nfsset/rdisk/d201
```

此步骤使主群集上的主卷将被复制到同一群集上的影子卷中。主卷、影子卷和时间点位图卷必须在同一设备组中。在本示例中，主卷为 d200，影子卷为 d201，时间点位图卷为 d203。

8. 将时间点快照附加到远程镜像集。

对 Availability Suite 软件使用以下命令：

```
nodeA# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

此步骤将时间点快照与远程镜像卷集相关联。Availability Suite 软件可以确保在进行远程镜像复制之前创建时间点快照。

接下来的步骤 请转到[如何在辅助群集上启用复制 \[297\]](#)。

▼ 如何在辅助群集上启用复制

开始之前 完成[如何在主群集上启用复制 \[296\]](#)中所述的过程。

1. 作为 root 角色访问 nodeC。
2. 刷新所有事务。

```
nodeC# lockfs -a -f
```

3. 启用从主群集到辅助群集的远程镜像复制。

对 Availability Suite 软件使用以下命令：

```
nodeC# /usr/sbin/sndradm -n -e lhost-reprg-prim \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 ip sync
```

主群集检测到存在辅助群集并启动同步。有关群集状态的信息，请参阅 Availability Suite 的系统日志文件 `/var/adm`。

4. 启用独立的时间点快照。

对 Availability Suite 软件使用以下命令：

```
nodeC# /usr/sbin/iidm -e ind \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d201 \  
/dev/md/nfsset/rdisk/d202  
nodeC# /usr/sbin/iidm -w \  
/dev/md/nfsset/rdisk/d201
```

5. 将时间点快照附加到远程镜像集。

对 Availability Suite 软件使用以下命令：

```
nodeC# /usr/sbin/sndradm -I a \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d201 \  
/dev/md/nfsset/rdisk/d202
```

接下来的步骤 请转到[“如何执行数据复制的示例” \[298\]](#)。

如何执行数据复制的示例

本节描述了如何为示例配置执行数据复制。本节使用了 Availability Suite 软件的命令 `sndradm` 和 `iidm`。有关这些命令的更多信息，请参见 Availability Suite 文档。

本节包含以下过程：

- [如何执行远程镜像复制 \[299\]](#)
- [如何执行时间点快照 \[300\]](#)
- [如何检验复制是否已正确配置 \[301\]](#)

▼ 如何执行远程镜像复制

在此过程中，将主磁盘的主卷复制到辅助磁盘上的主卷。主卷为 d200，远程镜像位图卷为 d203。

1. 作为 **root** 角色访问 **nodeA**。
2. 检验群集是否处于记录模式。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -P
```

输出应与以下所示类似：

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

在日志记录模式下，状态为 logging，自动同步的活动状态为 off。当磁盘上的数据卷被写入时，即更新同一磁盘上的位图文件。

3. 刷新所有事务。

nodeA# **lockfs -a -f**
4. 在 **nodeC** 上重复[步骤 1](#)到[步骤 3](#)。
5. 将 **nodeA** 的主卷复制到 **nodeC** 的主卷上。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -m lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

6. 完成复制和同步卷之前，请等待。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

7. 确认群集是否处于复制模式。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -P
```

输出应与以下所示类似：

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

在复制模式下，状态为 `replicating`，而自动同步的活动状态为 `on`。对主卷进行写入时，由 Availability Suite 软件对辅助卷进行更新。

接下来的步骤 请转到[如何执行时间点快照 \[300\]](#)。

▼ 如何执行时间点快照

在此过程中，时间点快照用于将主群集的影子卷同步到主群集的主卷中。主卷为 `d200`，位图卷为 `d203`，影子卷为 `d201`。

开始之前 完成[如何执行远程镜像复制 \[299\]](#)中所述的过程。

1. 作为可提供 `solaris.cluster.modify` 和 `solaris.cluster.admin` RBAC 授权的角色访问 `nodeA`。

2. 禁用 `nodeA` 上正在运行的资源。

```
nodeA# clresource disable nfs-rs
```

3. 将主群集更改为日志模式。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

当磁盘上的数据卷被写入时，即更新同一磁盘上的位图文件。系统不会执行任何复制操作。

4. 将主群集的影子卷同步到主群集的主卷。
对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/iiadm -u s /dev/md/nfsset/rdisk/d201
nodeA# /usr/sbin/iiadm -w /dev/md/nfsset/rdisk/d201
```

5. 将辅助群集的影子卷同步到辅助群集的主卷。

对 Availability Suite 软件运行以下命令：

```
nodeC# /usr/sbin/iiadm -u s /dev/md/nfsset/rdisk/d201
nodeC# /usr/sbin/iiadm -w /dev/md/nfsset/rdisk/d201
```

6. 在 nodeA 上重新启动该应用程序。

```
nodeA# clresource enable nfs-rs
```

7. 使辅助卷与主卷重新同步。

对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

接下来的步骤 请转到[如何检验复制是否已正确配置 \[301\]](#)。

▼ 如何检验复制是否已正确配置

开始之前 完成[如何执行时间点快照 \[300\]](#)中所述的过程。

1. 作为可提供 `solaris.cluster.admin` RBAC 授权的角色访问 nodeA 和 nodeC。
2. 检验主群集是否处于复制模式并已启用自动同步。

对 Availability Suite 软件使用以下命令：

```
nodeA# /usr/sbin/sndradm -P
```

输出应与以下所示类似：

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

在复制模式下，状态为 `replicating`，而自动同步的活动状态为 `on`。对主卷进行写入时，由 Availability Suite 软件对辅助卷进行更新。

3. 如果主群集未处于复制模式下，则使其处于复制模式下。
- 对 Availability Suite 软件使用以下命令：

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
```

```
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 ip sync
```

4. 在客户机上创建目录。

- a. 作为 **root** 角色登录到客户机。

您将看到类似下面的提示：

```
client-machine#
```

- b. 在客户机上创建目录。

```
client-machine# mkdir /dir
```

5. 将主卷挂载到应用程序目录上，并显示已挂载的目录。

- a. 将主卷挂载到应用程序目录上。

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

- b. 显示已挂载的目录。

```
client-machine# ls /dir
```

6. 从应用程序目录卸载主卷。

- a. 从应用程序目录卸载主卷。

```
client-machine# umount /dir
```

- b. 使主群集上的应用程序资源组脱机。

```
nodeA# clresource disable -g nfs-rg +  
nodeA# clresourcegroup offline nfs-rg
```

- c. 将主群集更改为日志模式。

对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \  
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 ip sync
```

当磁盘上的数据卷被写入时，即更新同一磁盘上的位图文件。系统不会执行任何复制操作。

- d. 确保 **PathPrefix** 目录可用。

```
nodeC# mount | grep /global/etc
```

- e. 确认文件系统适合挂载到辅助群集上。

```
nodeC# fsck -y /dev/md/nfsset/rdisk/d200
```

- f. 将应用程序置于受管状态，然后使其在辅助群集上联机。

```
nodeC# clresourcegroup online -eM nodeC nfs-rg
```

- g. 作为 **root** 角色访问客户机。

您将看到类似下面的提示：

```
client-machine#
```

- h. 将您在[步骤 4](#)中创建的应用程序目录挂载到辅助卷上的应用程序目录中。

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

- i. 显示已挂载的目录。

```
client-machine# ls /dir
```

7. 确保[步骤 5](#)中显示的目录与[步骤 6](#)中显示的目录相同。

8. 使主卷上的应用程序返回到已挂载应用程序目录。

- a. 使应用程序资源组在辅助卷上脱机。

```
nodeC# clresource disable -g nfs-rg +
nodeC# clresourcegroup offline nfs-rg
```

- b. 确保从辅助卷上卸载全局卷。

```
nodeC# umount /global/mountpoint
```

- c. 将应用程序资源组置于受管状态，然后使其在主群集上联机。

```
nodeA# clresourcegroup online -eM nodeA nfs-rg
```

- d. 将主卷更改为复制模式。

对 Availability Suite 软件运行以下命令：

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

对主卷进行写入时，由 Availability Suite 软件对辅助卷进行更新。

另请参见 [“用以管理接管的方法示例” \[304\]](#)

用以管理接管的方法示例

本节介绍了如何更新 DNS 条目。有关其他信息，请参见[“管理接管的准则” \[278\]](#)。

▼ 如何更新 DNS 条目

有关说明 DNS 如何将客户机映射到群集的图示，请参见[图 7](#)。

1. 启动 `nsupdate` 命令。
有关更多信息，请参见 [nsupdate\(1M\)](#) 手册页。
2. 针对两个群集（主群集和辅助群集），删除当前在应用程序资源组的逻辑主机名与群集 IP 地址之间存在的 DNS 映射。

```
> update delete lhost-nfsrg-prim A
> update delete lhost-nfsrg-sec A
> update delete ipaddress1rev.in-addr.arpa ttl PTR lhost-nfsrg-prim
> update delete ipaddress2rev.in-addr.arpa ttl PTR lhost-nfsrg-sec
```

ipaddress1rev

主群集的 IP 地址，按照逆序。

ipaddress2rev

辅助群集的 IP 地址，按照逆序。

ttl

有效时间（秒）。典型值为 3600。

3. 针对两个群集（主群集和辅助群集），创建应用程序资源组的逻辑主机名与群集 IP 地址之间的新 DNS 映射。

将主逻辑主机名映射到辅助群集的 IP 地址，将辅助逻辑主机名映射到主群集的 IP 地址。

```
> update add lhost-nfsrg-prim ttl A ipaddress2fwd
> update add lhost-nfsrg-sec ttl A ipaddress1fwd
> update add ipaddress2rev.in-addr.arpa ttl PTR lhost-nfsrg-prim
```



```
> update add ipaddress1rev.in-addr.arpa ttl PTR lhost-nfsrg-sec
```

ipaddress2fwd

辅助群集的 IP 地址，按照正序。

ipaddress1fwd

主群集的 IP 地址，按照正序。

索引

A

安全密钥

配置, 267

AffinityOn 属性

用于数据复制的扩展属性, 275

Availability Suite

用于数据复制, 271

B

备份

群集, 29, 259

镜像联机, 259

备用 BE

修补 solaris10 标记区域群集, 256

本地镜像 见 基于存储的复制

boot 命令, 68

BUI 见 Oracle Solaris Cluster Manager

C

查找

全局群集的节点 ID, 207

区域群集的节点 ID, 207

超时

更改法定设备的默认值, 169

传输电缆

启用, 182

添加, 37, 40, 175, 180

禁用, 183

传输电缆, 添加, 40, 175, 180

传输交换机, 添加, 37, 40, 175, 180

传输适配器, 添加, 37

磁盘路径

取消监视, 147

监视, 99, 145, 146

显示故障磁盘路径, 148

解决状态错误, 148

从 IPMP 组取消激活网络接口, 189

存储阵列

删除, 202

错误消息

/var/adm/messages 文件, 91

删除节点, 204

cconsole 命令 见 pconsole 实用程序

重命名节点

在全局群集中, 214

在区域群集中, 214

重新启动

全局群集节点, 86

区域群集节点, 86

重新引导

全局群集, 72

全局群集节点, 86

区域群集, 72

区域群集节点, 86

claccess 命令, 26

cldevice 命令, 26

cldevicegroup 命令, 26

clinterconnect 命令, 26

clnasdevice 命令, 26

clnode 命令, 229, 230

clquorum 命令, 26

clreslogicalhostname 命令, 26

clresource 命令, 26

删除资源和资源组, 236

删除资源组, 236

clresourcegroup 命令, 26, 230

clresourcetype 命令, 26

clressharedaddress 命令, 26

clsetup, 26, 26, 31

创建区域群集, 24

- 向区域群集添加网络地址, 234
- 管理传输交换机, 175
- 管理区域群集, 231
- 管理法定设备, 153
- 管理设备组, 112
- clsetup 实用程序
 - 创建区域群集, 23
- clsnmp host 命令, 26
- clsnmpmib 命令, 26
- clsnmpuser 命令, 26
- cltelemattribute 命令, 26
- cluster check
 - 命令, 49
- cluster check 命令, 26
 - vfstab 文件检查, 142
- cluster shutdown 命令, 65
- clzonecluster 命令, 26
 - 修补 solaris10 标记区域群集, 255
 - 停止, 65
 - 引导, 68
 - 描述, 31
- Common Agent Container
 - 配置安全密钥, 267
- CPU
 - 配置, 248
- CPU 份额
 - 全局群集节点, 248
 - 控制, 247
 - 配置, 247

D

- 登录
 - 远程, 30
- 电缆, 传输, 180
- 电源管理, 205
- 动态重新配置, 100, 100
 - 公共网络接口, 189
 - 法定设备, 154
 - 群集互连, 176
- DID 信息
 - 手动更新, 148
- DLPI, 179

E

- /etc/vfstab 文件, 54
 - 检验配置, 142
 - 添加挂载点, 141
- EMC SRDF
 - Domino 模式, 96
 - 最佳做法, 98
 - 校园群集的主工作间彻底故障转移后进行恢复, 110
 - 检验配置, 105
 - 管理, 101
 - 要求, 97
 - 配置 DID 设备, 103
 - 配置复制组, 101
 - 配置示例, 105
 - 限制, 97

F

- 发行版本信息, 32
- 法定
 - 概述, 153
 - 管理, 153
- 法定服务器 见 法定设备
- 法定服务器法定设备
 - 安装要求, 158
 - 排除删除操作引起的故障, 161
 - 添加, 158
- 法定设备
 - 修复, 168
 - 修改节点列表, 164
 - 列出配置, 167
 - 删除, 154, 160
 - 最后一个法定设备, 161
 - 动态重新配置设备, 154
 - 和基于存储的复制, 98
 - 支持的类型, 155
 - 更改默认超时时间, 169
 - 替换, 163
 - 添加, 156
 - Oracle ZFS Storage Appliance NAS, 157
 - 法定服务器, 158
 - 直接连接的共享磁盘法定设备, 156
 - 维护状态, 使设备脱离, 166
 - 维护状态, 将设备置于, 164
- 访问

- Oracle Solaris Cluster Manager, 266
- 非群集模式引导, 89
- 辅助功能
 - Oracle Solaris Cluster Manager, 269
- 辅助节点
 - 设置所需数目, 130
 - 默认编号, 128
- 负载限制
 - concentrate_load 属性, 229
 - preemption_mode 属性, 229
 - 在节点上配置, 229, 230
- 复制 见 数据复制
- 复制, 基于存储, 95
- failback 属性, 128
- fence_level 见 复制过程中

G

- 概述
 - 法定, 153
- 更改
 - numsecondaries 属性, 130
 - SNMP 事件 MIB 协议, 225
 - 专用主机名, 212
 - 主节点, 133
 - 属性, 128
 - 群集名称, 206
- 更新, 254
 - 参见 修补
 - AI 服务器, 256
 - Oracle GlassFish 的限制, 265
 - Oracle Solaris Cluster Manager 的限制, 265
 - solaris 标记区域群集, 255
 - 技巧, 253
 - 概述, 251
 - 法定服务器, 256
 - 特定软件包, 253
- 更新全局名称空间, 114
- 公共网络
 - 动态重新配置, 189
 - 管理, 175, 188
- 公平份额调度器
 - CPU 份额配置, 248
- 共享 SCSI 磁盘
 - 支持用作法定设备, 155
- 共享磁盘路径

- 启用自动重新引导功能, 150
- 监视, 145
- 禁用自动重新引导功能, 151
- 故障排除
 - network-bind-address, 268
 - Oracle Solaris Cluster Manager, 267
 - 传输电缆, 适配器, 和交换机, 187
- 挂载点
 - 修改 /etc/vfstab 文件, 141
 - 全局, 54
- 关闭
 - 全局群集, 65
 - 全局群集节点, 79
 - 区域群集, 65
 - 区域群集节点, 79
 - 节点, 79
- 关联性切换
 - 为数据复制配置, 288
- 管理
 - EMC SRDF 复制设备, 101
 - IPMP, 188
 - PNM, 175
 - 全局群集, 24
 - 全局群集设置, 205
 - 区域群集, 24, 231
 - 基于存储的复制设备, 100
 - 法定, 153
 - 群集互连和公共网络, 175
 - 群集文件系统, 112
- 管理工具
 - Oracle Solaris Cluster Manager, 265
- 管理控制台, 28
- GlobalDevicePaths
 - 用于数据复制的扩展属性, 275
- GUI 见 Oracle Solaris Cluster Manager

H

- 互连
 - 启用, 187
 - 故障排除, 187
- 恢复
 - 使用基于存储的数据复制的群集, 98
 - 根文件系统, 262
 - 群集文件, 261
 - 群集节点

- 从统一归档文件, 194
- 使用 `scinstall`, 194
- 回送挂载
 - 将文件系统导出到区域群集, 236

I

IPMP

- 检查状态, 39
- 管理, 188

iSCSI 存储

- 用作法定设备
 - 与基于探测的 IPMP, 188
 - 与基于链路的 IPMP, 188

J

基于存储的复制设备

- 管理, 100
- 配置, 100

基于存储的数据复制, 95

- 和法定设备, 98
- 定义, 94
- 恢复, 98
- 最佳做法, 98
- 要求, 97
- 限制, 97

基于角色的访问控制 见 RBAC

基于主机的数据复制

- 定义, 93
- 示例, 271

监视

- 共享磁盘路径, 150
- 磁盘路径, 146

监视群集

- 使用 Oracle Solaris Cluster Manager 拓扑, 269

检查

- 全局挂载点, 54, 144

检验

- `vfstab` 配置, 142
- 数据复制配置, 301
- 群集节点状态, 197

交换机, 传输, 180

角色

添加定制角色, 61

- 添加角色, 61
- 设置, 59

节点

- 主, 100, 128
- 关闭, 79
- 删除
 - 从全局群集中, 199
 - 从区域群集中, 198
 - 从设备组中, 124
 - 错误消息, 204
- 在全局群集中重命名, 214
- 在区域群集中重命名, 214
- 引导, 79
- 查找 ID, 207
- 添加, 191
- 置于维护状态, 216
- 辅助, 128
- 连接到, 30
- 配置负载限制, 230
- 验证, 208
- 禁用传输电缆, 183
- 镜像, 联机备份, 259
- 卷 见 基于存储的复制

K

/kernel/drv/

- `md.conf` 文件, 118

控制台

- 连接到, 30

快照 见 基于存储的复制

- 时间点, 273

L

来宾域, 82

列出

- 法定设备, 167
- 设备组配置, 132

逻辑主机名资源

- 在数据复制接管中的作用, 275

labeled 标记区域群集, 23

lofi 文件

- 卸载, 222

M

名称空间

- 全局, 99
- 迁移, 115

命令, 26

- boot, 68
- cconsole, 30
- claccess, 26
- cldevice, 26
- cldevicegroup, 26
- clinterconnect, 26
- clnasdevice, 26
- clquorum, 26
- clreslogicalhostname, 26
- clresource, 26
- clresourcegroup, 26
- clresourcetype, 26
- clressharedaddress, 26
- clsetup, 26
- clsnmphot, 26
- clsnmpmib, 26
- clsnmpuser, 26
- cltelemetryattribute, 26
- cluster check, 26, 29, 49, 54
- cluster shutdown, 65
- clzonecluster, 26, 65
- clzonecluster boot, 68
- clzonecluster verify, 49
- metaset, 99

命令行管理工具, 25

命名约定

- 原始磁盘设备, 141
- 复制资源组, 275

md.tab 文件, 29

metaset 命令, 99

MIB

- 启用和禁用 SNMP 事件, 224, 224
- 更改 SNMP 事件协议, 225

N

network-bind-address

- 检查, 268

ntp.conf.sc 文件, 213

numsecondaries 属性, 130

O

OpenBoot PROM (OBP), 211

Oracle GlassFish

- 更新限制, 265

Oracle Solaris Cluster 支持用作法定设备

- 支持用作法定设备, 155

Oracle Solaris Cluster Manager, 26, 265

- 使用, 265

可以执行的任务

- 关闭区域群集, 25
- 创建区域群集, 24

如何登录, 265

您可以执行的任务

- 从区域群集中删除存储设备, 239
- 从区域群集中删除文件系统, 237
- 从区域群集节点卸载软件, 25
- 使设备组联机, 113, 133
- 使设备组脱机, 113, 134
- 关闭全局群集节点, 80
- 关闭区域群集节点, 80, 199
- 创建区域群集, 231
- 创建法定服务器, 158
- 创建法定设备, 156
- 删除专用适配器, 180
- 删除传输适配器, 180
- 删除区域群集, 235
- 删除文件系统, 143
- 删除法定设备, 161
- 删除电缆, 180
- 向区域群集添加共享存储, 192
- 向区域群集添加网络地址, 234
- 启用法定设备, 166
- 启用电缆, 182
- 启用磁盘路径监视, 146
- 启用群集互连, 187
- 在全局群集节点上创建负载限制, 230
- 在区域群集节点上创建负载限制, 230
- 复位法定设备, 165
- 将存储设备添加到区域群集, 231
- 将文件系统添加到区域群集, 231
- 将节点置于维护状态, 217
- 引导区域群集节点, 83
- 查看区域配置, 25

- 查看法定信息, 167
- 查看群集配置, 40
- 查看节点的状态, 39
- 查看节点系统消息, 30, 31
- 查看资源和资源组, 34
- 检查群集互连的状态, 177
- 检查群集组件状态, 36
- 添加 Oracle ZFS Storage Appliance NAS 设备, 157
- 添加专用适配器, 178
- 添加传输适配器, 178
- 添加共享存储, 206
- 添加文件系统, 140
- 添加电缆, 178
- 清除节点, 65, 216
- 禁用法定设备, 165
- 禁用电缆, 183
- 禁用磁盘路径监视, 147
- 编辑区域群集的 "Resource Security" (资源安全性) 属性, 231
- 编辑节点属性, 151
- 访问群集配置实用程序, 31
- 重新引导区域群集节点, 86
- 拓扑, 265
- 拓扑视图, 269
- 故障排除, 265
- 更新限制, 265
- 设置辅助功能控件, 269
- 访问, 266
- Oracle Solaris OS
 - CPU 控制, 247
 - svcadm 命令, 212
 - 全局群集定义, 23
 - 全局群集管理任务, 24
 - 区域群集定义, 23
 - 基于主机的复制, 94
 - 基于存储的复制, 94
 - 特殊说明
 - 引导节点, 82
 - 重新引导节点, 86
- Oracle ZFS Storage Appliance
 - 支持用作法定设备, 155
 - 添加为法定设备, 157

P

配置

- 基于存储的复制设备, 100
- 复制的 ZFS 设备组, 121
- 安全密钥, 267
- 数据复制, 271
- 无 HAStoragePlus 的本地 ZFS 存储池, 122
- 节点上的负载限制, 230
- 配置示例 (校园群集)
 - 双工作间, 基于存储的数据复制, 95
- 配置文件
 - RBAC 权限, 59
- patchadd
 - 修补 solaris10 标记区域群集, 256
- pconsole 实用程序, 30
 - 使用, 196
 - 安全连接, 31

Q

启动

- pconsole 实用程序, 196
- 全局群集, 68
- 全局群集节点, 79
- 区域群集节点, 79
- 启用传输电缆, 182
- 启用和禁用 SNMP 事件 MIB, 224, 224
- 迁移
 - 全局设备名称空间, 115
- 切换
 - 设备组的主节点, 133
- 切换回原状态
 - 数据复制中的执行准则, 279
- 清单
 - 自动化安装程序, 195
- 区域路径
 - 移动, 231
- 区域群集
 - solaris 标记
 - 更新, 255
 - solaris10 标记
 - 使用 clzonecluster 修补, 255
 - 使用 patchadd 修补, 256
 - 修补备用 BE, 256
 - 为应用程序做准备, 231

- 从统一归档文件安装, 233
 - 从统一归档文件配置, 232
 - 克隆, 231
 - 关闭, 65
 - 创建, 23
 - 删除文件系统, 231
 - 定义, 23
 - 引导, 65
 - 支持的直接挂载, 236
 - 查看配置信息, 47
 - 标记, 23
 - 添加网络地址, 234
 - 移动区域路径, 231
 - 管理, 205
 - 组件状态, 36
 - 重新引导, 72
 - 验证配置, 49
 - 区域群集节点
 - 关闭, 79
 - 引导, 79
 - 指定 IP 地址和 NIC, 191
 - 重新引导, 86
 - 取消监视
 - 磁盘路径, 147
 - 取消注册
 - Solaris Volume Manager 设备组, 124
 - 全局
 - 名称空间, 99, 114
 - 挂载点, 检查, 54, 144
 - 设备, 99
 - 动态重新配置, 100
 - 设置权限, 100
 - 全局群集
 - 关闭, 65
 - 删除节点, 199
 - 定义, 23
 - 引导, 65
 - 查看配置信息, 40
 - 管理, 205
 - 组件状态, 36
 - 重新引导, 72
 - 验证配置, 49
 - 全局群集节点
 - CPU 份额, 248
 - 关闭, 79
 - 引导, 79
 - 检验
 - 状态, 197
 - 重新引导, 86
 - 全局设备名称空间
 - 迁移, 115
 - 权限, 全局设备, 100
 - 权限配置文件
 - RBAC, 59
 - 群集
 - 备份, 29, 259
 - 恢复文件, 261
 - 更改名称, 206
 - 节点验证, 208
 - 范围, 32
 - 设置时间, 209
 - 群集互连
 - 动态重新配置, 176
 - 管理, 175
 - 群集文件系统, 99
 - 删除, 142
 - 挂载选项, 141
 - 检验配置, 142
 - 添加, 140
 - 管理, 112
 - 群集文件系统的挂载选项
 - 要求, 141
- ## R
- 容灾性
 - 定义, 271
 - 软件包
 - 卸载, 257
 - 更新特定软件包, 253
 - 软件更新
 - 概述, 251
 - RBAC, 59
 - 任务
 - 使用, 59
 - 修改用户, 62
 - 添加定制角色, 61
 - 添加角色, 61
 - 设置, 59
 - 权限配置文件 (说明), 59

S

删除

SNMP

- 主机, 227

- 用户, 228

- Solaris Volume Manager 设备组, 124

- 从区域群集中, 198

- 传输电缆, 适配器, 和交换机, 180

- 区域群集中的资源和资源组, 236

- 存储阵列, 202

- 最后一个法定设备, 161

- 法定设备, 154, 160

- 群集文件系统, 142

- 节点, 198, 199

- 节点, 从所有设备组中, 124

设备

- 全局, 99

设备组

- Solaris Volume Manager

- 添加, 118

- 主所有权, 128

- 列出配置, 132

- 删除

- 和取消注册, 124

- 原始磁盘

- 添加, 120

- 更改属性, 128

- 添加, 120

- 管理概述, 112

- 维护状态, 134

- 针对数据复制进行配置, 283

- 设备组的主节点切换, 133

- 设备组的主所有权, 128

设置

- 角色 (RBAC), 59

- 设置群集时间, 209

升级

- 标记类型为 solaris 的故障转移区域, 251

- 标记类型为 solaris10 的故障转移区域, 251

- 概述, 251

时间点快照

- 定义, 273

- 执行, 300

使用

- 角色 (RBAC), 59

示例

- 列出交互式验证检查, 51

- 运行功能验证检查, 52

- 示例配置 (校园群集)

- 双工作间, 基于存储的复制, 95

- 事件 MIB

- 启用和禁用 SNMP, 224, 224

- 更改 log_number, 225

- 更改 min_severity, 225

- 更改 SNMP 协议, 225

- 适配器, 传输, 180

- 手动更新 DID 信息, 148

- 属性

- failback, 128

- numsecondaries, 130

- preferenced, 128

- 数据复制, 93

- 介绍, 271

- 准则

- 管理切换, 278

- 管理接管, 278

- 配置资源组, 274

- 同步, 272

- 启用, 295

- 基于主机, 93, 271

- 基于存储的, 94, 95

- 定义, 93

- 异步, 272

- 必需的硬件和软件, 281

- 时间点快照, 273, 300

- 更新 DNS 条目, 304

- 检验配置, 301

- 示例, 298

- 管理接管, 304

- 资源组

- 共享地址, 277

- 创建, 287

- 可伸缩应用程序, 277

- 命名约定, 275

- 应用程序, 275

- 故障转移应用程序, 276

- 配置, 275

- 远程镜像, 272, 299

- 配置

- NFS 应用程序资源组, 291

- 为 NFS 应用程序配置文件系统, 285

- 关联性切换, 275, 288

- 设备组, 283
- 配置示例, 280
- SATA 存储, 156
 - 支持用作法定设备, 155
- scinstall 命令
 - 恢复群集节点, 194
- showrev -p 命令, 32
- SNMP
 - 删除用户, 228
 - 启用主机, 226
 - 启用和禁用事件 MIB, 224, 224
 - 更改协议, 225
 - 添加用户, 228
 - 禁用主机, 227
- solaris 标记区域群集, 23
 - 更新, 255
- Solaris OS 见 Oracle Solaris OS
- Solaris Volume Manager
 - 原始磁盘设备名称, 141
- solaris10 标记区域群集, 23
 - 使用 clzonecluster 修补, 255
 - 使用 patchadd 修补, 256
 - 修补备用 BE, 256
- SRDF 见 EMC SRDF
- ssh, 31

T

- 拓扑
 - 在 Oracle Solaris Cluster Manager 中使用, 269
- 特性 见 属性
- 替换法定设备, 163
- 添加
 - SNMP
 - 主机, 226
 - 用户, 228
 - Solaris Volume Manager 设备组, 120
 - ZFS
 - 复制的设备组, 121
 - 无 HAStoragePlus 的存储池, 122
 - 传输电缆, 适配器和交换机, 37, 40, 175
 - 定制角色 (RBAC), 61
 - 法定设备, 156
 - Oracle ZFS Storage Appliance NAS, 157
 - 法定服务器, 158

- 直接连接的共享磁盘法定设备, 156
- 网络地址到区域群集, 234
- 群集文件系统, 140
- 节点, 191
 - 向全局群集, 192
 - 向区域群集, 192
- 角色 (RBAC), 61
- 设备组, 118, 120
- 停止
 - 全局群集, 72
 - 全局群集节点, 79
 - 区域群集, 72
 - 区域群集节点, 79
- 停止服务状态
 - 法定设备, 164
- 同步数据复制, 96, 272
- 统一归档文件
 - 安装区域群集, 233
 - 恢复群集节点, 194
 - 配置区域群集, 232

V

- /var/adm/messages 文件, 91
- vfstab 文件
 - 检验配置, 142
 - 添加挂载点, 141

W

- 网络地址
 - 向区域群集中添加, 234
- 网络文件系统 (Network File System, NFS)
 - 为数据复制配置应用程序文件系统, 285
 - 主群集上的资源, 292
- 维护
 - 法定设备, 164
- 维护状态
 - 使法定设备脱离, 166
 - 将法定设备置于, 164
 - 节点, 216
- 位图
 - 时间点快照, 273
 - 远程镜像复制, 272
- 文件

- /etc/vfstab, 54
 - md.conf, 118
 - md.tab, 29
 - ntp.conf.sc, 213
- 文件系统
 - NFS 应用程序
 - 为数据复制配置, 285
 - 在区域群集中删除, 231
 - 恢复根
 - 描述, 262
- X
- 显示
 - 故障磁盘路径, 148
- 校园群集
 - 基于存储的数据复制, 95
 - 基于存储的数据复制的恢复, 98
- 卸载
 - lofi 设备文件, 222
 - Oracle Solaris Cluster 软件, 220
 - 软件包, 257
- 修补, 254
 - 参见 更新
- solaris10 标记区域群集
 - 使用 clzonecluster, 255
 - 使用 patchadd, 256
 - 备用 BE, 256
- 修复
 - 法定设备, 168
- 修复已满的 /var/adm/messages 文件, 91
- 修改
 - 法定设备节点列表, 164
 - 用户 (RBAC), 62
- Y
- 验证
 - 全局群集配置, 49
 - 区域群集配置, 49
- 异步数据复制, 96, 272
- 引导
 - 全局群集, 65
 - 全局群集节点, 79
 - 区域群集, 65, 68
 - 区域群集节点, 79
 - 非群集模式, 89
- 应用程序资源组
 - 为数据复制配置, 291
 - 准则, 275
- 用户
 - 修改属性, 62
 - 删除 SNMP, 228
 - 添加 SNMP, 228
- 用于数据复制的共享地址资源组, 277
- 用于数据复制的故障转移应用程序
 - AffinityOn 属性, 275
 - GlobalDevicePaths, 275
 - ZPoolsSearchDir, 275
 - 准则
 - 管理接管, 278
 - 资源组, 276
 - 管理, 304
- 用于数据复制的可伸缩应用程序, 277
- 用于数据复制的扩展属性
 - 复制资源, 289
 - 应用程序资源, 291, 294
- 用于数据复制的切换
 - 关联性切换, 275
 - 执行, 304
- 与群集控制台的安全连接, 31
- 域名系统 (Domain Name System, DNS)
 - 在数据复制中更新, 304
 - 更新准则, 278
- 原始磁盘设备
 - 命名约定, 141
- 原始磁盘设备组
 - 添加, 120
- 远程登录, 30
- 远程复制 见 基于存储的复制
- 远程镜像 见 基于存储的复制
- 远程镜像复制
 - 定义, 272
 - 执行, 299
- Z
- 支持的法定设备类型, 155
- 直接挂载
 - 将文件系统导出到区域群集, 236
- 直接连接的共享磁盘法定设备

- 添加, 156
- 主机
 - 添加和删除 SNMP, 226, 227
- 专用主机名
 - 更改, 212
- 状态
 - 全局群集组件, 36
 - 区域群集组件, 36
- 资源
 - 删除, 236
 - 显示配置信息, 35
- 资源组
 - 数据复制
 - 在故障转移中的角色, 275
 - 配置, 275
 - 配置准则, 274
- 自动化安装程序
 - 清单, 195
- 最后一个法定设备
 - 删除, 161
- 最佳做法, 98
 - EMC SRDF, 98
 - 基于存储的数据复制, 98
- ZFS
 - 删除文件系统, 236
 - 复制, 121
 - 根文件系统的限制, 112
 - 添加
 - 复制的设备组, 121
 - 无 HAStoragePlus 的存储池, 122
 - 配置
 - 复制的设备组, 121
 - 无 HAStoragePlus 的本地存储池, 122
- ZFS Storage Appliance 见 法定设备
- zpool
 - 配置无 HAStoragePlus 的本地 ZFS 存储池, 122
 - 配置无 HAStoragePlus 的本地存储池, 121
- ZPoolsSearchDir
 - 用于数据复制的扩展属性, 275

