

Guía de administración del sistema de Oracle Solaris Cluster

Copyright © 2000, 2010, Oracle y/o sus subsidiarias. Todos los derechos reservados.

Este software y la documentación relacionada están sujetos a un contrato de licencia que incluye restricciones de uso y revelación, y se encuentran protegidos por la legislación sobre la propiedad intelectual. A menos que figure explícitamente en el contrato de licencia o esté permitido por la ley, no se podrá utilizar, copiar, reproducir, traducir, emitir, modificar, conceder licencias, transmitir, distribuir, exhibir, representar, publicar ni mostrar ninguna parte, de ninguna forma, por ningún medio. Queda prohibida la ingeniería inversa, desensamblaje o descompilación de este software, excepto en la medida en que sean necesarios para conseguir interoperabilidad según lo especificado por la legislación aplicable.

La información contenida en este documento puede someterse a modificaciones sin previo aviso y no se garantiza que se encuentre exenta de errores. Si detecta algún error, le agradeceremos que nos lo comuniqué por escrito.

Si este software o la documentación relacionada se entrega al Gobierno de EE.UU. o a cualquier entidad que adquiera licencias en nombre del Gobierno de EE.UU. se aplicará la siguiente disposición:

U.S. GOVERNMENT RIGHTS Programs, software, databases, and related documentation and technical data delivered to U.S. Government customers are "commercial computer software" or "commercial technical data" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, the use, duplication, disclosure, modification, and adaptation shall be subject to the restrictions and license terms set forth in the applicable Government contract, and, to the extent applicable by the terms of the Government contract, the additional rights set forth in FAR 52.227-19, Commercial Computer Software License (December 2007). Oracle America, Inc., 500 Oracle Parkway, Redwood City, CA 94065

Este software o hardware se ha desarrollado para uso general en diversas aplicaciones de gestión de la información. No se ha diseñado ni está destinado para utilizarse en aplicaciones de riesgo inherente, incluidas las aplicaciones que pueden causar daños personales. Si utiliza este software o hardware en aplicaciones de riesgo, usted será responsable de tomar todas las medidas apropiadas de prevención de fallos, copia de seguridad, redundancia o de cualquier otro tipo para garantizar la seguridad en el uso de este software o hardware. Oracle Corporation y sus subsidiarias declinan toda responsabilidad derivada de los daños causados por el uso de este software o hardware en aplicaciones de riesgo.

Oracle y Java son marcas comerciales registradas de Oracle y/o sus subsidiarias. Todos los demás nombres pueden ser marcas comerciales de sus respectivos propietarios.

AMD, Opteron, el logotipo de AMD y el logotipo de AMD Opteron son marcas comerciales o marcas comerciales registradas de Advanced Micro Devices. Intel e Intel Xeon son marcas comerciales o marcas comerciales registradas de Intel Corporation. Todas las marcas comerciales de SPARC se utilizan con licencia y son marcas comerciales o marcas comerciales registradas de SPARC International, Inc. UNIX es una marca comercial registrada con acuerdo de licencia de X/Open Company, Ltd.

Este software o hardware y la documentación pueden ofrecer acceso a contenidos, productos o servicios de terceros o información sobre los mismos. Ni Oracle Corporation ni sus subsidiarias serán responsables de ofrecer cualquier tipo de garantía sobre el contenido, los productos o los servicios de terceros y renuncian explícitamente a ello. Oracle Corporation y sus subsidiarias no se harán responsables de las pérdidas, los costos o los daños en los que se incurra como consecuencia del acceso o el uso de contenidos, productos o servicios de terceros.

Contenido

Prefacio	11
1 Introducción a la administración de Oracle Solaris Cluster	17
Introducción a la administración de Oracle Solaris Cluster	18
Trabajo con un clúster de zona	18
Restricciones de las funciones del sistema operativo Oracle Solaris	19
Herramientas de administración	20
Interfaz gráfica de usuario	20
Interfaz de línea de comandos	20
Preparación para administrar el clúster	22
Documentación de las configuraciones de hardware de Oracle Solaris Cluster	22
Uso de la consola de administración	22
Copia de seguridad del clúster	23
Procedimiento para empezar a administrar el clúster	24
▼ Inicio de sesión remota en el clúster	25
▼ Conexión segura a las consolas del clúster	27
▼ Obtención de acceso a las utilidades de configuración del clúster	27
▼ Visualización de la información de parches de Oracle Solaris Cluster	28
▼ Visualización de la información de versión de Oracle Solaris Cluster	29
▼ Visualización de tipos de recursos configurados, grupos de recursos y recursos	31
▼ Comprobación del estado de los componentes del clúster	32
▼ Comprobación del estado de la red pública	35
▼ Visualización de la configuración del clúster	36
▼ Validación de una configuración básica de clúster	45
▼ Comprobación de los puntos de montaje globales	47
▼ Visualización del contenido de los registros de comandos de Oracle Solaris Cluster	48

2	Oracle Solaris Cluster y RBAC	51
	Instalación y uso de RBAC con Oracle Solaris Cluster	51
	Perfiles de derechos de RBAC en Oracle Solaris Cluster	52
	Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster	53
	▼ Creación de una función con la herramienta Administrative Roles	53
	▼ Creación de una función desde la línea de comandos	55
	Modificación de las propiedades de RBAC de un usuario	57
	▼ Modificación de las propiedades de RBAC de un usuario con la herramienta User Accounts	57
	▼ Modificación de las propiedades de RBAC de un usuario desde la línea de comandos	58
3	Cierre y arranque de un clúster	59
	Descripción general sobre el cierre y el arranque de un clúster	59
	▼ Cierre de un clúster	61
	▼ Arranque de un clúster	63
	▼ Rearranque de un clúster	65
	Cierre y arranque de un solo nodo de un clúster	69
	▼ Cierre de un nodo	70
	▼ Arranque de un nodo	73
	▼ Rearranque de un nodo	76
	▼ Rearranque de un nodo en un modo que no sea de clúster	79
	Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad	82
	▼ Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad	83
4	Métodos de replicación de datos	85
	Comprensión de la replicación de datos	86
	Métodos admitidos de replicación de datos	87
	Uso de replicación de datos basada en almacenamiento dentro de un clúster	88
	Requisitos y restricciones al usar replicación de datos basada en almacenamiento dentro de un clúster	90
	Problemas de recuperación manual al usar una replicación de datos basada en almacenamiento dentro de un clúster	91
	Prácticas recomendadas al usar la replicación de datos basada en almacenamiento	91

5 Administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de clúster	93
Información general sobre la administración de dispositivos globales y el espacio de nombre global	93
Permisos de dispositivos globales en Solaris Volume Manager	94
Reconfiguración dinámica con dispositivos globales	94
Aspectos que tener en cuenta en la administración de Veritas Volume Manager	96
Administración de dispositivos replicados basados en almacenamiento	97
Administración de dispositivos replicados mediante Hitachi TrueCopy	97
Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility	109
Información general sobre la administración de sistemas de archivos de clústers	121
Restricciones del sistema de archivos de clúster	121
Directrices para admitir VxFS	122
Administración de grupos de dispositivos	123
▼ Actualización del espacio de nombre de dispositivos globales	126
▼ Cómo cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales	127
Migración del espacio de nombre de dispositivos globales	128
▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de <code>lofi</code>	129
▼ Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada	130
Adición y registro de grupos de dispositivos	132
▼ Adición y registro de un grupo de dispositivos (Solaris Volume Manager)	132
▼ Adición y registro de un grupo de dispositivos (disco básico)	134
▼ Adición y registro de un grupo de dispositivos replicado (ZFS)	135
▼ Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager) ..	137
Mantenimiento de grupos de dispositivos	137
Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)	138
▼ Eliminación de un nodo de todos los grupos de dispositivos	138
▼ Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)	139
▼ Creación de un nuevo grupo de discos al encapsular discos (Veritas Volume Manager) .	141
▼ Adición de un volumen nuevo a un grupo de dispositivos ya creado (Veritas Volume Manager)	143
▼ Conversión de un grupo de discos en un grupo de dispositivos (Veritas Volume Manager)	144

▼ Asignación de un número menor nuevo a un grupo de dispositivos (Veritas Volume Manager)	144
▼ Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)	145
▼ Registro de modificaciones en la configuración de grupos de discos (Veritas Volume Manager)	148
▼ Conversión de un grupo de discos local en un grupo de dispositivos (VxVM)	149
▼ Conversión de un grupo de dispositivos en un grupo de discos local (VxVM)	150
▼ Eliminación de un volumen de un grupo de dispositivos (Veritas Volume Manager)	151
▼ Eliminación y anulación del registro de un grupo de dispositivos (Veritas Volume Manager)	152
▼ Adición de un nodo a un grupo de dispositivos (Veritas Volume Manager)	153
▼ Eliminación de un nodo de un grupo de dispositivos (Veritas Volume Manager)	155
▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos	157
▼ Cambio de propiedades de los grupos de dispositivos	158
▼ Establecimiento del número de secundarios para un grupo de dispositivos	160
▼ Enumeración de la configuración de grupos de dispositivos	163
▼ Conmutación al nodo primario de un grupo de dispositivos	165
▼ Colocación de un grupo de dispositivos en estado de mantenimiento	166
Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento	168
▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento	168
▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento	169
▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento	170
▼ Modificación del protocolo de protección en un solo dispositivo de almacenamiento	172
Administración de sistemas de archivos de clúster	174
▼ Adición de un sistema de archivos de clúster	174
▼ Eliminación de un sistema de archivos de clúster	178
▼ Comprobación de montajes globales en un clúster	180
Administración de la supervisión de rutas de disco	180
▼ Supervisión de una ruta de disco	181
▼ Anulación de la supervisión de una ruta de disco	183
▼ Impresión de rutas de disco erróneas	184
▼ Resolución de un error de estado de ruta de disco	184
▼ Supervisión de rutas de disco desde un archivo	185
▼ Habilitación del rearranque automático de un nodo si fallan todas las rutas de disco	

compartido supervisadas	187
▼ Inhabilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	188
6 Administración de quórum	189
Administración de dispositivos de quórum	189
Reconfiguración dinámica con dispositivos de quórum	191
Adición de un dispositivo de quórum	192
Eliminación o sustitución de un dispositivo de quórum	203
Mantenimiento de dispositivos de quórum	207
Administración de servidores de quórum de Oracle Solaris Cluster	215
Información general sobre el archivo de configuración del servidor de quórum	215
Inicio y detención del software del servidor del quórum	216
▼ Inicio de un servidor de quórum	216
▼ Detención de un servidor de quórum	217
Visualización de información sobre el servidor de quórum	218
Limpieza de la información caducada sobre clústers del servidor de quórum	220
7 Administración de interconexiones de clústers y redes públicas	223
Administración de interconexiones de clústers	223
Reconfiguración dinámica con interconexiones de clústers	225
▼ Comprobación del estado de la interconexión de clúster	225
▼ Adición de dispositivos de cable de transporte de clúster, adaptadores o conmutadores de transporte	226
▼ Eliminación de cable de transporte de clúster, adaptadores de transporte y conmutadores de transporte	229
▼ Habilitación de un cable de transporte de clúster	232
▼ Inhabilitación de un cable de transporte de clúster	233
▼ Determinación del número de instancia de un adaptador de transporte	235
▼ Modificación de la dirección de red privada o del intervalo de direcciones de un clúster	236
Administración de redes públicas	238
Administración de grupos de varias rutas de red IP en un clúster	239
Reconfiguración dinámica con interfaces de red pública	240

8 Adición y eliminación de un nodo	243
Adición de nodos a un clúster	243
▼ Adición de un nodo a la lista de nodos autorizados	244
Creación de un nodo sin votación (zona) en un clúster global	246
Eliminación de nodos de un clúster	249
▼ Eliminación de un nodo de un clúster de zona	250
▼ Eliminación de un nodo de la configuración de software del clúster	251
▼ Eliminación de un nodo sin votación (zona) de un clúster global	254
▼ Eliminación de la conectividad entre una matriz y un único nodo, en un clúster con conectividad superior a dos nodos	255
▼ Corrección de mensajes de error	257
9 Administración del clúster	259
Información general sobre la administración del clúster	259
▼ Cambio del nombre del clúster	260
▼ Asignación de un ID de nodo a un nombre de nodo	262
▼ Uso de la autenticación del nodo del clúster nuevo	262
▼ Restablecimiento de la hora del día en un clúster	264
▼ SPARC: Visualización de la PROM OpenBoot en un nodo	266
▼ Cambio del nombre de host privado de nodo	267
▼ Adición un nombre de host privado a un nodo sin votación en un clúster global	270
▼ Cambio del nombre de host privado en un nodo sin votación de un clúster global	270
▼ Eliminación del nombre de host privado para un nodo sin votación en un clúster global	272
▼ Cómo cambiar el nombre de un nodo	272
▼ Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes	273
▼ Colocación de un nodo en estado de mantenimiento	274
▼ Procedimiento para sacar un nodo del estado de mantenimiento	276
Configuración de los límites de carga	278
Ejecución de tareas administrativas del clúster de zona	281
▼ Eliminación de un clúster de zona	282
▼ Eliminación de un sistema de archivos desde un clúster de zona	283
▼ Eliminación de un dispositivo de almacenamiento desde un clúster de zona	285
▼ Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster	287
Solución de problemas de la desinstalación de nodos	289
Creación, configuración y administración de MIB de eventos de SNMP de Oracle Solaris	

Cluster	291
Solución de problemas	297
Ejecución de una aplicación fuera del clúster global	297
Restauración de un conjunto de discos dañado	299
10 Configuración del control del uso de la CPU	303
Introducción al control de la CPU	303
Elección de una situación hipotética	303
Planificador por reparto equitativo	304
Configuración del control de la CPU	305
▼ Control del uso de la CPU en el nodo de votación de un clúster global	305
▼ Control del uso de la CPU en un nodo sin votación de clúster global con el conjunto de procesadores predeterminado	307
▼ Control del uso de la CPU en un nodo sin votación de clúster global con un conjunto de procesadores dedicado	310
11 Aplicación de parches en el software y el firmware de Oracle Solaris Cluster	315
Información general sobre aplicación de parches de Oracle Solaris Cluster	315
Sugerencias sobre los parches de Oracle Solaris Cluster	316
Aplicación de parches del software Oracle Solaris Cluster	317
▼ Aplicación de un parche de rearranque (nodo)	317
▼ Aplicación de un parche de rearranque (clúster)	322
▼ Aplicación de un parche de Oracle Solaris Cluster que no sea de rearranque	325
▼ Aplicación de parches en modo monousuario con nodos de migración tras error	326
Cambio de un parche de Oracle Solaris Cluster	330
12 Copias de seguridad y restauraciones de clústers	333
Copia de seguridad de un clúster	333
▼ Búsqueda de los nombres de sistema de archivos para hacer copias de seguridad	334
▼ Determinación del número de cintas necesarias para una copia de seguridad completa	334
▼ Copias de seguridad del sistema de archivos root (/)	335
▼ Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)	338
▼ Copias de seguridad en línea para volúmenes (Veritas Volume Manager)	341
▼ Copias de seguridad de la configuración del clúster	345
Restauración de los archivos del clúster	345

▼ Restauración de archivos individuales de forma interactiva (Solaris Volume Manager) .	346
▼ Restauración del sistema de archivos root (/) (Solaris Volume Manager)	346
▼ Cómo restaurar un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager	349
▼ Restauración de un sistema de archivos root no encapsulado (/) (Veritas Volume Manager)	354
▼ Restauración de un sistema de archivos root encapsulado (/) (Veritas Volume Manager)	356
13 Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario	361
Información general sobre Oracle Solaris Cluster Manager	361
SPARC: Información general sobre Sun Management Center	362
Configuración de Oracle Solaris Cluster Manager	363
Configuración de funciones de RBAC	363
▼ Uso de Common Agent Container para cambiar los números de puerto de los agentes de servicios o de administración	364
▼ Cambio de la dirección del servidor para Oracle Solaris Cluster Manager	365
▼ Regeneración de las claves de seguridad de Common Agent Container	365
Inicio del software Oracle Solaris Cluster Manager	366
▼ Inicio de Oracle Solaris Cluster Manager	367
A Ejemplo	371
Configuración de replicación de datos basada en host con el software Sun StorageTek Availability Suite	371
Comprensión del software Sun StorageTek Availability Suite en un clúster	372
Directrices para la configuración de replicación de datos basada en host entre clústers ..	375
Mapa de tareas: ejemplo de configuración de una replicación de datos	380
Conexión e instalación de clústers	381
Ejemplo de configuración de grupos de dispositivos y de recursos	383
Ejemplo de cómo habilitar la replicación de datos	397
Ejemplo de cómo efectuar una replicación de datos	400
Ejemplo de cómo controlar una migración tras error	406
Índice	409

Prefacio

El documento *Guía de administración del sistema de Oracle Solaris Cluster* proporciona procedimientos para administrar una configuración de Oracle Solaris Cluster en los sistemas de SPARC y x86.

Nota – Esta versión de Oracle Solaris Cluster admite sistemas que usen arquitecturas de las familias de procesadores SPARC y x86: UltraSPARC, SPARC64, AMD64 e Intel 64. En este documento, x86 hace referencia a la familia más amplia de productos compatibles con 64 y 32 bits. La información de este documento se aplica a todas las plataformas a menos que se especifique lo contrario.

Este documento está destinado a administradores de sistemas con amplios conocimientos de software y hardware de Oracle. Este documento no se puede usar como una guía de planificación ni de preventas.

Las instrucciones de este documento presuponen un conocimiento del sistema operativo Oracle Solaris y el dominio del software del administrador de volúmenes que se utiliza con Oracle Solaris Cluster.

Uso de los comandos de UNIX

Este documento contiene información sobre comandos específica para administrar una configuración de Oracle Solaris Cluster. Este documento podría no contener información completa sobre los comandos y procedimientos básicos de UNIX.

Consulte al menos una de las fuentes de información siguientes:

- Documentación en línea del software de Oracle Solaris
- Otra documentación de software recibida con el sistema
- Páginas de comando man del sistema operativo Oracle Solaris

Convenciones tipográficas

La siguiente tabla describe las convenciones tipográficas utilizadas en este manual.

TABLA P-1 Convenciones tipográficas

Tipo de letra	Significado	Ejemplo
AaBbCc123	Los nombres de comandos, archivos y directorios, así como la salida del equipo en pantalla.	Edite el archivo <code>.login</code> . Utilice el comando <code>ls -a</code> para mostrar todos los archivos. <code>machine_name%</code> tiene correo.
AaBbCc123	Lo que se escribe en contraposición con la salida del equipo en pantalla.	<code>machine_name% su</code> Contraseña:
<i>aabbcc123</i>	Marcador de posición: debe sustituirse por un valor o nombre real.	El comando necesario para eliminar un archivo es <code>rm filename</code> .
<i>AaBbCc123</i>	Títulos de manuales, términos nuevos y palabras destacables.	Consulte el capítulo 6 de la <i>Guía del usuario</i> . Una copia en <i>caché</i> es la que se almacena localmente. <i>No</i> guarde el archivo. Nota: algunos elementos destacados aparecen en negrita en línea.

Indicadores de los shells en los ejemplos de comandos

La tabla siguiente muestra los indicadores predeterminados de sistema UNIX y de superusuario para los shells que se incluyen en el sistema operativo Oracle Solaris. Tenga en cuenta que el indicador del sistema predeterminado que se visualiza en los ejemplos de comando varía en función de la versión de Oracle Solaris.

TABLA P-2 Indicadores del shell

Shell	Indicador
Shell Bash, Shell Korn y Shell Bourne	\$
Shell Bash, Shell Korn y Shell Bourne para superusuario	#
Shell C	machine_name%
Shell C para superusuario	machine_name# (núm_nombre_equipo)

Documentación relacionada

Puede encontrar información sobre temas referentes a Oracle Solaris Cluster en la documentación enumerada en la tabla siguiente. Toda la documentación sobre Oracle Solaris Cluster está disponible en <http://docs.sun.com>.

Tema	Documentación
Introducción	<i>Oracle Solaris Cluster Overview</i>
Conceptos	<i>Oracle Solaris Cluster Concepts Guide</i>
Administración e instalación de software	<i>Oracle Solaris Cluster 3.3 Hardware Administration Manual</i> Guías de administración de hardware individual
Instalación de software	<i>Guía de instalación del software Oracle Solaris Cluster</i>
Administración e instalación de servicio de datos	<i>Oracle Solaris Cluster Data Services Planning and Administration Guide</i> Guías de servicio de datos individual
Desarrollo de servicio de datos	<i>Oracle Solaris Cluster Data Services Developer's Guide</i>
Administración del sistema	<i>Guía de administración del sistema de Oracle Solaris Cluster</i>
Mensajes de error	<i>Oracle Solaris Cluster Error Messages Guide</i>
Referencias de comandos y funciones	<i>Oracle Solaris Cluster Reference Manual</i>

Para obtener una lista completa de la documentación de Oracle Solaris Cluster, consulte las notas de su versión de Oracle Solaris Cluster en <http://docs.sun.com>.

Documentación, asistencia o formación

Consulte los siguientes sitios web para obtener recursos adicionales:

- [Documentation \(http://docs.sun.com\)](http://docs.sun.com)
- [Support \(http://www.oracle.com/us/support/systems/index.html\)](http://www.oracle.com/us/support/systems/index.html)
- [Training \(http://education.oracle.com\)](http://education.oracle.com). Haga clic en el enlace de Sun de la barra de navegación izquierda.

Oracle agradece sus comentarios

Oracle agradece sus comentarios y sugerencias sobre la calidad y la utilidad de la documentación. Si encuentra errores o tiene otras sugerencias de mejora, entre en <http://docs.sun.com> y haga clic en Feedback (Comentarios). Indique el título y el número de referencia de la documentación junto con el capítulo, la sección y el número de la página, si dispone de esos datos. Indíquenos si desea que le respondamos.

La [Oracle Technology Network](http://www.oracle.com/technetwork/index.html) (<http://www.oracle.com/technetwork/index.html>) ofrece una amplia variedad de recursos relacionados con el software de Oracle:

- Si desea abordar problemas técnicos y soluciones, utilice los [Discussion Forums](http://forums.oracle.com) (<http://forums.oracle.com>).
- Consulte tutoriales paso a paso en [Oracle by Example](http://www.oracle.com/technology/obe/start/index.html) (<http://www.oracle.com/technology/obe/start/index.html>).
- Descargue [Sample Code](http://www.oracle.com/technology/sample_code/index.html) (http://www.oracle.com/technology/sample_code/index.html).

Obtención de ayuda

Póngase en contacto con su proveedor de servicios si tiene problemas con la instalación o el uso de Oracle Solaris Cluster. Proporcione a su proveedor de servicios la información que se especifica a continuación.

- El nombre y la dirección de correo electrónico
- El nombre, la dirección y el número de teléfono de su empresa
- Los números de serie y de modelo de sus sistemas
- El número de versión del sistema operativo, por ejemplo Oracle Solaris 10
- El número de versión de Oracle Solaris Cluster, por ejemplo Oracle Solaris Cluster 3.3

Use los comandos siguientes para recopilar información sobre su sistema para el proveedor de servicios:

Comando	Función
<code>prtconf -v</code>	Muestra el tamaño de la memoria del sistema y ofrece información sobre los dispositivos periféricos.
<code>psrinfo -v</code>	Muestra información acerca de los procesadores.
<code>showrev -p</code>	Indica los parches instalados.
<code>SPARC: prtdiag -v</code>	Muestra información de diagnóstico del sistema.

Comando	Función
<code>/usr/cluster/bin/clnode show-rev</code>	Muestra la información de la versión de Oracle Solaris Cluster y de los paquetes

Tenga también disponible el contenido del archivo `/var/adm/messages`.

Introducción a la administración de Oracle Solaris Cluster

Este capítulo ofrece la información siguiente sobre la administración de clústeres globales y de zona. Además, brinda procedimientos para utilizar las herramientas de administración de Oracle Solaris Cluster:

- “Introducción a la administración de Oracle Solaris Cluster” en la página 18
- “Restricciones de las funciones del sistema operativo Oracle Solaris” en la página 19
- “Herramientas de administración” en la página 20
- “Preparación para administrar el clúster” en la página 22
- “Procedimiento para empezar a administrar el clúster” en la página 24

Todos los procedimientos indicados en esta guía son para el uso en el sistema operativo Oracle Solaris 10.

Los clústeres globales se componen sólo de uno o varios nodos de votación de clúster global y, de forma opcional, por cero o más nodos sin votación de clúster global. Un clúster global puede incluir de forma opcional el sistema operativo LINUX o marca nativa, zonas no globales que no son nodos, pero que contienen alta disponibilidad (como recursos). Un clúster de zona precisa de un clúster global. Si desea obtener información general sobre los clústeres de zona, consulte *Oracle Solaris Cluster Concepts Guide*.

Un clúster de zona se compone de uno o más nodos de votación, de marca clúster. Un clúster de zona depende de, y por tanto requiere, un clúster global. Un clúster global no contiene un clúster de zona. No puede configurar un clúster de zona sin un clúster global. Un clúster de zona tiene, como máximo, un nodo de clúster de zona en un equipo. Un nodo de clúster de zona sigue funcionando únicamente si también lo hace el nodo de votación de clúster global en el mismo equipo. Si falla un nodo de votación de clúster global en un equipo, también fallan todos los nodos de clúster de zona de ese equipo.

Introducción a la administración de Oracle Solaris Cluster

El entorno de alta disponibilidad Oracle Solaris Cluster garantiza que los usuarios finales puedan disponer de las aplicaciones fundamentales. La tarea del administrador del sistema consiste en asegurarse de que la configuración de Oracle Solaris Cluster sea estable y operativa.

Antes de comenzar las tareas de administración, familiarícese con el contenido de planificación que se proporciona en *Guía de instalación del software Oracle Solaris Cluster* y en *Oracle Solaris Cluster Concepts Guide*. Si desea obtener instrucciones sobre cómo crear un clúster de zona, consulte “Configuración de un clúster de zona” de *Guía de instalación del software Oracle Solaris Cluster*. La administración de Oracle Solaris Cluster se organiza en tareas, distribuidas en la documentación siguiente.

- Tareas estándar para administrar y mantener el clúster global o el de zona con regularidad a diario. Estas tareas se describen en la presente guía.
- Tareas de servicios de datos como instalación, configuración y modificación de las propiedades. Dichas tareas se describen en *Oracle Solaris Cluster Data Services Planning and Administration Guide*.
- Tareas de servicios como agregar o reparar hardware de almacenamiento o de red. Dichas tareas se describen en *Oracle Solaris Cluster 3.3 Hardware Administration Manual*.

En términos generales, se pueden llevar a cabo las tareas de administración de Oracle Solaris Cluster mientras el clúster está operativo. Si necesita retirar un nodo del clúster o incluso cerrar dicho nodo, puede hacerse mientras el resto de los nodos continúan desarrollando las operaciones del clúster. Salvo que se indique lo contrario, las tareas administrativas de Oracle Solaris Cluster deben efectuarse en el nodo de votación del clúster global. En los procedimientos que necesitan el cierre de todo el clúster, el impacto sobre el sistema se reduce al mínimo si los periodos de inactividad se programan fuera del horario de trabajo. Si piensa cerrar el clúster o un nodo del clúster, notifíquelo con antelación a los usuarios.

Trabajo con un clúster de zona

En un clúster de zona también se pueden ejecutar dos comandos administrativos de Oracle Solaris Cluster: `cluster` y `clnode`. Sin embargo, el ámbito de estos comandos se limita al clúster de zona donde se ejecute dicho comando. Por ejemplo, utilizar el comando `cluster` en el nodo de votación del clúster global recupera toda la información del clúster global de votación y de todos los clústers de zona. Al utilizar el comando `cluster` en un clúster de zona, se recupera la información de ese mismo clúster de zona.

Si el comando `clzonecluster` se utiliza en un nodo de votación, afecta a todos los clústers de zona del clúster global. Los comandos de clústers de zona también afectan a todos los nodos del clúster de zona, aunque un nodo esté inactivo al ejecutarse el comando.

Los clústers de zona admiten la administración delegada de los recursos situados jerárquicamente bajo el control del administrador de grupos de recursos. Así, los administradores de clústers de zona pueden ver las dependencias de dichos clústers de zona que superan los límites de los clústers de zona, pero no modificarlas. El administrador de un nodo de votación es el único facultado para crear, modificar o eliminar dependencias que superen los límites de los clústers de zona.

La lista siguiente contiene las tareas administrativas principales que se realizan en un clúster de zona.

- Creación de un clúster de zona: use el comando `clzonecluster configure` para crear un clúster de zona. Consulte las instrucciones de [“Configuración de un clúster de zona” de Guía de instalación del software Oracle Solaris Cluster](#).
- Inicio y arranque de un clúster de zona: consulte el [Capítulo 3, “Cierre y arranque de un clúster”](#).
- Adición de un nodo a un clúster de zona: consulte el [Capítulo 8, “Adición y eliminación de un nodo”](#).
- Eliminación de un nodo de un clúster de zona: consulte [“Eliminación de un nodo de un clúster de zona” en la página 250](#).
- Visualización de la configuración de un clúster de zona: consulte [“Visualización de la configuración del clúster” en la página 36](#).
- Validación de la configuración de un clúster de zona: consulte [“Validación de una configuración básica de clúster” en la página 45](#).
- Detención de un clúster de zona: consulte el [Capítulo 3, “Cierre y arranque de un clúster”](#).

Restricciones de las funciones del sistema operativo Oracle Solaris

No se deben habilitar ni inhabilitar los siguientes servicios de Oracle Solaris Cluster mediante la interfaz de la Utilidad de gestión de servicios (SMF).

TABLA 1-1 Oracle Solaris ClusterServices

Servicios de Oracle Solaris Cluster	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

TABLA 1-1 Oracle Solaris ClusterServices (Continuación)

Servicios de Oracle Solaris Cluster	FMRI
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

Herramientas de administración

Las tareas administrativas de una configuración de Oracle Solaris Cluster se pueden efectuar en una interfaz gráfica de usuario (GUI) o mediante la línea de comandos. La sección siguiente ofrece información general sobre las herramientas de la GUI y la línea de comandos.

Interfaz gráfica de usuario

El software de Oracle Solaris Cluster admite herramientas de GUI para realizar diversas tareas administrativas en los clústeres. Dichas herramientas de la interfaz gráfica o GUI son Oracle Solaris Cluster Manager; si emplea software Oracle Solaris Cluster en un sistema basado en SPARC, Sun Management Center. Consulte el [Capítulo 13, “Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario”](#), si desea obtener más información o para ver los procedimientos de configuración de Oracle Solaris Cluster Manager y Sun Management Center. Si desea obtener información específica sobre el uso de Oracle Solaris Cluster Manager, consulte la ayuda en línea de la GUI.

Interfaz de línea de comandos

La mayoría de las tareas de administración de Oracle Solaris Cluster se pueden efectuar de forma interactiva con la utilidad `clsetup(1CL)`. Siempre que sea posible, los procedimientos de administración detallados en esta guía emplean la utilidad `clsetup`.

La utilidad `clsetup` permite administrar los elementos siguientes del menú principal.

- Quórum
- Grupos de recursos

- Servicios de datos
- interconexión de clúster
- Grupos de dispositivos y volúmenes
- Nombres de host privados
- Nuevos nodos
- Otras tareas del clúster

En la lista que se muestra a continuación figuran otros comandos utilizados para la administración de una configuración de Oracle Solaris Cluster. Consulte las páginas de comando man si desea obtener información más detallada.

<code>ccp(1M)</code>	Inicia el acceso de la consola remota al clúster.
<code>if_mpadm(1M)</code>	Conmuta las direcciones IP de un adaptador a otro en un grupo de varias rutas de red IP.
<code>claccess(1CL)</code>	Administra políticas de acceso de Oracle Solaris Cluster para agregar nodos.
<code>cldevice(1CL)</code>	Administra dispositivos de Oracle Solaris Cluster.
<code>cldevicegroup(1CL)</code>	Administra grupos de dispositivos de Oracle Solaris Cluster.
<code>clinterconnect(1CL)</code>	Administra la interconexión de Oracle Solaris Cluster.
<code>clnasdevice(1CL)</code>	Administra el acceso a los servicios NAS de una configuración de Oracle Solaris Cluster.
<code>clnode(1CL)</code>	Administra nodos de Oracle Solaris Cluster.
<code>clquorum(1CL)</code>	Administra el quórum de Oracle Solaris Cluster.
<code>clreslogicalhostname(1CL)</code>	Administra recursos de Oracle Solaris Cluster para nombres de host lógicos.
<code>clresource(1CL)</code>	Administra recursos de servicios de datos de Oracle Solaris Cluster.
<code>clresourcegroup(1CL)</code>	Administra recursos de servicios de datos de Oracle Solaris Cluster.
<code>clresourcetype(1CL)</code>	Administra recursos de servicios de datos de Oracle Solaris Cluster.
<code>clressharedaddress(1CL)</code>	Administra recursos de Oracle Solaris Cluster para direcciones compartidas.
<code>clsetup(1CL)</code>	Configura una configuración de Oracle Solaris Cluster de forma interactiva.
<code>clsnmphost(1CL)</code>	Administra hosts SNMP de Oracle Solaris Cluster.
<code>clsnmpmib(1CL)</code>	Administra MIB de SNMP de Oracle Solaris Cluster.

<code>clsnmpuser(1CL)</code>	Administra usuarios de SNMP de Oracle Solaris Cluster.
<code>cltelemetryattribute(1CL)</code>	Configura la supervisión de los recursos del sistema.
<code>cluster(1CL)</code>	Administra la configuración y el estado global de la configuración de Oracle Solaris Cluster.
<code>clvxvm(1CL)</code>	Inicia Veritas Volume Manager (VxVM) en un nodo de Oracle Solaris Cluster y, de forma opcional, realiza una encapsulación del disco raíz.
<code>clzonecluster(1CL)</code>	Crea y modifica un clúster de zona.

Además, puede utilizar comandos para administrar la porción del administrador de volúmenes de una configuración de Oracle Solaris Cluster. Estos comandos dependen del administrador de volúmenes que utilice el clúster, ya sea Veritas Volume Manager o Solaris Volume Manager.

Preparación para administrar el clúster

En esta sección se explica cómo prepararse para administrar el clúster.

Documentación de las configuraciones de hardware de Oracle Solaris Cluster

Documente los aspectos del hardware que sean únicos de su sitio a medida que se escala la configuración de Oracle Solaris Cluster. Para reducir la administración, consulte la documentación del hardware al modificar o actualizar el clúster. Otra forma de facilitar la administración es etiquetar los cables y las conexiones entre los diversos componentes del clúster.

Guardar un registro de la configuración original del clúster y de los cambios posteriores reduce el tiempo que necesitan otros proveedores de servicios a la hora de realizar tareas de mantenimiento del clúster.

Uso de la consola de administración

Para administrar el clúster activo, puede optar por utilizar una estación de trabajo dedicada o una estación de trabajo conectada a través de una red de administración como *consola de administración*. En la consola de administración suelen instalarse y ejecutarse las herramientas del panel de control del clúster y la GUI. Para obtener más información sobre el panel de control del clúster, consulte [“Inicio de sesión remota en el clúster” en la página 25](#). Si necesita instrucciones sobre cómo instalar el módulo del panel de control del clúster para las

herramientas Sun Management Center y la GUI de Oracle Solaris Cluster Manager, consulte *Guía de instalación del software Oracle Solaris Cluster*.

La consola de administración no es un nodo del clúster. La consola de administración se utiliza para disponer de acceso remoto a los nodos del clúster, ya sea a través de la red pública o mediante un concentrador de terminales basado en red.

Si el clúster SPARC consta de un servidor Sun Enterprise 10000 de Oracle, debe iniciar sesión en el procesador de servicios del sistema desde la consola de administración. Establezca la conexión con el comando `net con`. El método predeterminado para que `net con` se conecte con un dominio de Sun Enterprise 10000 es través de la interfaz de red. Si no se tiene acceso a la red, puede utilizar `net con` en modo "exclusivo" mediante la opción `-f`. También puede escribir `~*` durante una sesión normal de `net con`. Cualquiera de estas dos soluciones brinda la opción de alternar a la interfaz serie si no se puede tener acceso a la red.

Oracle Solaris Cluster no necesita una consola de administración dedicada, pero emplearla aporta las ventajas siguientes:

- Permite administrar el clúster de manera centralizada agrupando en un mismo equipo herramientas de administración y de consola.
- Ofrece soluciones potencialmente más rápidas mediante servicios empresariales o del proveedor de servicios.

Copia de seguridad del clúster

Realice copias de seguridad del clúster de forma periódica. Aunque el software Oracle Solaris Cluster ofrece un entorno de alta disponibilidad con copias duplicadas de los datos en los dispositivos de almacenamiento, el software Oracle Solaris Cluster no sirve como sustituto de las copias de seguridad realizadas a intervalos periódicos. Una configuración de Oracle Solaris Cluster puede sobrevivir a multitud de errores, pero no ofrece protección frente a los errores de los usuarios o del programa ni un error fatal del sistema. Por lo tanto, debe disponer de un procedimiento de copia de seguridad para contar con protección frente a posibles pérdidas de datos.

Se recomienda que la información forme parte de la copia de seguridad.

- Todas las particiones del sistema de archivos
- Todos los datos de las bases de datos, en el caso de que se ejecuten servicios de datos DBMS
- Información sobre las particiones de los discos correspondiente a todos los discos del clúster

Procedimiento para empezar a administrar el clúster

La [Tabla 1–2](#) proporciona un punto de partida para administrar el clúster.

Nota – Los comandos de Oracle Solaris Cluster que sólo ejecuta desde el nodo de votación del clúster global no son válidos para los clústeres de zona. Consulte la página de comando man de Oracle Solaris Cluster para obtener información sobre la validez del uso de un comando en zonas.

TABLA 1–2 Herramientas de administración de Oracle Solaris Cluster

Tarea	Herramienta	Instrucciones
Iniciar sesión en el clúster de forma remota	Utilice el comando <code>ccp</code> para iniciar la ejecución del panel de control del clúster. A continuación seleccione uno de los siguientes iconos: <code>cconsole</code> , <code>crlogin</code> , <code>cssh</code> o <code>ctelnet</code> .	“Inicio de sesión remota en el clúster” en la página 25 “Conexión segura a las consolas del clúster” en la página 27
Configurar el clúster de forma interactiva	Inicie la unidad <code>clzonecluster(1CL)</code> o la utilidad <code>clsetup(1CL)</code> .	“Obtención de acceso a las utilidades de configuración del clúster” en la página 27
Mostrar información sobre el número y la versión de Oracle Solaris Cluster	Utilice el comando <code>clnode(1CL)</code> con el subcomando y la opción <code>show-rev -v -nodo</code> .	“Visualización de la información de versión de Oracle Solaris Cluster” en la página 29
Mostrar los recursos, grupos de recursos y tipos de recursos instalados	Utilice los comandos siguientes para mostrar la información sobre recursos: ■ <code>clresource(1CL)</code> ■ <code>clresourcegroup(1CL)</code> ■ <code>clresourcetype(1CL)</code>	“Visualización de tipos de recursos configurados, grupos de recursos y recursos” en la página 31
Supervisar los componentes del clúster de forma gráfica	Utilice Oracle Solaris Cluster Manager.	Consulte la ayuda en línea.
Administrar gráficamente algunos componentes del clúster	Utilice el módulo de Oracle Solaris Cluster Manager o Oracle Solaris Cluster con Sun Management Center, disponible sólo con Oracle Solaris Cluster en los sistemas basados en SPARC.	Para Oracle Solaris Cluster Manager, consulte la ayuda en línea Para Sun Management Center, consulte la documentación de Sun Management Center
Comprobar el estado de los componentes del clúster	Utilice el comando <code>cluster(1CL)</code> con el subcomando <code>status</code> .	“Comprobación del estado de los componentes del clúster” en la página 32

TABLA 1–2 Herramientas de administración de Oracle Solaris Cluster (Continuación)

Tarea	Herramienta	Instrucciones
Comprobar el estado de los grupos de varias rutas IP en la red pública	En un clúster global, utilice el comando <code>clnode(1CL)</code> status con la opción -m. En un clúster de zona, utilice el comando <code>clzonecluster(1CL)</code> show.	“Comprobación del estado de la red pública” en la página 35
Ver la configuración del clúster	En un clúster global, utilice el comando <code>cluster(1CL)</code> con el subcomando show. En un clúster de zona, utilice el comando <code>clzonecluster(1CL)</code> con el subcomando show.	“Visualización de la configuración del clúster” en la página 36
Ver y mostrar los dispositivos NAS configurados	En un clúster de zona o en un clúster global, utilice el comando <code>clzonecluster(1CL)</code> con el subcomando show.	<code>clnasdevice(1CL)</code>
Comprobar los puntos de montaje globales o verificar la configuración del clúster	En un clúster global, utilice el comando <code>cluster(1CL)cluster (1CL)</code> con el subcomando check. En un clúster de zona, utilice el comando <code>clzonecluster(1CL)</code> verify.	“Validación de una configuración básica de clúster” en la página 45
Examinar el contenido de los registros de comandos de Oracle Solaris Cluster	Examine el archivo <code>/var/cluster/logs/ commandlog</code> .	“Visualización del contenido de los registros de comandos de Oracle Solaris Cluster” en la página 48
Examinar los mensajes del sistema del Oracle Solaris Cluster	Examine el archivo <code>/var/adm/messages</code> .	“Viewing System Messages” de <i>System Administration Guide: Advanced Administration</i>
Supervisar el estado de Solaris Volume Manager	Utilice el comando <code>metastat</code> .	<i>Solaris Volume Manager Administration Guide</i>

▼ Inicio de sesión remota en el clúster

El panel de control del clúster ofrece un launchpad para las herramientas `cconsole`, `crlogin`, `cssh` y `ctelnet`. Todas las herramientas inician una conexión de varias ventanas con un conjunto de nodos especificados. La conexión de varias ventanas consta de una ventana host

para cada uno de los nodos indicados y una ventana común. Las entradas recibidas en la ventana común se envían a cada una de las ventanas del host; de este modo, se pueden ejecutar comandos simultáneamente en todos los nodos del clúster.

También puede iniciar sesiones de `cconsole`, `crlogin`, `cssh` o `ctelnet` desde la línea de comandos.

La utilidad `cconsole` emplea de forma predeterminada una conexión `telnet` con las consolas de los nodos. Para establecer conexiones de shell seguro con las consolas, en el menú Opciones de la ventana `cconsole` habilite la casilla de verificación Usar SSH. También puede indicar la opción `-s` al ejecutar el comando `ccp` o `cconsole`.

Para obtener más información, consulte las páginas de comando [man ccp\(1M\)](#) y [cconsole\(1M\)](#).

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Antes de empezar

Antes de iniciar el panel de control del clúster, compruebe que se cumplan los requisitos siguientes:

- Instale el paquete `SUNWccn` en la consola de administración.
- Compruebe que la variable `PATH` de la consola de administración incluya los directorios de herramientas de Oracle Solaris Cluster, `/opt/SUNWcluster/bin` y `/usr/cluster/bin`. Puede indicar una ubicación alternativa para el directorio de herramientas si establece la variable de entorno `$CLUSTER_HOME`.
- Configure el archivo `clusters`, el archivo `serialports` y el archivo `nsswitch.conf` si utiliza un concentrador de terminales. Los archivos pueden ser archivos `/etc`, o bases de datos NIS o NIS+. Consulte las páginas de comando [man clusters\(4\)](#) y [serialports\(4\)](#) si desea obtener más información.

1 Si dispone de una plataforma Servidor Sun Enterprise 10000, inicie sesión en el procesador de servicios del sistema.

a. Establezca la conexión con el comando `netcon`.

b. Una vez establecida la conexión, escriba `Shift~@` para desbloquear la consola y obtener acceso de escritura.

2 Desde la consola de administración, inicie el launchpad del panel de control del clúster.

`phys-schost# ccp clustername`

Se muestra el launchpad del panel de control del clúster.

- 3 Para iniciar una sesión remota con el clúster, haga clic en el icono de `cconsole`, `crlogin`, `cssh` o `ctelnet` en el launchpad del panel de control del clúster.

▼ Conexión segura a las consolas del clúster

Lleve a cabo este procedimiento para establecer conexiones de shell seguro con las consolas de los nodos del clúster.

Antes de empezar

Configure los archivos `clusters`, `serialports` y `nsswitch.conf` si utiliza un concentrador de terminales. Los archivos pueden ser archivos `/etc`, o bases de datos NIS o NIS+.

Nota – En el archivo `serialports`, asigne el número de puerto que se va a utilizar para la conexión segura a cada dispositivo con acceso a la consola. El número de puerto predeterminado para la conexión con shell seguro es el 22.

Si desea obtener más información, consulte las páginas de comando `man clusters(4)` y `serialports(4)`.

- 1 Conviértase en superusuario en la consola de administración.

- 2 Inicie la utilidad `cconsole` en modo seguro.

```
# cconsole -s [-l username] [-p ssh-port]
```

<code>-s</code>	Habilita la conexión con shell seguro.
<code>-l nombre_usuario</code>	Especifica el nombre de usuario para las conexiones remotas. Si la opción <code>-l</code> no está indicada, se utiliza el nombre del usuario que ejecutó la utilidad <code>cconsole</code> .
<code>-p puerto_ssh</code>	Indica el número de puerto que se usa para el shell seguro. Si no se indica la opción <code>-p</code> , de forma predeterminada se utiliza el puerto número 22 para las conexiones seguras.

▼ Obtención de acceso a las utilidades de configuración del clúster

La utilidad `clsetup` permite configurar de forma interactiva el quórum, los grupos de recursos, el transporte del clúster, los nombres de host privados, los grupos de dispositivos y las nuevas

opciones de nodos del clúster global. La utilidad `clzonecluster` realiza tareas de configuración similares para los clústers de zona. Si desea más información, consulte las páginas de comando `man clsetup(1CL)` y `clzonecluster(1CL)`.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario de un nodo de miembro activo en un clúster global. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

- 2 **Inicie la utilidad de configuración.**

```
phys-schost# clsetup
```

- **Para un clúster global, inicie la utilidad con el comando `clsetup`.**

```
phys-schost# clsetup
```

Aparece el menú principal.

- **En el caso de un clúster de zona, inicie la utilidad con el comando `clzonecluster`. El clúster de zona de este ejemplo es `zona_sc`.**

```
phys-schost# clzonecluster configure szzone
```

Con la opción siguiente puede ver las acciones disponibles en la utilidad:

```
clzc:szzone> ?
```

- 3 **En el menú, elija la configuración. Siga las instrucciones en pantalla para finalizar una tarea. Si desea ver más detalles, consulte las instrucciones de [“Configuración de un clúster de zona” de Guía de instalación del software Oracle Solaris Cluster](#).**

Véase también Consulte la ayuda en línea de `clsetup` o `clzonecluster` para obtener más información.

▼ Visualización de la información de parches de Oracle Solaris Cluster

Para realizar este procedimiento no es necesario iniciar sesión como superusuario.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- **Visualice la información sobre parches de Oracle Solaris Cluster:**

```
phys-schost# showrev -p
```

Las versiones de las actualizaciones de Oracle Solaris Cluster se identifican mediante el número de parche del producto principal y la versión de la actualización.

Ejemplo 1–1 Visualización de la información sobre los parches de Oracle Solaris Cluster

En el ejemplo siguiente se muestra información sobre el parche 110648-05.

```
phys-schost# showrev -p | grep 110648
Patch: 110648-05 Obsoletes: Requires: Incompatibles: Packages:
```

▼ **Visualización de la información de versión de Oracle Solaris Cluster**

Para realizar este procedimiento no es necesario iniciar sesión como superusuario. Siga todos los pasos de este procedimiento desde un nodo del clúster global.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- **Visualice la información de versión de Oracle Solaris Cluster:**

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

Este comando muestra el número de versión y las cadenas de caracteres de versión de todos los paquetes de Oracle Solaris Cluster.

Ejemplo 1–2 Visualización de la información de versión de Oracle Solaris Cluster

El ejemplo siguiente muestra la información de versión del clúster y de versión de todos los paquetes.

```
phys-schost# clnode show-rev
3.2
```

```

phys-schost#% clnode show-rev -v
Oracle Solaris Cluster 3.3 for Solaris 10 sparc
SUNWscu: 3.3.0,REV=2010.06.14.03.44
SUNWsccomu: 3.3.0,REV=2010.06.14.03.44
SUNWsczr: 3.3.0,REV=2010.06.14.03.44
SUNWsccomzu: 3.3.0,REV=2010.06.14.03.44
SUNWsczu: 3.3.0,REV=2010.06.14.03.44
SUNWscsckr: 3.3.0,REV=2010.06.14.03.44
SUNWscscku: 3.3.0,REV=2010.06.14.03.44
SUNWscr: 3.3.0,REV=2010.06.14.03.44
SUNWscrtlh: 3.3.0,REV=2010.06.14.03.44
SUNWscnmr: 3.3.0,REV=2010.06.14.03.44
SUNWscnmu: 3.3.0,REV=2010.06.14.03.44
SUNWscdev: 3.3.0,REV=2010.06.14.03.44
SUNWscgds: 3.3.0,REV=2010.06.14.03.44
SUNWscsmf: 3.3.0,REV=2010.06.14.03.44
SUNWscman: 3.3.0,REV=2010.05.21.18.40
SUNWscsal: 3.3.0,REV=2010.06.14.03.44
SUNWscsam: 3.3.0,REV=2010.06.14.03.44
SUNWscvm: 3.3.0,REV=2010.06.14.03.44
SUNWmdmr: 3.3.0,REV=2010.06.14.03.44
SUNWmdmu: 3.3.0,REV=2010.06.14.03.44
SUNWscmasa: 3.3.0,REV=2010.06.14.03.44
SUNWscmasar: 3.3.0,REV=2010.06.14.03.44
SUNWscmasasen: 3.3.0,REV=2010.06.14.03.44
SUNWscmasazu: 3.3.0,REV=2010.06.14.03.44
SUNWscmasau: 3.3.0,REV=2010.06.14.03.44
SUNWscmautil: 3.3.0,REV=2010.06.14.03.44
SUNWscmautilr: 3.3.0,REV=2010.06.14.03.44
SUNWjfreechart: 3.3.0,REV=2010.06.14.03.44
SUNWscspmr: 3.3.0,REV=2010.06.14.03.44
SUNWscspmu: 3.3.0,REV=2010.06.14.03.44
SUNWscderby: 3.3.0,REV=2010.06.14.03.44
SUNWsc telemetry: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepavs: 3.2.3,REV=2009.10.23.12.12
SUNWscgrepsrdf: 3.2.3,REV=2009.10.23.12.12
SUNWscgreptc: 3.2.3,REV=2009.10.23.12.12
SUNWscghb: 3.2.3,REV=2009.10.23.12.12
SUNWscgctl: 3.2.3,REV=2009.10.23.12.12
SUNWscims: 6.0,REV=2003.10.29
SUNWscics: 6.0,REV=2003.11.14
SUNWscapc: 3.2.0,REV=2006.12.06.18.32
SUNWscdns: 3.2.0,REV=2006.12.06.18.32
SUNWschadb: 3.2.0,REV=2006.12.06.18.32
SUNWschtt: 3.2.0,REV=2006.12.06.18.32
SUNWscslas: 3.2.0,REV=2006.12.06.18.32
SUNWscsrb5: 3.2.0,REV=2006.12.06.18.32
SUNWscnfs: 3.2.0,REV=2006.12.06.18.32
SUNWscor: 3.2.0,REV=2006.12.06.18.32
SUNWscslmq: 3.2.0,REV=2006.12.06.18.32
SUNWscsap: 3.2.0,REV=2006.12.06.18.32
SUNWscslc: 3.2.0,REV=2006.12.06.18.32
SUNWscsapdb: 3.2.0,REV=2006.12.06.18.32
SUNWscsapenq: 3.2.0,REV=2006.12.06.18.32
SUNWscsaprepl: 3.2.0,REV=2006.12.06.18.32
SUNWscsapscs: 3.2.0,REV=2006.12.06.18.32
SUNWscsapwebas: 3.2.0,REV=2006.12.06.18.32
SUNWscsbl: 3.2.0,REV=2006.12.06.18.32

```

```

SUNWscsyb:      3.2.0,REV=2006.12.06.18.32
SUNWscwls:      3.2.0,REV=2006.12.06.18.32
SUNWsc9ias:      3.2.0,REV=2006.12.06.18.32
SUNWscPostgreSQL: 3.2.0,REV=2006.12.06.18.32
SUNWsczone:      3.2.0,REV=2006.12.06.18.32
SUNWscdhc:      3.2.0,REV=2006.12.06.18.32
SUNWscsbebs:      3.2.0,REV=2006.12.06.18.32
SUNWscmqi:      3.2.0,REV=2006.12.06.18.32
SUNWscmqms:      3.2.0,REV=2006.12.06.18.32
SUNWscmys:      3.2.0,REV=2006.12.06.18.32
SUNWscsge:      3.2.0,REV=2006.12.06.18.32
SUNWscsaa:      3.2.0,REV=2006.12.06.18.32
SUNWscsag:      3.2.0,REV=2006.12.06.18.32
SUNWscsmb:      3.2.0,REV=2006.12.06.18.32
SUNWscsps:      3.2.0,REV=2006.12.06.18.32
SUNWscsctomcat: 3.2.0,REV=2006.12.06.18.32

```

▼ Visualización de tipos de recursos configurados, grupos de recursos y recursos

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte el [Capítulo 13, “Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario”](#), o consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Antes de empezar

Para utilizar subcomando, todos los usuarios necesitan la autorización RBAC `solaris.cluster.read` menos el superusuario.

- **Visualice los tipos de recursos, los grupos de recursos y los recursos configurados para el clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

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

Si desea obtener información sobre un determinado recurso, los grupos de recursos y los tipos de recursos, utilice el subcomando `show` con uno de los comandos siguientes:

- `resource`
- `resource group`
- `resourcetype`

Ejemplo 1-3 Visualización de tipos de recursos, grupos de recursos y recursos configurados

En el ejemplo siguiente se muestran los tipos de recursos (RT Name), los grupos de recursos (RG Name) y los recursos (RS Name) configurados para el clúster schost.

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

=== Registered Resource Types ===

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

=== 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:                                default
Resource_project_name:                        default
Enabled{phys-schost-2}:                       True
Enabled{phys-schost-1}:                       True
Monitored{phys-schost-2}:                     True
Monitored{phys-schost-1}:                     True
```

▼ Comprobación del estado de los componentes del clúster

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Nota – El comando `cluster status` también muestra el estado de un clúster de zona.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Antes de empezar Todos los usuarios necesitan la autorización RBAC `solaris.cluster.read` para utilizar el subcomando `status` menos el superusuario.

- **Compruebe el estado de los componentes del clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

`phys-schost# cluster status`

Ejemplo 1–4 Comprobación del estado de los componentes del clúster

El ejemplo siguiente es una muestra de la información del estado de los componentes del clúster devuelta por `cluster(1CL) status`.

```
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:qfe1                      phys-schost-4:qfe1                      Path online
phys-schost-1:hme1                      phys-schost-4:hme1                      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/rdisk/d2   phys-schost-1    Ok
/dev/did/rdisk/d3   phys-schost-1    Ok
                    phys-schost-2    Ok
/dev/did/rdisk/d4   phys-schost-1    Ok
                    phys-schost-2    Ok
/dev/did/rdisk/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      Online  Running
```

▼ **Comprobación del estado de la red pública**

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Para comprobar el estado de los grupos de varias rutas de red IP, utilice el comando `clnode(1CL)` con el subcomando `status`.

Antes de empezar Todos los usuarios necesitan la autorización RBAC `solaris.cluster.read` utilizar este subcomando menos el superusuario.

- **Compruebe el estado de los componentes del clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

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

Ejemplo 1-5 Comprobación del estado de la red pública

En el ejemplo siguiente se puede ver una muestra de la información del estado de los componentes del clúster devuelta por el comando `clnode status`.

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

Node Name      Group Name      Status  Adapter  Status
-----
phys-schost-1   test-rg         Online  qfe1     Online
phys-schost-2   test-rg         Online  qfe1     Online
```

▼ Visualización de la configuración del clúster

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Antes de empezar Todos los usuarios necesitan la autorización de RBAC `solaris.cluster.read` para utilizar el subcomando `status` menos el superusuario.

- **Visualice la configuración de un clúster global o un clúster de zona. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

```
% cluster show
```

Al ejecutar el comando `cluster show` desde un nodo de votación de clúster global, se muestra información detallada sobre la configuración del clúster e información sobre los clústers de zona, si es que están configurados.

También puede usar el comando `clzonecluster show` para visualizar la información de configuración sólo de los clústers de zona. Entre las propiedades de un clúster de zona están el nombre, el tipo de IP, el arranque automático y la ruta de zona. El subcomando `show` se ejecuta dentro de un clúster de zona y se aplica sólo a ese clúster de zona. Al ejecutar el comando `clzonecluster show` desde un nodo del clúster de zona, sólo se recupera el estado de los objetos visibles en ese clúster de zona en concreto.

Para visualizar más información acerca del comando `cluster`, utilice las opciones para obtener más detalles. Si desea obtener más detalles, consulte la página de comando `man cluster(1CL)`. Consulte la página de comando `man clzonecluster(1CL)` si desea obtener más información sobre `clzonecluster`.

Ejemplo 1–6 Visualización de la configuración del clúster global

En el ejemplo siguiente figura información de configuración sobre el clúster global. Si tiene configurado un clúster de zona, también se enumera la pertinente información.

```
phys-schost# cluster show
```

```
=== Cluster ===
```

```
Cluster Name:                cluster-1
installmode:                 disabled
heartbeat_timeout:          10000
heartbeat_quantum:          1000
private_netaddr:            172.16.0.0
private_netmask:            255.255.248.0
max_nodes:                   64
max_privatenets:            10
global_fencing:             Unknown
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
Type:                         cluster
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:       qfe3, hme0
```

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

```
Transport Adapter:           qfe3
Adapter State:               Enabled
Adapter Transport Type:      dlpi
Adapter Property(device_name): qfe
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): 172.16.1.1
Adapter Property(netmask): 255.255.255.128
Adapter Port Names: 0
Adapter Port State(0): Enabled

Transport Adapter: hme0
Adapter State: Enabled
Adapter Transport Type: dlpi
Adapter Property(device_name): hme
Adapter Property(device_instance): 0
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
defaulttpsetmin: 2
quorum_vote: 1
quorum_defaultvote: 1
quorum_resv_key: 0x43CB1E1800000002
Transport Adapter List: hme0, qfe3

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

Transport Adapter: hme0
Adapter State: Enabled
Adapter Transport Type: dlpi
Adapter Property(device_name): hme

```

```

Adapter Property(device_instance):      0
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:                      qfe3
Adapter State:                          Enabled
Adapter Transport Type:                  dlpi
Adapter Property(device_name):           qfe
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:qfe3,switch2@1
Cable Endpoint1:                        phys-schost-1:qfe3
Cable Endpoint2:                        switch2@1
Cable State:                            Enabled

Transport Cable:                         phys-schost-1:hme0,switch1@1
Cable Endpoint1:                        phys-schost-1:hme0
Cable Endpoint2:                        switch1@1
Cable State:                            Enabled

Transport Cable:                         phys-schost-2:hme0,switch1@2
Cable Endpoint1:                        phys-schost-2:hme0
Cable Endpoint2:                        switch1@2
Cable State:                            Enabled

Transport Cable:                         phys-schost-2:qfe3,switch2@2
Cable Endpoint1:                        phys-schost-2:qfe3
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:	scsi
Access Mode:	scsi2
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:	2
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:                 SUNWscu
RT_system:              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:                SUNWscu
RT_system:              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:                SUNWscu
RT_system:              False

```

```

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:                SUNWscderby
RT_system:              False

```

```

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:                SUNWsc telemetry
RT_system:              False

```

=== Resource Groups and Resources ===

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

--- Resources for Group HA_RG ---

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

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

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

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

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

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

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

Nas Device:                                nas2
Type:                                       netapp
User ID:                                   llai

```

Ejemplo 1-7 Visualización de la información del clúster de zona

En el ejemplo siguiente figuran las propiedades de la configuración del clúster de zona.

```

% 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:                    /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

```

También puede ver los dispositivos NAS que se han configurado para clústeres globales o de zona, utilizando el subcomando `clnasdevice show` o Oracle Solaris Cluster Manager. Para obtener más información, consulte la página de comando `man clnasdevice(1CL)`.

▼ Validación de una configuración básica de clúster

El comando `cluster(1CL)` se sirve del comando `check` para validar la configuración básica que necesita un clúster global para funcionar correctamente. Si ninguna comprobación arroja un resultado incorrecto, `cluster check` devuelve al indicador del shell. Si falla alguna de las comprobaciones, `cluster check` emite informes que aparecen en el directorio de salida que se haya especificado o en el predeterminado. Si ejecuta `cluster check` con más de un nodo, `cluster check` emite un informe para cada nodo y un informe para las comprobaciones que comprenden varios nodos. También puede utilizar el comando `cluster list-checks` para que se muestre una lista con todas las comprobaciones disponibles para el clúster.

Puede ejecutar el comando `cluster check` en modo detallado con el indicador `-v` para que se muestre la información de progreso.

Nota – Ejecute `cluster check` después de realizar un procedimiento de administración que pueda provocar modificaciones en los dispositivos, en los componentes de administración de volúmenes o en la configuración de Oracle Solaris Cluster.

Al ejecutar el comando `clzonecluster(1CL)` en el nodo de votación de clúster global, se lleva a cabo un conjunto de comprobaciones con el fin de validar la configuración necesaria para que un clúster de zona funcione correctamente. Si todas las comprobaciones son correctas, `clzonecluster verify` devuelve al indicador de shell y el clúster de zona se puede instalar con seguridad. Si falla alguna de las comprobaciones, `clzonecluster verify` informa sobre los nodos del clúster global en los que la verificación no obtuvo un resultado correcto. Si ejecuta `clzonecluster verify` respecto a más de un nodo, se emite un informe para cada nodo y un informe para las comprobaciones que comprenden varios nodos. No se permite utilizar el subcomando `verify` dentro de los clústers de zona.

- 1 **Conviértase en superusuario de un nodo de miembro activo en un clúster global. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

```
phys-schost# su
```

- 2 **Compruebe la configuración del clúster.**

- **Compruebe la configuración del clúster global.**

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

- **Compruebe la configuración del clúster de zona para ver si es posible instalar un clúster de zona.**

```
phys-schost# clzonecluster verify zoneclustername
```

Ejemplo 1–8 Comprobación de la configuración del clúster global con resultado correcto en todas las comprobaciones

El ejemplo siguiente muestra cómo la ejecución de `cluster check` en modo detallado respecto a los nodos `phys-schost-1` y `phys-schost-2` supera correctamente todas las comprobaciones.

```
phys-schost# cluster check -v -h 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
#
```

Ejemplo 1–9 Comprobación de la configuración del clúster global con una comprobación con resultado no satisfactorio

El ejemplo siguiente muestra el nodo `phys-schost-2`, perteneciente al clúster denominado `suncluster`, excepto el punto de montaje `/global/phys-schost-1`. Los informes se crean en el directorio de salida `/var/cluster/logs/cluster_check/<timestamp>`.

```
phys-schost# cluster check -v -h 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 3.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.
...
#

```

▼ Comprobación de los puntos de montaje globales

El comando `cluster(1CL)` incluye comprobaciones que verifican el archivo `/etc/vfstab` para detectar posibles errores de configuración con el sistema de archivos del clúster y sus puntos de montaje globales.

Nota – Ejecute `cluster check` después de efectuar cambios en la configuración del clúster que hayan afectado a los dispositivos o a los componentes de administración de volúmenes.

- 1 **Conviértase en superusuario de un nodo de miembro activo en un clúster global. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

```
% su
```

- 2 **Compruebe la configuración del clúster global.**

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

Ejemplo 1–10 Comprobación de puntos de montaje globales

El ejemplo siguiente muestra el nodo `phys-schost-2` del clúster denominado `suncluster`, excepto el punto de montaje `/global/schost-1`. Los informes se envían al directorio de salida, `/var/cluster/logs/cluster_check/<timestamp>/`.

```

phys-schost# cluster check -v1 -h 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 3.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 3.x node.
ANALYSIS : This server may not been qualified to be used as an Oracle Solaris Cluster 3.x node.
Only servers that have been qualified with Oracle Solaris Cluster 3.x are supported as
Oracle Solaris Cluster 3.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 3.x.
...
#
```

▼ Visualización del contenido de los registros de comandos de Oracle Solaris Cluster

El archivo de texto ASCII /var/cluster/logs/commandlog contiene registros de comandos de Oracle Solaris Cluster seleccionados que se ejecutan en un clúster. Los comandos comienzan a

registrarse automáticamente al configurarse el clúster y la operación se finaliza al cerrarse el clúster. Los comandos se registran en todos los nodos activos y que se han arrancado en modo de clúster.

Entre los comandos que no quedan registrados en este archivo están los encargados de mostrar la configuración y el estado actual del clúster.

Entre los comandos que quedan registrados en este archivo están los encargados de configurar y modificar el estado actual del clúster:

- `claccess`
- `cldevice`
- `cldevicegroup`
- `clinterconnect`
- `clnasdevice`
- `clnode`
- `clquorum`
- `clreslogicalhostname`
- `clresource`
- `clresourcegroup`
- `clresourcetype`
- `clressharedaddress`
- `clsetup`
- `clsnmp host`
- `clsnmp mib`
- `clsnmp user`
- `cltelemetryattribute`
- `cluster`
- `clzonecluster`
- `scdidadm`

Los registros del archivo `commandlog` pueden contener los elementos siguientes:

- Fecha y marca de tiempo
- Nombre del host desde el cual se ejecutó el comando
- ID de proceso del comando
- Nombre de inicio de sesión del usuario que ejecutó el comando
- Comando ejecutado por el usuario, con todas las opciones y los operandos

Nota – Las opciones del comando se recogen en el archivo `commandlog` para poderlas identificar, copiar, pegar y ejecutar en el shell.

- Estado de salida del comando ejecutado

Nota – Si un comando interrumpe su ejecución de forma anómala con resultados desconocidos, el software Oracle Solaris Cluster *no* muestra un estado de salida en el archivo `commandlog`.

De forma predeterminada, el archivo `commandlog` se archiva periódicamente una vez por semana. Para modificar las directrices de archivado del archivo `commandlog`, utilice el comando `crontab` en todos los nodos del clúster. Para obtener más información, consulte la página de comando `man crontab(1)`.

El software Oracle Solaris Cluster conserva en todo momento hasta un máximo de ocho archivos `commandlog` archivados anteriormente en cada nodo del clúster. El archivo `commandlog` de la semana actual se denomina `commandlog`. El archivo semanal completo más reciente se denomina `commandlog.0`. El archivo semanal completo más antiguo se denomina `commandlog.7`.

- **El contenido del archivo `commandlog`, correspondiente a la semana actual, se muestra una sola pantalla a la vez.**

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

Ejemplo 1–11 Visualización del contenido de los registros de comandos de Oracle Solaris Cluster

El ejemplo siguiente muestra el contenido del archivo `commandlog` que se visualiza mediante el comando `more`.

```
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 y RBAC

Este capítulo describe el control de acceso basado en funciones (RBAC) en relación con Oracle Solaris Cluster. Los temas que se tratan son:

- “Instalación y uso de RBAC con Oracle Solaris Cluster” en la página 51
- “Perfiles de derechos de RBAC en Oracle Solaris Cluster” en la página 52
- “Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster” en la página 53
- “Modificación de las propiedades de RBAC de un usuario” en la página 57

Instalación y uso de RBAC con Oracle Solaris Cluster

Utilice la tabla siguiente para determinar la documentación que debe consultar sobre la instalación y el uso de RBAC. Más adelante en este mismo capítulo, se indican los pasos específicos para instalar y utilizar RBAC con el software Oracle Solaris Cluster.

Tarea	Instrucciones
Más información sobre RBAC	Capítulo 8, “Using Roles and Privileges (Overview)” de <i>System Administration Guide: Security Services</i>
Instalar, administrar elementos y utilizar RBAC	Capítulo 9, “Using Role-Based Access Control (Tasks)” de <i>System Administration Guide: Security Services</i>
Más información sobre elementos y herramientas de RBAC	Capítulo 10, “Role-Based Access Control (Reference)” de <i>System Administration Guide: Security Services</i>

Perfiles de derechos de RBAC en Oracle Solaris Cluster

Las opciones y los comandos seleccionados de Oracle Solaris Cluster Manager y Oracle Solaris Cluster que se introducen en la línea de comandos usan RBAC para la autorización. Los comandos y las opciones de Oracle Solaris Cluster que requieran la autorización de RBAC necesitan al menos uno de los niveles de autorización siguientes. Los perfiles de derechos de RBAC de Oracle Solaris Cluster se aplican a los nodos de votación y a los que no son de votación de clúster global.

- `solaris.cluster.read`

Autorización para operaciones de enumeración, visualización y otras operaciones de lectura.
- `solaris.cluster.admin`

Autorización para cambiar el estado de un objeto de clúster.
- `solaris.cluster.modify`

Autorización para cambiar las propiedades de un objeto de clúster.

Para obtener más información sobre la autorización de RBAC que necesita un comando de Oracle Solaris Cluster, consulte la página de comando `man`.

Los perfiles de derechos de RBAC tienen al menos una autorización de RBAC. Puede asignar estos perfiles de derechos a los usuarios o a las funciones para darles distintos niveles de acceso a Oracle Solaris Cluster. Oracle proporciona los perfiles de derechos siguientes con el software Oracle Solaris Cluster.

Nota – Los perfiles de derechos de RBAC que aparecen en la tabla siguiente continúan admitiendo las autorizaciones antiguas de RBAC, tal y como se define en las versiones anteriores de Oracle Solaris Cluster.

Perfil de derechos	Incluye autorizaciones	Permiso de identidad de función
Comandos de Oracle Solaris Cluster	Ninguna, pero incluye una lista de comandos de Oracle Solaris Cluster que se ejecutan con <code>euclid=0</code> .	Ejecutar los comandos seleccionados de Oracle Solaris Cluster que se usen para configurar y administrar un clúster, incluidos los subcomandos siguientes para todos los comandos de Oracle Solaris Cluster: ■ list ■ show ■ status <code>scha_control(1HA)</code> <code>scha_resource_get(1HA)</code> <code>scha_resource_setstatus(1HA)</code> <code>scha_resourcegroup_get(1HA)</code> <code>scha_resourcetype_get(1HA)</code>

Perfil de derechos	Incluye autorizaciones	Permiso de identidad de función
Usuario de Oracle Solaris básico	Este perfil de derechos de Oracle Solaris contiene autorizaciones de Oracle Solaris, además de la siguiente: <code>solaris.cluster.read</code>	Enumerar, mostrar y realizar otras operaciones de lectura para los comandos de Oracle Solaris Cluster, así como acceder a la GUI de Oracle Solaris Cluster Manager.
Operación del clúster	El perfil de derechos es específico del software Oracle Solaris Cluster y cuenta con las autorizaciones siguientes: <code>solaris.cluster.read</code> <code>solaris.cluster.admin</code>	Enumerar, mostrar, exportar, mostrar el estado y realizar otras operaciones de lectura, así como acceder a la GUI de Oracle Solaris Cluster Manager. Cambiar el estado de los objetos de clúster.
Administrador del sistema	Este perfil de derechos de Oracle Solaris contiene las mismas autorizaciones que el perfil de administración del clúster.	Realizar las mismas operaciones que la identidad de función de administración del clúster, además de otras operaciones de administración del sistema.
Administración del clúster	Este perfil de derechos contiene las mismas autorizaciones que el perfil de operaciones del clúster, además de la autorización siguiente: <code>solaris.cluster.modify</code>	Realizar las mismas operaciones que la identidad de función de operaciones del clúster, además de cambiar las propiedades de un objeto del clúster.

Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster

Use esta tarea para crear una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster y para asignar a los usuarios esta función nueva.

▼ Creación de una función con la herramienta Administrative Roles

Antes de empezar

Para crear una función, debe asumir una que disponga del perfil de derechos del administrador principal asignado o ejecutarla como usuario root.

1 Inicie la herramienta Administrative Roles.

Para ejecutar la herramienta Administrative Roles, inicie Solaris Management Console, como se describe en [“How to Assume a Role in the Solaris Management Console” de System Administration Guide: Security Services](#). Abra User Tool Collection y haga clic en el icono de funciones administrativas.

2 Inicie el asistente Add Administrative Role (Agregar funciones administrativas) para agregar una función administrativa.

Seleccione Add Administrative Role (Agregar función administrativa) en el menú Action (Acción) para iniciar el asistente Add Administrative Role (Agregar funciones administrativas) para la configuración de funciones.

3 Configure una función que tenga asignado el perfil de derechos de administración del clúster.

Utilice los botones Siguiente y Atrás para desplazarse entre los cuadros de diálogo. El botón Siguiente no se activa mientras no se hayan completado los campos obligatorios. El último cuadro de diálogo permite revisar los datos introducidos y retroceder para cambiarlos o hacer clic en Finish (Finalizar) para guardar la nueva función. La lista siguiente resume los campos y botones del cuadro de diálogo.

Role Name	Nombre corto de la función.
Full Name	Versión larga del nombre.
Description	Descripción de la función.
Role ID Number	UID para la función, incrementado de forma automática.
Role Shell	Los shells de perfil disponibles para las funciones: los shells C, Bourne y Korn del administrador.
Create a role mailing list	Realiza una lista de correo para los usuarios asignados a esta función.
Available Rights / Granted Rights	Asigna o elimina los perfiles de derechos de una función. El sistema no impide escribir varias apariciones del mismo comando. Los atributos asignados a la primera aparición de un comando en un perfil de derechos tienen preferencia y se ignoran todas las siguientes apariciones. Use las flechas Arriba y Abajo para cambiar el orden.
Server	Servidor para el directorio de inicio.
Path	Ruta del directorio de inicio.
Add	Agrega usuarios que puedan asumir esta función. Deben estar en el mismo ámbito.
Delete	Elimina los usuarios asignados a esta función.

Nota – Este perfil se debe colocar en primer lugar en la lista de perfiles asignados a la función.

4 Agregue los usuarios que necesiten utilizar las funciones de Oracle Solaris Cluster Manager o los comandos de Oracle Solaris Cluster a la función recién creada.

Use el comando `useradd(1M)` para agregar una cuenta de usuario al sistema. La opción `-P` asigna una función a una cuenta de usuario.

5 Haga clic en Finish (Finalizar).

6 Abra una ventana de terminal y conviértase en usuario root.

7 Inicie y detenga el daemon de memoria caché del servicio de nombres.

La función nueva no se activará hasta que se reinicie el daemon de memoria caché del servicio de nombres. Después de convertirse en usuario root, escriba el texto siguiente:

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

▼ Creación de una función desde la línea de comandos

1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.admin`.

2 Seleccione un método para crear una función:

- Para las funciones de ámbito local, use el comando `roleadd(1M)` para especificar una función local nueva y sus atributos.
- Asimismo, para las funciones de ámbito local, edite el archivo `user_attr(4)` para agregar un usuario con `type=role`.

Use este método sólo en caso de emergencia.

- Para las funciones de un servicio de nombres, use el comando `smrole(1M)` para especificar la función nueva y sus atributos.

Este comando requiere autenticación por parte del usuario o una función que, a su vez, pueda crear otras. Puede aplicar el comando `smrole` a todos los servicios de nombres. Este comando se ejecuta como cliente del servidor de Solaris Management Console.

3 Inicie y detenga el daemon de memoria caché del servicio de nombres.

Las funciones nuevas no se activan hasta que se reinicie el daemon de memoria caché del servicio de nombres. Como usuario root, escriba el texto siguiente:

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

Ejemplo 2-1 Creación de una función de operador personalizado con el comando smrole

La siguiente secuencia muestra cómo se crea una función con el comando smrole. En este ejemplo, se crea una versión nueva de la función de operador a la que tiene asignado el perfil de derechos del operador estándar, así como el perfil de restauración de soporte.

```
% 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 ::      <type primaryadmin password>

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 ::      <type oper2 password>
```

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

Para ver la función recién creada (y cualquier otra), use el comando smrole con la opción list, tal y como se indica a continuación:

```
# /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 ::      <type primaryadmin password>

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
```


Modificación de las propiedades de RBAC de un usuario

Puede modificar las propiedades de RBAC de un usuario con la herramienta User Accounts o la línea de comandos. Para modificar las propiedades de RBAC de un usuario, elija uno de los procedimientos siguientes.

- [“Modificación de las propiedades de RBAC de un usuario con la herramienta User Accounts” en la página 57](#)
- [“Modificación de las propiedades de RBAC de un usuario desde la línea de comandos” en la página 58](#)

▼ Modificación de las propiedades de RBAC de un usuario con la herramienta User Accounts

Antes de empezar

Para modificar las propiedades de un usuario, debe ejecutar User Tool Collection como superusuario o asumir una función que disponga del perfil de derechos de administrador principal.

1 Inicie la herramienta User Accounts.

Para ejecutar la herramienta User Accounts, inicie Solaris Management Console, tal y como se describe en [“How to Assume a Role in the Solaris Management Console” de *System Administration Guide: Security Services*](#). Abra User Tool Collection y haga clic en el icono User Accounts (Cuentas de usuario).

Una vez iniciada la herramienta User Accounts, los iconos correspondientes a las cuentas de usuario se muestran en el panel de visualización.

2 Haga clic en el icono de la cuenta de usuario que se debe modificar; a continuación, seleccione Properties (Propiedades) en el menú Action (Acción) o haga doble clic en el icono.

3 Haga clic en la ficha correspondiente del cuadro de diálogo según la propiedad que se desee modificar, como se indica a continuación:

- Para cambiar las funciones asignadas al usuario, haga clic en la ficha Roles (Funciones) y mueva la asignación de función que se debe modificar a la columna apropiada: Available Roles (Funciones disponibles) o Assigned Roles (Funciones asignadas).
- Para cambiar los perfiles de derechos asignados al usuario, haga clic en la ficha Rights (Derechos) y mueva los perfiles a la columna apropiada: Available Rights (Derechos disponibles) o Assigned Rights (Derechos asignados).

Nota – No es conveniente asignar perfiles de derechos directamente a usuarios. Es más adecuado indicarles que asuman funciones para poder ejecutar aplicaciones con privilegios. Esta estrategia impide que los usuarios normales abusen de los privilegios.

▼ Modificación de las propiedades de RBAC de un usuario desde la línea de comandos

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.
- 2 Elija el comando adecuado:
 - Para cambiar las autorizaciones, las funciones o los perfiles de derechos asignados a un usuario definido en el ámbito local, utilice el comando `usermod(1M)`.
 - Otra posibilidad es editar el archivo `user_attr`.
Use este método sólo en caso de emergencia.
 - Para cambiar las autorizaciones, las funciones o los perfiles de derechos asignados a un usuario definido en un servicio de nombres, utilice el comando `smuser(1M)`.
Este comando requiere la autenticación por parte de un superusuario o de una función que pueda modificar archivos de usuario. Puede aplicar el comando `smuser` a todos los servicios de nombres. El comando `smuser` se ejecuta como cliente del servidor de Solaris Management Console.

Cierre y arranque de un clúster

En este capítulo se ofrece información sobre las tareas de cierre y arranque de un clúster global, un clúster de zona y determinados nodos. Para obtener información sobre cómo iniciar una zona no global, consulte el [Capítulo 18, “Planificación y configuración de zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

- “Descripción general sobre el cierre y el arranque de un clúster” en la página 59
- “Cierre y arranque de un solo nodo de un clúster” en la página 69
- “Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad” en la página 82

Para consultar una descripción pormenorizada de los procedimientos tratados en este capítulo, consulte [“Rearranque de un nodo en un modo que no sea de clúster”](#) en la página 79 y la [Tabla 3–2](#).

Descripción general sobre el cierre y el arranque de un clúster

El comando de Oracle Solaris Cluster `cluster(1CL) shutdown` detiene los servicios del clúster global de forma correcta y controlada, para cerrar sin sobresaltos un clúster global por completo. Puede utilizar el comando `cluster shutdown` al desplazar la ubicación de un clúster global o para cerrar el clúster global si un error de aplicación está dañando los datos. El comando `clzonecluster halt` detiene un clúster de zona que se ejecuta en un determinado nodo o en un clúster de zona completo en todos los nodos configurados. También puede utilizar el comando `cluster shutdown` con los clústers de zona.

En los procedimientos de este capítulo, `phys - schost#` refleja una solicitud de clúster global. El indicador de shell interactivo de `clzonecluster` es `clzc:schost>`.

Nota – Use el comando `cluster shutdown` para asegurarse de que todo el clúster global se cierre correctamente. El comando `shutdown` de Oracle Solaris se emplea con el comando `clnode(1CL)` evacúate para cerrar nodos independientes. Si desea obtener más información, consulte “Cierre de un clúster” en la página 61 o “Cierre y arranque de un solo nodo de un clúster” en la página 69.

Los comandos `cluster shutdown` y `clzonecluster halt` detienen todos los nodos comprendidos en un clúster global o un clúster de zona respectivamente, al ejecutar las acciones siguientes:

1. Pone fuera de línea todos los grupos de recursos en ejecución.
2. Desmonta todos los sistemas de archivos del clúster correspondientes a un clúster global o uno de zona.
3. El comando `cluster shutdown` cierra los servicios de dispositivos activos en el clúster global o de zona.
4. El comando `cluster shutdown` ejecuta `init 0`; en los sistemas basados en SPARC, lleva a todos los nodos del clúster al indicador de solicitud OpenBoot PROM `ok` o al mensaje `Press any key to continue` (Pulse cualquier tecla para continuar) del menú de GRUB en el caso de los sistemas basados en x86. Los menús de GRUB se describen más pormenorizadamente en “[Booting an x86 Based System by Using GRUB \(Task Map\)](#)” de *System Administration Guide: Basic Administration*. El comando `clzonecluster halt` ejecuta la operación del comando `zoneadm -z nombre_clúster_zona halt` para detener (pero no cerrar) las zonas del clúster de zona.

Nota – Si es necesario, tiene la posibilidad de arrancar un nodo en un modo que no sea de clúster para que el nodo no participe como miembro en el clúster. Los modos que no son de clúster son útiles al instalar software de clúster o para efectuar determinados procedimientos administrativos. Si desea obtener más información, consulte “[Rearranque de un nodo en un modo que no sea de clúster](#)” en la página 79.

TABLA 3–1 Lista de tareas: cerrar e iniciar un clúster

Tarea	Instrucciones
Detener el clúster	“Cierre de un clúster” en la página 61
Iniciar el clúster arrancando todos los nodos. Los nodos deben contar con una conexión operativa con la interconexión de clúster para conseguir convertirse en miembros del clúster.	“Arranque de un clúster” en la página 63
Rearranque el clúster.	“Rearranque de un clúster” en la página 65

▼ Cierre de un clúster

Puede cerrar un clúster global, uno de zona o todos los clústers de zona.



Precaución – No use el comando `send brk` en la consola de un clúster para cerrar un nodo del clúster global ni un nodo de un clúster de zona. El comando no puede utilizarse dentro de un clúster.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Si el clúster global o el de zona ejecutan Oracle Real Application Clusters (RAC), cierre todas las instancias de base de datos del clúster que va a cerrar.
Consulte la documentación del producto de Oracle RAC para ver los procedimientos de cierre.
- 2 Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` en todos los nodos del clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.
- 3 Cierre el clúster global, el de zona o todos los clústers de zona.

- Cierre el clúster global. Esta acción cierra también todos los clústers de zona.

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

- Cierre un clúster de zona concreto.

```
phys-schost# clzonecluster halt zoneclustername
```

- Cierre todos los clústers de zona.

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

El comando `cluster shutdown` también puede usarse dentro de un clúster de zona para cerrar todos los clústers de zona.

- 4 **Compruebe que todos los nodos del clúster global o el clúster de zona muestren el indicador ok en los sistemas basados en SPARC o un menú de GRUB en los sistemas basados en x86.**

No cierre ninguno de los nodos hasta que todos muestren el indicador ok en los sistemas basados en SPARC o se hallen en un subsistema de arranque en los sistemas basados en x86.

- **Compruebe que los nodos del clúster global muestren el indicador ok en los sistemas basados en SPARC o el mensaje Press any key to continue (Pulse cualquier tecla para continuar) del menú de GRUB en los sistemas basados en x86.**

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

- **Use el subcomando status para comprobar que el clúster de zona se haya cerrado.**

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

- 5 **Si es necesario, cierre los nodos del clúster global.**

Ejemplo 3-1 Cierre de un clúster de zona

En el ejemplo siguiente se cierra un clúster de zona denominado *zona_sc_dispersa*.

```
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 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 2 of cluster 'sparse-sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 4 of cluster 'sparse-sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone' died.
Sep  5 19:06:01 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster 'sparse-sczone' died.
phys-schost#
```

Ejemplo 3-2 SPARC: Cierre de un clúster global

En el ejemplo siguiente se muestra la salida de una consola cuando se detiene el funcionamiento normal de un clúster global y se cierran todos los nodos, lo que permite que se muestre el indicador ok. La opción `-g 0` establece el período de gracia para el cierre en cero; la opción `-y` proporciona una respuesta yes automática para la pregunta de confirmación. Los mensajes de cierre también aparecen en las consolas de los otros nodos del clúster global.

```
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
```

Ejemplo 3-3 x86: Cierre de un clúster global

En el ejemplo siguiente se muestra la salida de la consola al detenerse el funcionamiento normal del clúster global y cerrarse todos los nodos. En este ejemplo, el indicador ok no se aparece en todos los nodos. La opción `-g 0` establece el período de gracia para el cierre en cero; la opción `-y` proporciona una respuesta `yes` automática para la pregunta de confirmación. Los mensajes de cierre también aparecen en las consolas de los otros nodos del clúster global.

```
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
```

Véase también Consulte [“Arranque de un clúster” en la página 63](#) para reiniciar un clúster global o un clúster de zona que se ha cerrado.

▼ Arranque de un clúster

Este procedimiento explica cómo arrancar un clúster global o uno de zona cuyos nodos se han cerrado. En los nodos de un clúster global, el sistema muestra el indicador ok en los sistemas SPARC o el mensaje `Press any key to continue` (Pulse cualquier tecla para continuar) en los sistemas x86 basados en GRUB.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Nota – Para crear un clúster de zona, siga las instrucciones de [“Configuración de un clúster de zona” de Guía de instalación del software Oracle Solaris Cluster](#).

1 Arranque todos los nodos en modo de clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.

- En los sistemas basados en SPARC, ejecute el comando siguiente.

```
ok boot
```

- En los sistemas basados en x86, ejecute los comandos siguientes.

Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                     |
|                                                       |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

Nota – Los nodos deben disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembros del clúster.

- Si tiene un clúster de zona, puede arrancar el clúster de zona completo.

```
phys-schost# clzonecluster boot zoneclustername
```

- Si tiene más de un clúster de zona, puede arrancar todos los clústers de zona. Use + en lugar de *nombre_clúster_zona*.

2 Compruebe que los nodos se hayan arrancado sin errores y que estén en línea.

El comando de estado `cluster(1CL)` informa sobre el estado de los nodos del clúster global.

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

Al ejecutar el comando de estado `clzonecluster(1CL)` desde un nodo del clúster global, este comando informa sobre el estado del nodo del clúster de zona.

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


Nota – Si el sistema de archivos /var de un sistema alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en dicho nodo. Si surge este problema, consulte [“Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad” en la página 83.](#)

Ejemplo 3–4 SPARC: Arranque de un clúster global

El ejemplo siguiente muestra la salida de la consola al arrancarse el nodo `phys-schost-1` en el clúster global. Aparecen mensajes similares en las consolas de los otros nodos del clúster global. Si la propiedad de arranque automático de un clúster de zona se establece en `true`, el sistema arranca el nodo del clúster de zona de forma automática tras haber arrancado el nodo del clúster global en ese equipo.

Al rearrancarse un nodo del clúster global, todos los nodos de clúster de zona presentes en ese equipo se detienen. Todos los nodos de clúster de zona de ese equipo con la propiedad de arranque automático establecida en `true` se arrancan tras volver a arrancarse el nodo del clúster global.

```
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.
...
```

▼ Rearranque de un clúster

Para cerrar un clúster global, ejecute el comando `cluster shutdown` y luego arranque el clúster global aplicando el comando `boot` en todos los nodos. Para cerrar un clúster de zona, use el comando `clzonecluster halt`; después, ejecute el comando `clzonecluster boot` para arrancar el clúster de zona. También puede usar el comando `clzonecluster reboot`. Si desea obtener más información, consulte las páginas de comando `man cluster(1CL)`, `boot(1M)` y `clzonecluster(1CL)`.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Si el clúster ejecuta Oracle RAC, cierre todas las instancias de base de datos del clúster que va a cerrar.**

Consulte la documentación del producto de Oracle RAC para ver los procedimientos de cierre.

- 2 **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` en todos los nodos del clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

- 3 **Cierre el clúster.**

- **Cierre el clúster global.**

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

- **Si tiene un clúster de zona, ciérrelo desde un nodo del clúster global.**

```
phys-schost# clzonecluster halt zoneclustername
```

Se cierran todos los nodos. Para cerrar el clúster de zona también puede usar el comando `cluster shutdown` dentro de un clúster de zona.

Nota – Los nodos deben disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembros del clúster.

- 4 **Arranque todos los nodos.**

No importa el orden en que se arranquen los nodos, a menos que haga modificaciones de configuración entre operaciones de cierre. Si modifica la configuración entre operaciones de cierre, inicie primero el nodo con la configuración más actual.

- Para un nodo del clúster global que esté en un sistema basado en SPARC, ejecute el comando siguiente.

```
ok boot
```

- Para un nodo del clúster global que esté en un sistema basado en x86, ejecute los comandos siguientes.

Cuando aparezca el menú GRUB, seleccione la entrada del sistema operativo Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
+-----+
```

```
| Solaris failsafe |
|-----|
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

Nota – Los nodos deben disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembros del clúster.

Si desea obtener más información sobre las operaciones de inicio basadas en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- En el caso de un clúster de zona, para arrancar el clúster de zona, escriba el comando siguiente en un único nodo del clúster global.

```
phys-schost# clzonecluster boot zoneclustername
```

A medida que se activan los componentes del clúster, aparecen mensajes en las consolas de los nodos que se han arrancado.

5 Compruebe que los nodos se hayan arrancado sin errores y que estén en línea.

- El comando `clnode status` informa sobre el estado de los nodos del clúster global.

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

- Si ejecuta el comando `clzonecluster status` en un nodo del clúster global, se informa sobre el estado de los nodos de los clústers de zona.

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

También puede ejecutar el comando `cluster status` en un clúster de zona para ver el estado de los nodos.

Nota – Si el sistema de archivos `/var` de un sistema alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en dicho nodo. Si surge este problema, consulte [“Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad”](#) en la página 83.

Ejemplo 3–5 Rearranque de un clúster de zona

El ejemplo siguiente muestra cómo detener y arrancar un clúster de zona denominado `zona_sc_dispersa`. También puede usar el comando `clzonecluster reboot`.

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

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

=== Zone Clusters ===

--- Zone Cluster Status ---

Name	Node Name	Zone HostName	Status	Zone Status
-----	-----	-----	-----	-----
sparse-szone	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#

Ejemplo 3-6 SPARC: Rearranque de un clúster global

El ejemplo siguiente muestra la salida de la consola al detenerse el funcionamiento normal del clúster global, todos los nodos se cierran y muestran el indicador ok, y se reinicia el clúster global. La opción -g 0 establece el período de gracia en cero; la opción -y proporciona una respuesta yes automática para la pregunta de confirmación. Los mensajes de cierre también aparecen en las consolas de otros nodos del clúster global.

```
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:

```

Cierre y arranque de un solo nodo de un clúster

Puede cerrar un nodo del clúster global, un nodo de un clúster de zona o una zona no global. Esta sección proporciona instrucciones para cerrar nodos del clúster global y nodos de clústers de zona.

Para cerrar un nodo del clúster global, use el comando `clnode evacuate` con el comando `shutdown` de Oracle Solaris. Use el comando `cluster shutdown` sólo cuando vaya a cerrar todo un clúster global.

En el caso de un nodo de un clúster de zona, use el comando `clzonecluster halt` en un clúster global para cerrar sólo un nodo de un clúster de zona o todo un clúster de zona. También puede usar los comandos `clnode evacuate` y `shutdown` para cerrar nodos de un clúster de zona.

Si desea obtener más información sobre operaciones de cierre e inicio de una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*. Consulte también `clnode(1CL)`, `shutdown(1M)` y `clzonecluster(1CL)`.

En los procedimientos tratados en este capítulo, `phys-schost#` refleja una solicitud de clúster global. El indicador de solicitud de shell interactivo de `clzonecluster` es `clzc:schost>`.

TABLA 3-2 Mapa de tareas: cerrar y arrancar un nodo

Tarea	Herramienta	Instrucciones
Detener un nodo	Para los nodos del clúster global, use <code>clnode(1CL)</code> evacuate y shutdown. Para los nodos de un clúster de zona, use <code>clzonecluster(1CL)</code> halt.	“Cierre de un nodo” en la página 70
Arrancar un nodo	Para los nodos del clúster global, use boot o b. Para los nodos de un clúster de zona, use <code>clzonecluster(1CL)</code> boot.	“Arranque de un nodo” en la página 73
El nodo debe disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembro de este último.		
Detener y reiniciar (rearrancar) un nodo de un clúster	Para un nodo del clúster global, use <code>clnode</code> evacuate y shutdown, seguidos de boot o b. Para un nodo de un clúster de zona, use <code>clzonecluster(1CL)</code> reboot.	“Rearranque de un nodo” en la página 76
El nodo debe disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembro de este último.		
Arrancar un nodo de manera que no participe como miembro en el clúster	Para un nodo del clúster global, use los comandos <code>clnode</code> evacuate y shutdown, seguidos de <code>boot -x</code> en SPARC o en la edición de entradas del menú de GRUB en x86. Si el clúster global subyacente se arranca en un modo que no sea de clúster, el nodo del clúster de zona pasa automáticamente al modo que no sea de clúster.	“Rearranque de un nodo en un modo que no sea de clúster” en la página 79

▼ Cierre de un nodo

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.



Precaución – No use `send brk` en una consola de clúster para cerrar un nodo del clúster global ni de un clúster de zona. El comando no puede utilizarse dentro de un clúster.

- 1 **Si el clúster ejecuta Oracle RAC, cierre todas las instancias de base de datos del clúster que va a cerrar.**

Consulte la documentación del producto de Oracle RAC para ver los procedimientos de cierre.

- 2 **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` sobre el nodo del clúster que se va a cerrar. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

- 3 **Si desea detener un determinado miembro de un clúster de zona, omita los pasos del 4 al 6 y ejecute el comando siguiente desde un nodo del clúster global:**

```
phys-schost# clzonecluster halt -n physical-name zoneclustername
```

Al especificar un determinado nodo de un clúster de zona, sólo se detiene ese nodo. El comando `halt` detiene de forma predeterminada los clústers de zona en todos los nodos.

- 4 **Conmute todos los grupos de recursos, recursos y grupos de dispositivos del nodo que se va a cerrar a otros miembros del clúster global.**

En el nodo del clúster global que se va a cerrar, escriba el comando siguiente. El comando `clnode evacuate` conmuta todos los grupos de recursos y de dispositivos (incluidas las zonas no globales) del nodo especificado al siguiente nodo por orden de preferencia. También puede ejecutar `clnode evacuate` dentro de un nodo de un clúster de zona.

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

nodo Especifica el nodo desde el que esté conmutando los grupos de recursos y de dispositivos.

- 5 **Cierre el nodo.**

Especifique el nodo del clúster global que quiera cerrar.

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

Compruebe que el nodo del clúster global muestre el indicador ok en los sistemas basados en SPARC o el mensaje Press any key to continue (Pulse cualquier tecla para continuar) en el menú de GRUB de los sistemas basados en x86.

- 6 **Si es necesario, cierre el nodo.**

Ejemplo 3-7 SPARC: Cierre de nodos del clúster global

El ejemplo siguiente muestra la salida de la consola cuando se cierra el nodo `phys-schost-1`. La opción `-g0` establece el período de gracia en cero; la opción `-y` proporciona una respuesta `yes` automática para la pregunta de confirmación. Los mensajes de cierre correspondientes a este nodo también aparecen en las consolas de los otros nodos del clúster global.

```
phys-schost# clnode evacuate nodename
phys-schost# 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.
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
```

Ejemplo 3-8 x86: Cierre de nodos del clúster global

El ejemplo siguiente muestra la salida de la consola cuando se cierra el nodo `phys-schost-1`. La opción `-g0` establece el período de gracia en cero; la opción `-y` proporciona una respuesta `yes` automática para la pregunta de confirmación. Los mensajes de cierre correspondientes a este nodo también aparecen en las consolas de los otros nodos del clúster global.

```
phys-schost# clnode evacuate phys-schost-1
phys-schost# shutdown -g0 -y
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
```

Ejemplo 3-9 Cierre de un nodo de un clúster de zona

El ejemplo siguiente muestra el uso de `clzonecluster halt` para cerrar un nodo presente en un clúster de zona denominado *zona_sc_dispersa*. En un nodo de un clúster de zona también se pueden ejecutar los comandos `clnode evacuate` y `shutdown`.

```
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#
```

Véase también Consulte [“Arranque de un nodo” en la página 73](#) para ver cómo reiniciar un nodo del clúster global que se haya cerrado.

▼ Arranque de un nodo

Si desea cerrar o reiniciar otros nodos activos del clúster global o del clúster de zona, espere a que aparezca en línea el estado guía de servidor multiusuario para el nodo que está iniciando.

De lo contrario, el nodo no estará disponible para hacerse cargo de los servicios de otros nodos del clúster que se cierran o rearranquen. Si desea obtener información sobre cómo iniciar una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

Nota – La configuración del quórum puede afectar el inicio de un nodo. En un clúster de dos nodos, debe tener el dispositivo de quórum configurado de manera que el número total de quórum correspondiente al clúster ascienda a tres. Es conveniente tener un número de quórum para cada nodo y un número de quórum para el dispositivo de quórum. De esta forma, si cierra el primer nodo, el segundo sigue disponiendo de quórum y funciona como miembro único del clúster. Para que el primer nodo retorne al clúster como nodo integrante, el segundo debe estar operativo y en ejecución. Debe estar el correspondiente número de quórum de clúster (dos).

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Nota – Los nodos deben disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembros del clúster.

1 Para iniciar un nodo del clúster global o un nodo de un clúster de zona que se haya cerrado, arranque el nodo. Siga todos los pasos de este procedimiento desde un nodo del clúster global.

- En los sistemas basados en SPARC, ejecute el comando siguiente.

ok **boot**

- En los sistemas basados en x86, ejecute los comandos siguientes.

Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                     |
|                                                       |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

A medida que se activan los componentes del clúster, aparecen mensajes en las consolas de los nodos que se han arrancado.

- Si tiene un clúster de zona, puede elegir un nodo para que arranque.

```
phys-schost# clzonecluster boot -n node zoneclustername
```

2 Compruebe que el nodo haya arrancado sin errores y esté en línea.

- Al ejecutarlo, el comando `cluster status` informa sobre el estado de un nodo del clúster global.

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

- Al ejecutarlo desde un nodo del clúster global, el comando `clzonecluster status` informa sobre el estado de todos los nodos de clústers de zona.

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

Un nodo de un clúster de zona sólo puede arrancarse en modo de clúster si el nodo que lo aloja arranca en modo de clúster.

Nota – Si el sistema de archivos `/var` de un sistema alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en dicho nodo. Si surge este problema, consulte [“Reparación de un sistema de archivos `/var` que haya alcanzado el límite de capacidad” en la página 83.](#)

Ejemplo 3–10 SPARC: Arranque de un nodo del clúster global

El ejemplo siguiente muestra la salida de la consola cuando se arranca el nodo `phys-schost-1` en el clúster global.

```
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:
```

▼ Rearranque de un nodo

Para cerrar o reiniciar otros nodos activos en el clúster global o en el clúster de zona, espere a que aparezca en línea el estado guía de servidor multiusuario para el nodo que está iniciando.

De lo contrario, el nodo no estará disponible para hacerse cargo de los servicios de otros nodos del clúster que se cierran o rearranquen. Si desea obtener información sobre cómo reiniciar una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Si el nodo del clúster global o del clúster de zona está ejecutando Oracle RAC, cierre todas las instancias de base de datos presentes en el nodo que va a cerrar.**

Consulte la documentación del producto de Oracle RAC para ver los procedimientos de cierre.

- 2 **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` sobre el nodo que vaya a cerrar. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**
- 3 **Cierre el nodo del clúster global con los comandos `clnode evacuate` y `shutdown`. Cierre el clúster de zona mediante la ejecución del comando `clzonecluster halt` desde un nodo del clúster global. Los comandos `clnode evacuate` y `shutdown` también funcionan en los clústers de zona.**

En el caso de un clúster global, escriba los comandos siguientes en el nodo que vaya a cerrar. El comando `clnode evacuate` conmuta todos los grupos de dispositivos desde el nodo especificado al siguiente nodo por orden de preferencia. Este comando también conmuta todos los grupos de recursos de las zonas globales y no globales del nodo especificado a las zonas globales o no globales de otros nodos que se sitúen a continuación en el orden de preferencia.

- Ejecute los comandos siguientes en un sistema basado en SPARC.

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

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

- Ejecute los comandos siguientes en un sistema basado en x86.

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

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

Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                |
| Solaris failsafe                      |
|                                     |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

- Indique el nodo del clúster de zona que se vaya a cerrar y rearmar.

```
phys-schost# clzonecluster reboot - node zoneclustername
```

Nota – Los nodos deben disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembros del clúster.

4 Compruebe que el nodo haya arrancado sin errores y esté en línea.

- Compruebe que el nodo del clúster global esté en línea.

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

- Compruebe que el nodo del clúster de zona esté en línea.

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

Ejemplo 3–11 SPARC: Rearranque de un nodo del clúster global

El ejemplo siguiente muestra la salida de la consola cuando se rearmar el nodo `phys-schost-1`. Los mensajes correspondientes a este nodo, como las notificaciones de cierre e inicio, aparecen en las consolas de los otros nodos del clúster global.

```
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:
```

Ejemplo 3–12 x86: Rearranque de un nodo del clúster global

El ejemplo siguiente muestra la salida de la consola al rearrancar el nodo `phys-schost-1`. Los mensajes correspondientes a este nodo, como las notificaciones de cierre e inicio, aparecen en las consolas de los otros nodos del clúster global.

```
phys-schost# clnode evacuate phys-schost-1
phys-schost # shutdown -y -g0 -i6

GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                     |
|                                                       |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
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:
```

Ejemplo 3–13 Rearranque de un nodo del clúster global

El ejemplo siguiente muestra cómo rearrancar un nodo de un clúster de zona.

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

phys-schost#
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---
Name           Node Name      Zone HostName   Status   Zone Status
-----
sparse-sczzone 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#
```

▼ Rearranque de un nodo en un modo que no sea de clúster

Puede arrancar un nodo del clúster global en un modo que no sea de clúster, en el cual el nodo no participa como miembro del clúster. El modo que no es de clúster resulta útil a la hora de instalar el software del clúster o en ciertos procedimientos administrativos como la aplicación de parches en un nodo. Un nodo de un clúster de zona no puede estar en un estado de arranque diferente del que tenga el nodo del clúster global subyacente. Si el nodo del clúster global se arranca en un modo que no sea de clúster, el nodo del clúster de zona asume automáticamente el modo que no es de clúster.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC solaris.cluster.admin sobre el clúster que se va a iniciar en un modo que no es de clúster. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**

- 2 **Cierre el nodo del clúster de zona con el comando `clzonecluster halt` en un nodo del clúster global. Cierre el nodo del clúster global con los comandos `clnode evacuate` y `shutdown`.**

El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. Este comando también conmuta todos los grupos de recursos de las zonas globales y no globales del nodo especificado y los lleva a las zonas globales o no globales de otros nodos que se sitúen a continuación en el orden de preferencia.

- **Cierre un clúster global en concreto.**

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

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

- **Cierre un nodo del clúster de zona en concreto a partir de un nodo del clúster global.**

```
phys-schost# clzonecluster halt -n node zoneclustername
```

También puede utilizar los comandos `clnode evacuate` y `shutdown` dentro del clúster de zona.

- 3 **Compruebe que el nodo del clúster global muestre el indicador ok en un sistema basado en Oracle Solaris o el mensaje Press any key to continue (Pulse cualquier tecla para continuar) en el menú de GRUB de un sistema basado en x86.**

- 4 **Arranque el nodo del clúster global en un modo que no sea de clúster.**

- En los sistemas basados en SPARC, ejecute el comando siguiente.

```
ok boot -xs
```

- En los sistemas basados en x86, ejecute los comandos siguientes.

- a. **En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba `e` para editar los comandos.**

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
```

```
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                       |
|                                                         |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)” de *System Administration Guide: Basic Administration*](#).

- b. En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba **e** para editarla.

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a)                                |
| kernel /platform/i86pc/multiboot              |
| module /platform/i86pc/boot_archive           |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.
```

- c. Agregue **-x** al comando para especificar que el sistema arranque en un modo que no sea de clúster.

[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/multiboot -x
```

- d. Pulse la tecla **Intro** para aceptar el cambio y volver a la pantalla de los parámetros de arranque.

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a)                                |
| kernel /platform/i86pc/multiboot -x          |
| module /platform/i86pc/boot_archive           |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

- e. Escriba **b** para arrancar el nodo en un modo que no sea de clúster.

Nota – Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en el modo que no es de clúster, realice estos pasos de nuevo para agregar la opción **-x** al comando del parámetro de arranque del núcleo.

Ejemplo 3–14 SPARC: Arranque de un nodo del clúster global en un modo que no sea de clúster

El ejemplo siguiente muestra la salida de la consola cuando el nodo `phys-schost-1` se cierra y se reinicia en un modo que no es de clúster. La opción `-g0` establece el período de gracia en cero; la opción `-y` proporciona una respuesta yes automática para la pregunta de confirmación y la opción `-i0` invoca el nivel de ejecución 0 (cero). Los mensajes de cierre correspondientes a este nodo también aparecen en las consolas de los otros nodos del clúster global.

```
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:
```

Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad

El software de Oracle Solaris y el de Oracle Solaris Cluster dejan constancia de los mensajes de error en el archivo `/var/adm/messages`, que con el paso del tiempo puede llegar a ocupar totalmente el sistema de archivos `/var`. Si el sistema de archivos `/var` alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en ese nodo. Además, es posible que no pueda iniciarse sesión en ese nodo.

▼ Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad

Si un nodo emite un informe que anuncia que el sistema de archivos /var ha alcanzado su límite de capacidad y continúa ejecutando servicios de Oracle Solaris Cluster, use este procedimiento para borrar el sistema de archivos. Si desea obtener más información, consulte [“Viewing System Messages” de System Administration Guide: Advanced Administration](#) en [System Administration Guide: Advanced Administration](#).

- 1 Conviértase en superusuario en el nodo del clúster cuyo sistema de archivos /var haya alcanzado el límite de capacidad.**

- 2 Borre todo el sistema de archivos.**

Por ejemplo, elimine los archivos prescindibles que haya en dicho sistema de archivos.

Métodos de replicación de datos

Este capítulo describe las tecnologías de replicación de datos que puede usar con el software Oracle Solaris Cluster. La *replicación de datos* es un procedimiento que consiste en copiar datos de un dispositivo de almacenamiento primario a un dispositivo secundario o de copia de seguridad. Si el dispositivo primario falla, los datos están disponibles en el secundario. La replicación de datos asegura alta disponibilidad y tolerancia ante errores graves del clúster.

El software Oracle Solaris Cluster es compatible con los tipos de replicación de datos siguientes:

- Entre clústeres: use Oracle Solaris Cluster Geographic Edition para recuperar datos en caso de problema grave.
- En un clúster: úselo como sustitución de la duplicación basada en host en un clúster de campus

Para llevar a cabo la replicación de datos, debe haber un grupo de dispositivos con el mismo nombre que el objeto que vaya a replicar. Un dispositivo sólo puede pertenecer a un grupo de dispositivos a la vez; por eso, si ya tiene un grupo de dispositivos de Oracle Solaris Cluster que contiene el dispositivo, debe eliminar este grupo antes de agregar el dispositivo al nuevo grupo de dispositivos. Para obtener instrucciones sobre cómo crear y administrar grupos de dispositivos de disco de Solaris Volume Manager, Veritas Volume Manager, ZFS o discos básicos, consulte [“Administración de grupos de dispositivos” en la página 123](#) en el Capítulo 5.

Antes de seleccionar el método de replicación más apropiado para el clúster, es necesario comprender la replicación de datos basada en host y en almacenamiento. Para obtener más información sobre cómo usar Oracle Solaris Cluster Geographic Edition para administrar la replicación de datos para recuperación de problemas graves, consulte [Oracle Solaris Cluster Geographic Edition Overview](#).

Este capítulo incluye las secciones siguientes:

- [“Comprensión de la replicación de datos” en la página 86](#)
- [“Uso de replicación de datos basada en almacenamiento dentro de un clúster” en la página 88](#)

Comprensión de la replicación de datos

Oracle Solaris Cluster admite los métodos de replicación de datos siguientes:

- *La replicación de datos basada en host* usa el software para replicar en tiempo real volúmenes de discos entre clústers ubicados en puntos geográficos distintos. La replicación por duplicación remota permite replicar los datos del volumen maestro del clúster primario en el volumen maestro del clúster secundario separado geográficamente. Un mapa de bits de duplicación remota controla las diferencias entre el volumen maestro del disco primario y el del secundario. Un ejemplo de software de replicación basada en host que se usa para la replicación entre clústeres (y entre un clúster y un host que no se encuentre en el clúster) es Sun StorageTek Availability Suite 4.

La replicación de datos basada en host es una solución de replicación más económica, porque usa recursos del host en lugar de matrices de almacenamiento especiales. No se admiten las bases de datos, las aplicaciones ni los sistemas de archivos configurados para permitir que varios hosts ejecuten el sistema operativo Oracle Solaris para escribir datos en un volumen compartido (por ejemplo, Oracle 9iRAC y Oracle Parallel Server). Para obtener más información sobre cómo usar una replicación de datos basada en host entre dos clústeres, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Sun StorageTek Availability Suite](#). Para ver un ejemplo de replicación basada en host que no use Oracle Solaris Cluster Geographic Edition, consulte el A, “[Configuración de replicación de datos basada en host con el software Sun StorageTek Availability Suite](#)” en la página 371.

- *La replicación de datos basada en almacenamiento* utiliza el software del controlador de almacenamiento para sacar el trabajo de la replicación de datos de los nodos del clúster y moverlo al dispositivo de almacenamiento. Este software libera capacidad de procesamiento del nodo para atender las solicitudes del clúster. Entre los ejemplos de software basado en almacenamiento que puede replicar datos dentro del clúster o entre clústeres figuran Hitachi TrueCopy, Hitachi Universal Replicator y EMC SRDF. La replicación de datos basada en almacenamiento puede ser especialmente importante en configuraciones de clústeres de campus y puede simplificar la infraestructura necesaria. Para obtener más información sobre el uso de la replicación de datos basada en almacenamiento en un entorno de clúster de campus, consulte “[Uso de replicación de datos basada en almacenamiento dentro de un clúster](#)” en la página 88.

Para obtener más información sobre el uso de la replicación basada en almacenamiento entre dos clústeres o más y el producto Oracle Solaris Cluster Geographic Edition que automatiza el proceso, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Hitachi TrueCopy and Universal Replicator](#) y [Oracle Solaris Cluster Geographic Edition Data Replication Guide for EMC Symmetrix Remote Data Facility](#). Consulte también el A, “[Configuración de replicación de datos basada en host con el software Sun StorageTek Availability Suite](#)” en la página 371 para obtener un ejemplo completo de este tipo de configuración del clúster.

Métodos admitidos de replicación de datos

El software Oracle Solaris Cluster admite los métodos siguientes de replicación de datos entre clústers o dentro de un mismo clúster:

1. Replicación entre clústers: para la recuperación de datos en caso de problema grave, puede usar una replicación basada en almacenamiento o en host con el objeto de realizar una replicación de datos entre clústers. Normalmente se opta por una replicación basada en host o en almacenamiento, en lugar de una combinación de las dos. Puede administrar ambos tipos de replicación con el software Oracle Solaris Cluster Geographic Edition.

- Replicación basada en host
 - Sun StorageTek Availability Suite , a partir del sistema operativo Oracle Solaris 10

Si desea utilizar la replicación basada en host con el software de Oracle Solaris Cluster Geographic Edition, consulte las instrucciones del [Apéndice A, “Ejemplo”, “Configuración de replicación de datos basada en host con el software Sun StorageTek Availability Suite” en la página 371.](#)

- Replicación basada en almacenamiento
 - Hitachi TrueCopy e Hitachi Universal Replicator, mediante Oracle Solaris Cluster Geographic Edition
 - EMC Symmetrix Remote Data Facility (SRDF), mediante Oracle Solaris Cluster Geographic Edition

Si desea utilizar una replicación basada en almacenamiento sin el software de Oracle Solaris Cluster Geographic Edition, consulte la documentación del software de replicación.

2. Replicación dentro de un clúster: este método se utiliza como sustitución para la duplicación basada en host.
 - Replicación basada en almacenamiento
 - Hitachi TrueCopy e Hitachi Universal Replicator
 - EMC Symmetrix Remote Data Facility (SRDF)
3. Replicación basada en aplicaciones: Oracle Data Guard es un ejemplo de software de replicación basada en aplicaciones. Este tipo de software se usa sólo para recuperar datos en caso de un problema grave. Para obtener más información, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard](#).

Uso de replicación de datos basada en almacenamiento dentro de un clúster

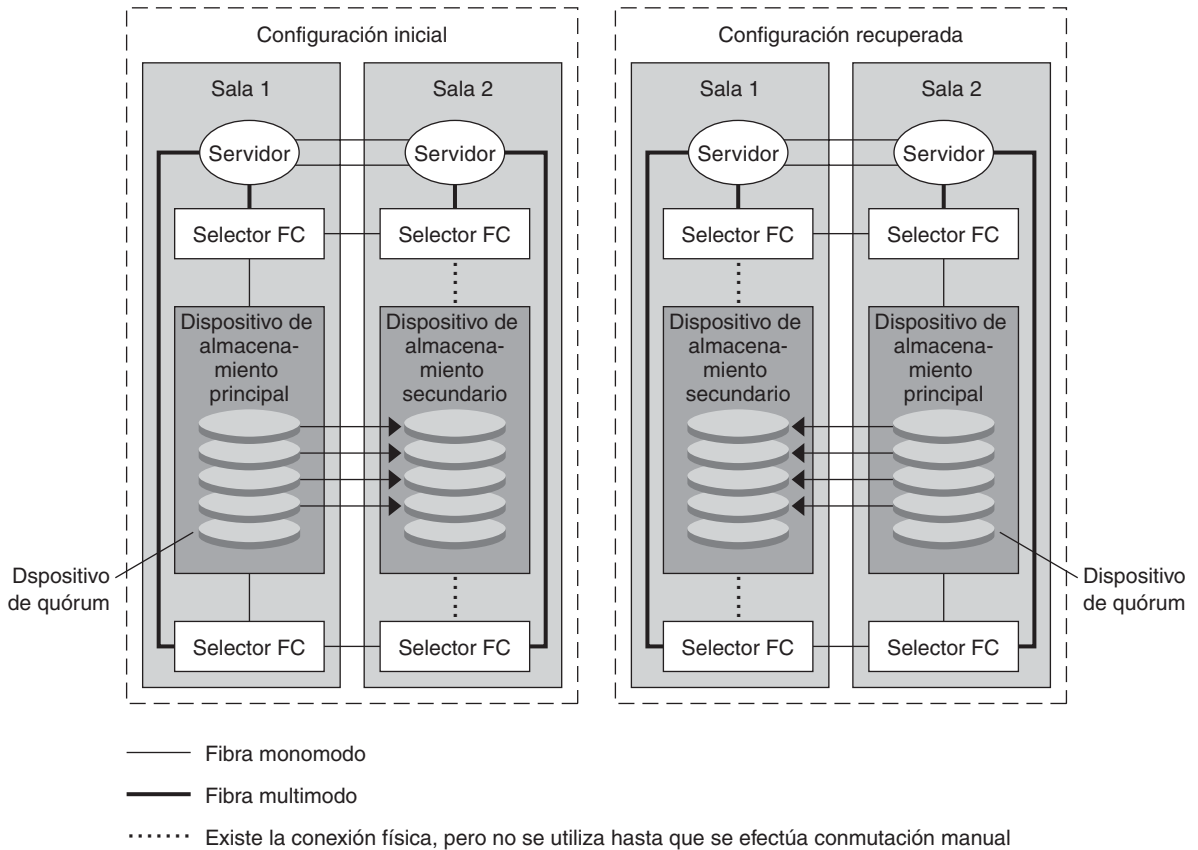
La replicación de datos basada en almacenamiento usa un software instalado en el dispositivo de almacenamiento para administrar la replicación dentro de un clúster o un clúster de campus. Dicho software es específico del dispositivo de almacenamiento particular y no se utiliza para recuperación en caso de un problema grave. Consulte la documentación suministrada con el dispositivo de almacenamiento al configurar la replicación de datos basada en almacenamiento.

En función del software que use, puede utilizar la migración tras error automática o manual con la replicación de datos basada en almacenamiento. Oracle Solaris Cluster admite la migración tras error automática o manual de los replicadores con Hitachi TrueCopy, Hitachi Universal Replicator y EMC SRDF.

Esta sección describe la replicación de datos basada en almacenamiento que se usa en un clúster de campus. La [Figura 4-1](#) muestra un ejemplo de configuración de dos salas donde los datos se replican entre dos matrices de almacenamiento. En esta configuración, la matriz de almacenamiento principal se encuentra en la primera sala, donde proporciona datos a los nodos de ambas salas. Esta matriz de almacenamiento principal también proporciona la matriz secundaria con datos para replicar.

Nota – La [Figura 4-1](#) muestra que el dispositivo de quórum está en un volumen sin replicar. Un volumen replicado no es válido como dispositivo de quórum.

FIGURA 4-1 Configuración de dos salas con replicación de datos basada en almacenamiento



Se puede llevar a cabo una replicación de datos basada en almacenamiento con Hitachi TrueCopy o Hitachi Universal Replicator de forma sincrónica o asincrónica en el entorno de Oracle Solaris Cluster, según el tipo de aplicación que utilice. Si desea efectuar una migración tras error automática en un clúster de campus, use TrueCopy de forma sincrónica. La replicación sincrónica basada en almacenamiento con EMC SRDF es compatible con Oracle Solaris Cluster; la replicación asincrónica no es compatible con EMC SRDF.

No use los modos Domino ni Adaptive Copy de EMC SRDF. El modo Domino hace que los volúmenes de SRDF locales y de destino no estén disponibles para el sistema cuando el destino no esté disponible. El modo Adaptive Copy se suele usar para migrar datos y desplazar centros de datos, pero no se recomienda para la recuperación de datos en caso de problema grave.

No use los modos Data ni Status en Hitachi TrueCopy ni Hitachi Universal Replicator. Si falla el dispositivo de almacenamiento secundario, podría tener problemas de escritura en el dispositivo de almacenamiento principal.

Requisitos y restricciones al usar replicación de datos basada en almacenamiento dentro de un clúster

Para asegurar la integridad de los datos, use múltiples rutas y el paquete adecuado de RAID. La lista siguiente incluye consideraciones sobre cómo implementar una configuración de clúster que utilice replicación de datos basada en almacenamiento.

- La distancia entre nodo y nodo está limitada por el canal de fibra de Oracle Solaris Cluster y la infraestructura de interconexión. Póngase en contacto con el proveedor de servicios de Oracle para obtener más información sobre las limitaciones actuales y las tecnologías admitidas.
- No configure un volumen replicado como dispositivo de quórum. Localice los dispositivos de quórum en un volumen compartido con replicación anulada o utilice un servidor de quórum.
- Compruebe que sólo la copia primaria de los datos esté visible para los nodos del clúster. De lo contrario, el administrador de volúmenes podría intentar acceder simultáneamente a las copias primaria y secundaria de los datos. Consulte la documentación suministrada con la matriz de almacenamiento para obtener información sobre el control de la visibilidad de las copias de datos.
- EMC SRDF, Hitachi TrueCopy e Hitachi Universal Replicator permiten que el usuario defina los grupos de dispositivos replicados. Todos los grupos de dispositivos de replicación requieren un grupo de dispositivos de Oracle Solaris Cluster con el mismo nombre.
- Los datos particulares específicos de una aplicación podrían no ser adecuados para una replicación de datos asincrónica. Según sus conocimientos sobre el comportamiento de la aplicación, determine la mejor forma de replicar los datos específicos de la aplicación en los dispositivos de almacenamiento.
- Si quiere configurar el clúster para que realice automáticamente migraciones tras error, use la replicación sincrónica.

Para obtener instrucciones sobre cómo configurar el clúster para que realice migraciones tras error automáticas de los volúmenes replicados, consulte [“Administración de dispositivos replicados basados en almacenamiento”](#) en la página 97.

- Oracle Real Application Clusters (RAC) no es compatible con SRDF, Hitachi TrueCopy ni Hitachi Universal Replicator cuando se replica dentro de un clúster. Los nodos conectados a réplicas que no sean la réplica primaria no tendrán acceso de escritura. Cualquier aplicación escalable que requiera acceso de escritura directo desde todos los nodos del clúster no puede ser compatible con dispositivos replicados.
- Veritas Cluster Volume Manager (CVM) y Solaris Volume Manager multipropietario no son compatibles con Oracle Solaris Cluster.
- No use los modos Domino o Adaptive Copy de EMC SRDF. Consulte [“Uso de replicación de datos basada en almacenamiento dentro de un clúster”](#) en la página 88 para obtener más información.

- No use los modos Data ni Status en Hitachi TrueCopy ni Hitachi Universal Replicator. Consulte “[Uso de replicación de datos basada en almacenamiento dentro de un clúster](#)” en la página 88 para obtener más información.

Problemas de recuperación manual al usar una replicación de datos basada en almacenamiento dentro de un clúster

Como sucede en todos los clústers de campus, los que utilizan replicación de datos basada en almacenamiento no suelen necesitar intervención cuando tienen un solo error. Sin embargo, si se usa migración tras error manual y se pierde el espacio que aloja el dispositivo de almacenamiento principal (como se muestra en la [Figura 4-1](#)), los problemas surgen en un nodo de dos clústers. El nodo que queda no puede reservar el dispositivo de quórum y no se puede arrancar como miembro del clúster. En este caso, el clúster requiere la intervención manual siguiente:

1. El proveedor de servicios de Oracle debe volver a configurar el nodo restante para que inicie como miembro del clúster.
2. El proveedor de servicios de Oracle debe configurar un volumen con replicación anulada del dispositivo de almacenamiento secundario como dispositivo de quórum.
3. El proveedor de servicios de Oracle debe configurar el nodo restante para que use el dispositivo de almacenamiento secundario como almacenamiento principal. Esta reconfiguración podría suponer la reconstrucción de volúmenes del administrador de volúmenes, la restauración de datos o el cambio de las asociaciones de aplicación con volúmenes de almacenamiento.

Prácticas recomendadas al usar la replicación de datos basada en almacenamiento

Cuando configure los grupos de dispositivos que usan Hitachi TrueCopy o Hitachi Universal Replicator para la replicación de datos basada en almacenamiento, siga estas prácticas:

- Utilice una replicación sincrónica para evitar la posibilidad de perder datos si falla la ubicación principal.
- Debe haber una relación unívoca entre el grupo de dispositivos globales de Oracle Solaris Cluster y el grupo de replications de TrueCopy definidos en el archivo de configuración `horcm`. De este modo, los dos grupos se mueven de un nodo a otro como una sola unidad.
- Los volúmenes de sistema de archivos globales y los del sistema de archivos de migración tras error no se pueden mezclar en el mismo grupo de dispositivos replicados porque se controlan de forma distinta. Los sistemas de archivos globales están controlados por un

sistema de configuración de dispositivos , mientras que los volúmenes de sistemas de archivos de migración tras error lo están por HAS+. El primario de cada uno podría ser un nodo diferente, lo que podría provocar conflictos sobre el nodo en el que debería estar la replicación primaria.

- Todas las instancias de RAID Manager deben estar siempre activadas y en funcionamiento.

Cuando utilice el software EMC SRDF para replicación de datos basada en almacenamiento, use dispositivos dinámicos en lugar de estáticos. Los dispositivos estáticos necesitan varios minutos para cambiar el nodo primario de replicación y pueden afectar al tiempo de migración tras error.

Administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de clúster

En este capítulo se ofrece información sobre los procedimientos de administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de clúster.

- “Información general sobre la administración de dispositivos globales y el espacio de nombre global” en la página 93
- “Administración de dispositivos replicados basados en almacenamiento” en la página 97
- “Información general sobre la administración de sistemas de archivos de clústers” en la página 121
- “Administración de grupos de dispositivos” en la página 123
- “Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento” en la página 168
- “Administración de sistemas de archivos de clúster” en la página 174
- “Administración de la supervisión de rutas de disco” en la página 180

Si desea ver una descripción detallada de los procedimientos relacionados con el presente capítulo, consulte la [Tabla 5–4](#).

Si desea obtener información conceptual sobre los dispositivos globales, el espacio de nombre global, los grupos de dispositivos, la supervisión de las rutas de disco y el sistema de archivos de clúster, consulte *Oracle Solaris Cluster Concepts Guide*.

Información general sobre la administración de dispositivos globales y el espacio de nombre global

La administración de grupos de dispositivos de Oracle Solaris Cluster depende del administrador de volúmenes que esté instalado en el clúster. Solaris Volume Manager está habilitado para trabajar con clústers para poder agregar, registrar y eliminar grupos de dispositivos mediante el comando `metaset(1M)` de Solaris Volume Manager. Si utiliza Veritas Volume Manager (VxVM), puede crear grupos de discos mediante los comandos de VxVM. Registre los grupos de discos como grupos de dispositivos de Oracle Solaris Cluster con la

utilidad `clsetup`. Al eliminar los grupos de dispositivos de VxVM, puede usar tanto el comando `clsetup` como los comandos de VxVM.

Nota – No es posible acceder directamente a los dispositivos globales desde los nodos del clúster global que no sean de votación.

El software Oracle Solaris Cluster crea automáticamente un grupo de dispositivos de discos básicos para cada dispositivo de disco y de cinta del clúster. Sin embargo, los grupos de dispositivos del clúster permanecen fuera de línea hasta que se acceda a dichos grupos como dispositivos globales. Al administrar grupos de dispositivos o grupos de discos del administrador de volúmenes, es necesario estar en el nodo del clúster que sirve como nodo primario del grupo.

En general, no es necesario administrar el espacio de nombre del dispositivo global. El espacio de nombre global se establece automáticamente durante la instalación y se actualiza, también de manera automática, al reiniciar el sistema operativo Oracle Solaris. Sin embargo, si el espacio de nombre global debe actualizarse, el comando `cldevice populate` puede ejecutarse desde cualquier nodo del clúster. Este comando hace que el espacio de nombre global se actualice en todos los demás nodos miembros del clúster y en los nodos que posteriormente tengan la posibilidad de asociarse al clúster.

Permisos de dispositivos globales en Solaris Volume Manager

Las modificaciones en los permisos del dispositivo global no se propagan automáticamente a todos los nodos del clúster para Solaris Volume Manager y los dispositivos de disco. Si desea modificar los permisos de los dispositivos globales, los cambios deben hacerse de forma manual en todos los nodos del clúster. Por ejemplo, para modificar los permisos del dispositivo global `/dev/global/dsk/d3s0` y establecer un valor de 644, debe ejecutarse el comando siguiente en todos los nodos del clúster:

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

VxVM no admite el comando `chmod`. Para modificar los permisos del dispositivo global en VxVM, consulte la guía del administrador de VxVM.

Reconfiguración dinámica con dispositivos globales

A la hora de concluir operaciones de reconfiguración dinámica en dispositivos de disco y cinta de un clúster, deben tener en cuenta una serie de aspectos.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la reconfiguración dinámica de Oracle Solaris también son válidos para la reconfiguración dinámica en Oracle Solaris Cluster. La única excepción consiste en la inactividad del sistema operativo. Por lo tanto, consulte la documentación sobre la reconfiguración dinámica de Oracle Solaris *antes* de utilizarla con el software Oracle Solaris Cluster. Preste especial atención a los problemas que afecten a los dispositivos de E/S no conectados a red durante las operaciones de desconexión de reconfiguración dinámica.
- Oracle Solaris Cluster rechaza operaciones de eliminación de tarjetas de reconfiguración dinámica en dispositivos activos del nodo primario. Las operaciones de reconfiguración dinámica son factibles en los dispositivos inactivos del nodo primario y en cualquier dispositivo de los nodos secundarios.
- Tras la operación de reconfiguración dinámica, el acceso a los datos del clúster sigue igual que antes.
- Oracle Solaris Cluster no acepta operaciones de reconfiguración dinámica que repercutan en la disponibilidad de los dispositivos de quórum. Si desea obtener más información, consulte [“Reconfiguración dinámica con dispositivos de quórum” en la página 191](#).



Precaución – Si el nodo primario actual falla durante la operación de reconfiguración dinámica, dicho error repercute en la disponibilidad del clúster. El nodo primario no podrá migrar tras error hasta que se proporcione un nuevo nodo secundario.

Para realizar operaciones de reconfiguración dinámica en dispositivos globales, aplique los pasos siguientes en el orden que se indica.

TABLA 5-1 Mapa de tareas: reconfiguración dinámica con dispositivos de disco y cinta

Tarea	Instrucciones
1. Si en el nodo primario actual debe realizarse una operación de reconfiguración dinámica que afecta a un grupo de dispositivos activo, conmutar los nodos primario y secundario antes de ejecutar la operación en el dispositivo	“Conmutación al nodo primario de un grupo de dispositivos” en la página 165
2. Efectuar la operación de eliminación de DR en el dispositivo que se va a eliminar	<i>Sun Enterprise 10000 Dynamic Reconfiguration Reference Manual de la recopilación de Solaris 10 on Sun Hardware.</i>

Aspectos que tener en cuenta en la administración de Veritas Volume Manager

- Para que el software Oracle Solaris Cluster mantenga el espacio de nombre de VxVM, debe registrar todas las modificaciones realizadas en los volúmenes o grupos de discos de VxVM como cambios en la configuración de grupos de dispositivos de Oracle Solaris Cluster. El registro de modificaciones asegura la actualización del espacio de nombre en todos los nodos del clúster. Algunos ejemplos de modificaciones que repercuten sobre el espacio de nombre son agregar, quitar o cambiar el nombre de un volumen. Modificar los permisos del volumen, su propietario o el ID de grupo afecta también al espacio de nombre.

Nota – No importe ni deporte nunca grupos de discos de VxVM con comandos de VxVM una vez registrado el grupo de discos con el clúster como grupo de dispositivos de Oracle Solaris Cluster. El software Oracle Solaris Cluster se ocupa de todos los casos en que haya grupos de discos que deben importarse o deportarse.

- Cada grupo de discos de VxVM debe tener un número menor exclusivo para todo el clúster. De forma predeterminada, al crear un grupo de discos, VxVM elige un número aleatorio que sea múltiplo de 1.000 para que sirva como número menor base del grupo de discos. En la mayoría de las configuraciones que sólo tienen un número reducido de grupos de discos, este número menor es suficiente para garantizar la singularidad. El número menor de un grupo de discos acabado de crear podría entrar en conflicto con el número menor de un grupo de discos que ya existiera que se hubiese importado de otro nodo. En este caso, si se intenta registrar el grupo de dispositivos de Oracle Solaris Cluster, se producirá un error. Para solucionar este problema, asigne un nuevo número menor exclusivo al grupo de discos nuevo; a continuación, regístrelo como grupo de dispositivos de Oracle Solaris Cluster.
- Si está configurando un volumen duplicado, puede utilizar DRL (Dirty Region Logging) para reducir el tiempo de recuperación del volumen tras error de nodo. Se recomienda especialmente usar DRL, si bien es posible que repercuta en el rendimiento de E/S.
- VxVM no admite el comando `chmod`. Para modificar los permisos del dispositivo global en VxVM, consulte la guía del administrador de VxVM.
- El software Oracle Solaris Cluster 3.3 no admite la administración DMP (Dynamic Multipathing) de VxVM de varias rutas desde un mismo nodo.
- Si utiliza VxVM para configurar grupos de discos compartidos para Oracle RAC, use la función de clúster de VxVM como se describe en el documento *Veritas Volume Manager Administrator's Reference Guide*. Crear grupos de discos compartidos para Oracle RAC es distinto de crear otros grupos de discos. Debe importar los grupos de discos compartidos de Oracle RAC mediante `vxvg -s`. No registre los grupos de discos compartidos de Oracle RAC con la estructura del clúster. Para crear otros grupos de discos de VxVM, consulte [“Creación de un nuevo grupo de discos al inicializar discos \(Veritas Volume Manager\)” en la página 137](#).

Administración de dispositivos replicados basados en almacenamiento

Puede configurar un grupo de dispositivos de Oracle Solaris Cluster para que contenga dispositivos que se repliquen mediante la replicación basada en almacenamiento. Oracle Solaris Cluster admite el uso de Hitachi TrueCopy y EMC Symmetrix Remote Data Facility para la replicación basada en almacenamiento.

Antes de poder replicar datos con Hitachi TrueCopy o EMC Symmetrix Remote Data Facility, debe familiarizarse con la documentación sobre replicación basada en almacenamiento, y tener instalados en el sistema el producto de replicación basada en almacenamiento y sus parches más actuales. Si desea obtener información sobre cómo instalar el software de replicación basada en almacenamiento, consulte la documentación del producto.

El software de replicación basada en almacenamiento configura un par de dispositivos como réplicas: un dispositivo como réplica primaria y el otro como réplica secundaria. En un determinado momento, el dispositivo conectado a un conjunto de nodos convierte en la réplica primaria. El dispositivo conectado a otro conjunto de nodos es la réplica secundaria.

En la configuración de Oracle Solaris Cluster, la réplica primaria se desplaza automáticamente cada vez que se desplaza el grupo de dispositivos de Oracle Solaris Cluster al que pertenece. Por lo tanto, en una configuración de Oracle Solaris Cluster en principio la réplica primaria nunca se debe desplazar directamente. La toma de control debería tener lugar mediante el desplazamiento del grupo de dispositivos de Oracle Solaris Cluster asociado.



Precaución – El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

Esta sección incluye los procedimientos siguientes:

- [“Administración de dispositivos replicados mediante Hitachi TrueCopy” en la página 97](#)
- [“Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility” en la página 109](#)

Administración de dispositivos replicados mediante Hitachi TrueCopy

En la tabla siguiente figuran las tareas necesarias para configurar con Hitachi TrueCopy un dispositivo replicado basado en almacenamiento.

TABLA 5-2 Mapa de tareas: administrar un dispositivo replicado con Hitachi TrueCopy

Tarea	Instrucciones
Instalar el software de TrueCopy en los nodos y dispositivos de almacenamiento	Consulte la documentación proporcionada con el dispositivo de almacenamiento Hitachi.
Configurar el grupo de replications de Hitachi	“Configuración de un grupo de replications de Hitachi TrueCopy” en la página 98
Configurar el dispositivo de DID	“Configuración de dispositivos de DID para replicación mediante Hitachi TrueCopy” en la página 100
Registrar el grupo replicado	“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132 o “Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)” en la página 145
Comprobar la configuración	“Comprobación de la configuración de un grupo de dispositivos globales replicado con Hitachi TrueCopy” en la página 102

▼ **Configuración de un grupo de replications de Hitachi TrueCopy**

Antes de empezar

En primer lugar, configure los grupos de dispositivos de Hitachi TrueCopy en discos compartidos del clúster primario. La información sobre la configuración se especifica en el archivo `/etc/horcm.conf` en cada uno de los nodos del clúster que tenga acceso a la matriz de Hitachi. Si desea más información sobre cómo configurar el archivo `/etc/horcm.conf`, consulte *Sun StorEdge SE 9900 V Series Command and Control Interface User and Reference Guide*.



Precaución – El nombre del grupo de dispositivos de Oracle Solaris Cluster que cree (Solaris Volume Manager, Veritas Volume Manager, ZFS o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos conectados a la matriz de almacenamiento.**
- 2 **Agregue la entrada `horcm` al archivo `/etc/services`.**
`horcm 9970/udp`
Asigne un número de puerto y un nombre de protocolo a la nueva entrada.

- 3 **Especifique la información de configuración del grupo de dispositivos de Hitachi TrueCopy en el archivo `/etc/horcm.conf`.**

Si desea obtener instrucciones, consulte la documentación incluida con el software de TrueCopy.

- 4 **Inicie el daemon CCI de TrueCopy. Para ello, ejecute el comando `horcmstart.sh` en todos los nodos.**

```
# /usr/bin/horcmstart.sh
```

- 5 **Si aún no ha creado los pares de réplicas, hágalo ahora.**

Use el comando `paircreate` para crear los pares de réplicas con el nivel de protección que desee. Para obtener instrucciones sobre cómo crear los pares de réplicas, consulte la documentación de TrueCopy.

- 6 **Con el comando `pairdisplay`, compruebe que la replicación de datos se haya configurado correctamente en todos los nodos configurados. Un grupo de dispositivos de Hitachi TrueCopy o Hitachi Universal Replicator con un nivel de protección (`fence_level`) de ASYNC (asíncrono) no puede compartir el `ctgid` con ningún otro grupo de dispositivos del sistema.**

```
# pairdisplay -g group-name
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#,P/S,Status,Fence,Seq#,P-LDEV# M
group-name pair1(L) (CL1-C , 0, 9) 54321 58..P-VOL PAIR NEVER ,12345 29 -
group-name pair1(R) (CL1-A , 0, 29)12345 29..S-VOL PAIR NEVER ,----- 58 -
```

- 7 **Compruebe que todos los nodos puedan controlar los grupos de replicación.**

- a. **Con el comando `pairdisplay`, determine los nodos que contendrán la réplica primaria y la secundaria.**

```
# pairdisplay -g group-name
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#,P/S,Status,Fence,Seq#,P-LDEV# M
group-name pair1(L) (CL1-C , 0, 9) 54321 58..P-VOL PAIR NEVER ,12345 29 -
group-name pair1(R) (CL1-A , 0, 29)12345 29..S-VOL PAIR NEVER ,----- 58 -
```

El nodo con el dispositivo local (L) en el estado P-VOL contiene la réplica primaria; el nodo con el dispositivo local (L) en el estado S-VOL contiene la réplica secundaria.

- b. **Convierta al nodo secundario en el primario; para ello, ejecute el comando `horctakeover` en el nodo que contenga la réplica secundaria.**

```
# horctakeover -g group-name
```

Espere a que finalice la copia inicial de los datos antes de continuar con el paso siguiente.

- c. **Compruebe que el nodo que ejecutó `horctakeover` tenga el dispositivo local (L) en estado P-VOL.**

```
# pairdisplay -g group-name
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#,P/S,Status,Fence,Seq#,P-LDEV# M
group-name pair1(L) (CL1-C , 0, 9) 54321 58..S-VOL PAIR NEVER ,12345 29 -
group-name pair1(R) (CL1-A , 0, 29)12345 29..P-VOL PAIR NEVER ,----- 58 -
```

- d. Ejecute el comando `horctakeover` en el nodo que contenía originalmente la réplica primaria.

```
# horctakeover -g group-name
```

- e. Compruebe, mediante la ejecución del comando `pairdisplay`, que el nodo primario haya vuelto a la configuración original.

```
# pairdisplay -g group-name
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#,P/S,Status,Fence,Seq#,P-LDEV# M
group-name pair1(L) (CL1-C , 0, 9) 54321 58..P-VOL PAIR NEVER ,12345 29 -
group-name pair1(R) (CL1-A , 0, 29)12345 29..S-VOL PAIR NEVER ,----- 58 -
```

Pasos siguientes Continúe la configuración del dispositivo replicado conforme a las instrucciones de “Configuración de dispositivos de DID para replicación mediante Hitachi TrueCopy” en la página 100.

▼ Configuración de dispositivos de DID para replicación mediante Hitachi TrueCopy

Antes de empezar Tras haber configurado un grupo de dispositivos para el dispositivo replicado, debe configurar el controlador del DID utilizado por el dispositivo replicado.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Compruebe que el daemon `horcmd` se ejecute en todos los nodos.**

El comando siguiente inicia el daemon si no todavía no está ejecutándose. Si el daemon ya se está ejecutando, el sistema muestra un mensaje.

```
# /usr/bin/horcmstart.sh
```

- 3 **Determine el nodo que contiene la réplica secundaria; para ello, ejecute el comando `pairdisplay`.**

```
# pairdisplay -g group-name
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#,P/S,Status,Fence,Seq#,P-LDEV# M
group-name pair1(L) (CL1-C , 0, 9) 54321 58..P-VOL PAIR NEVER ,12345 29 -
group-name pair1(R) (CL1-A , 0, 29)12345 29..S-VOL PAIR NEVER ,----- 58 -
```

El nodo con el dispositivo local (L) en estado S-VOL contiene la réplica secundaria.

4 En el nodo con la réplica secundaria (determinado por el paso anterior), configure los dispositivos de DID para usarlos con la replicación basada en almacenamiento.

Este comando combina las dos instancias de DID de los pares de réplicas de dispositivos en una sola instancia de DID lógica. La instancia única permite que el software de administración de volúmenes use el dispositivo desde ambas partes.



Precaución – Si hay varios nodos conectados a la réplica secundaria, ejecute este comando sólo en uno de ellos.

```
# cldevice replicate -D primary-replica-nodename -S secondary replica-nodename
```

nombre_nodo_réplica_primaria

Especifica el nombre del nodo remoto que contiene la réplica primaria.

-S

Especifica un nodo de origen distinto del nodo actual.

nombre_nodo_réplica_secundaria

Especifica el nombre del nodo remoto que contiene la réplica secundaria.

Nota – De forma predeterminada, el nodo actual es el nodo de origen. Use la opción -S para indicar otro nodo de origen.

5 Compruebe que se hayan combinado las instancias de DID.

```
# cldevice list -v logical_DID_device
```

6 Compruebe que esté establecida la replicación de TrueCopy.

```
# cldevice show logical_DID_device
```

En principio, la salida del comando indica que el tipo de replicación es TrueCopy.

7 Si la reasignación de DID no combina correctamente todos los dispositivos replicados, combine los dispositivos replicados uno por uno manualmente.



Precaución – Extreme las precauciones al combinar instancias de DID de forma manual. Una reasignación incorrecta de los dispositivos podría causar daños en los datos.

a. Ejecute el comando `cldevice combine` en todos los nodos que contengan la réplica secundaria.

```
# cldevice combine -d destination-instance source-instance
```

-d La instancia de DID remota, que corresponde a la réplica primaria.

instancia_destino

instancia_origen La instancia de DID local, que corresponde a la réplica secundaria.

b. Compruebe que la reasignación de DID haya sido correcta.

```
# cldevice list desination-instance source-instance
```

Una de las instancias de DID no debería figurar en la lista.

8 Compruebe que en todos los nodos se pueda tener acceso a los dispositivos de DID de todas las instancias de DID combinadas.

```
# cldevice list -v
```

Pasos siguientes

Para finalizar la configuración del grupo de dispositivos replicado, efectúe los pasos de los procedimientos siguientes.

- “Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132 o “Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)” en la página 145

Al registrar el grupo de dispositivos, debe asignársele el mismo nombre que al grupo de replications de TrueCopy.

- “Comprobación de la configuración de un grupo de dispositivos globales replicado con Hitachi TrueCopy” en la página 102

▼ Comprobación de la configuración de un grupo de dispositivos globales replicado con Hitachi TrueCopy

Antes de empezar

Antes de poder comprobarse es preciso crear el grupo de dispositivos globales. Puede usar grupos de dispositivos de Solaris Volume Manager, Veritas Volume Manager, ZFS o disco básico. Si desea obtener más información, consulte:

- “Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132
- “Adición y registro de un grupo de dispositivos (disco básico)” en la página 134
- “Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 135
- “Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager)” en la página 137



Precaución – El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que el grupo de dispositivos primario se corresponda con el mismo nodo que contiene la réplica primaria.

```
# pairedisplay -g group-name
# cldevicegroup status -n nodename group-name
```
- 2 Compruebe que la propiedad de replicación esté configurada para el grupo de dispositivos.

```
# cldevicegroup show -n nodename group-name
```
- 3 Compruebe que la propiedad replicada esté configurada para el dispositivo.

```
# usr/cluster/bin/cldevice status [-s state] [-n node[,?]] [+| [disk-device ]]
```
- 4 Realice una conmutación de prueba para asegurarse de que los grupos de dispositivos estén bien configurados y que las réplicas puedan desplazarse de unos nodos a otros.

Si el grupo de dispositivos está fuera de línea, póngalo en línea.

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

-n *nombre_nodo* El nodo al que se conmuta el grupo de dispositivos. Este nodo se convierte en el nuevo nodo primario.
- 5 Compruebe que la conmutación haya sido correcta. Para ello, compare la salida de los comandos siguientes.

```
# pairedisplay -g group-name
# cldevicegroup status -n nodename group-name
```

Ejemplo: Configuración de un grupo de replications de TrueCopy para Oracle Solaris Cluster

En este ejemplo se realizan todos los pasos de Oracle Solaris Cluster necesarios para configurar una replicación de TrueCopy en el clúster. En el ejemplo se supone que ya se han realizado las tareas siguientes:

- Se han configurado los LUN de Hitachi
- Se ha instalado el software de TrueCopy en los nodos del clúster y el dispositivo de almacenamiento
- Se han configurado los pares de replicación de los nodos del clúster

Si desea obtener instrucciones sobre la configuración de los pares de replicación, consulte [“Configuración de un grupo de replications de Hitachi TrueCopy” en la página 98](#).

En este ejemplo aparece un clúster de tres nodos que usa TrueCopy. El clúster se distribuye en dos sitios remotos: dos nodos en el primero, y otro nodo en el segundo. Cada sitio dispone de su propio dispositivo de almacenamiento de Hitachi.

Los ejemplos siguientes muestran el archivo de configuración `/etc/horcm.conf` de TrueCopy en cada nodo.

EJEMPLO 5-1 Archivo de configuración de TrueCopy en el nodo 1

```

HORCM_DEV
#dev_group    dev_name    port#    TargetID    LU#    MU#
VG01          pair1      CL1-A    0           29
VG01          pair2      CL1-A    0           30
VG01          pair3      CL1-A    0           31
HORCM_INST
#dev_group    ip_address  service
VG01          node-3     horcm

```

EJEMPLO 5-2 Archivo de configuración de TrueCopy en el nodo 2

```

HORCM_DEV
#dev_group    dev_name    port#    TargetID    LU#    MU#
VG01          pair1      CL1-A    0           29
VG01          pair2      CL1-A    0           30
VG01          pair3      CL1-A    0           31
HORCM_INST
#dev_group    ip_address  service
VG01          node-3     horcm

```

EJEMPLO 5-3 Archivo de configuración de TrueCopy en el nodo 3

```

HORCM_DEV
#dev_group    dev_name    port#    TargetID    LU#    MU#
VG01          pair1      CL1-C    0           09
VG01          pair2      CL1-C    0           10
VG01          pair3      CL1-C    0           11
HORCM_INST
#dev_group    ip_address  service
VG01          node-1     horcm
VG01          node-2     horcm

```

En los ejemplos anteriores, se replican tres LUN entre los dos sitios. Todos los LUN están en un grupo de replicaciones denominado VG01. El comando `pairdisplay` comprueba esta información y muestra que el nodo 3 tiene la réplica primaria.

EJEMPLO 5-4 Salida del comando `pairdisplay` en el nodo 1

```

# pairdisplay -g VG01
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#.P/S,Status,Fence, Seq#,P-LDEV# M
VG01 pair1(L)      (CL1-A , 0, 29)61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R)      (CL1-C , 0, 9)20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair2(L)      (CL1-A , 0, 30)61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R)      (CL1-C , 0, 10)20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L)      (CL1-A , 0, 31)61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R)      (CL1-C , 0, 11)20064 60..P-VOL PAIR NEVER ,61114 31 -

```

EJEMPLO 5-5 Salida del comando `pairdisplay` en el nodo 2

```

# pairdisplay -g VG01
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#.P/S,Status,Fence, Seq#,P-LDEV# M
VG01 pair1(L)      (CL1-A , 0, 29)61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R)      (CL1-C , 0, 9)20064 58..P-VOL PAIR NEVER ,61114 29 -

```


EJEMPLO 5-5 Salida del comando `pairdisplay` en el nodo 2 (Continuación)

```

VG01 pair2(L) (CL1-A , 0, 30)61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R) (CL1-C , 0, 10)20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L) (CL1-A , 0, 31)61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R) (CL1-C , 0, 11)20064 60..P-VOL PAIR NEVER ,61114 31 -

```

EJEMPLO 5-6 Salida del comando `pairdisplay` en el nodo 3

```

# pairdisplay -g VG01
Group PairVol(L/R) (Port#,TID,LU),Seq#,LDEV#.P/S,Status,Fence, Seq#,P-LDEV# M
VG01 pair1(L) (CL1-C , 0, 9)20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair1(R) (CL1-A , 0, 29)61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair2(L) (CL1-C , 0, 10)20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair2(R) (CL1-A , 0, 30)61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair3(L) (CL1-C , 0, 11)20064 60..P-VOL PAIR NEVER ,61114 31 -
VG01 pair3(R) (CL1-A , 0, 31)61114 31..S-VOL PAIR NEVER ,----- 60 -

```

Para ver los discos que se están usando, emplee la opción `-fd` del comando `pairdisplay`, como se muestra en los ejemplos siguientes.

EJEMPLO 5-7 Salida del comando `pairdisplay` en el nodo 1 con los discos usados

```

# pairdisplay -fd -g VG01
Group PairVol(L/R) Device_File ,Seq#,LDEV#.P/S,Status,Fence,Seq#,P-LDEV# M
VG01 pair1(L) c6t500060E80000000000000000E00000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair2(L) c6t500060E80000000000000000E00000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R) c5t50060E80000000000000004E600000003Bd0s2 0064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L) c6t500060E80000000000000000E00000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -

```

EJEMPLO 5-8 Salida del comando `pairdisplay` en el nodo 2 con los discos usados

```

# pairdisplay -fd -g VG01
Group PairVol(L/R) Device_File ,Seq#,LDEV#.P/S,Status,Fence,Seq#,P-LDEV# M
VG01 pair1(L) c5t500060E80000000000000000E00000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair2(L) c5t500060E80000000000000000E00000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R) c5t50060E80000000000000004E600000003Bd0s2 20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L) c5t500060E80000000000000000E00000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -

```

EJEMPLO 5-9 Salida del comando `pairdisplay` en el nodo 3 con los discos usados

```

# pairdisplay -fd -g VG01
Group PairVol(L/R) Device_File ,Seq#,LDEV#.P/S,Status,Fence ,Seq#,P-LDEV# M
VG01 pair1(L) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair1(R) c6t500060E80000000000000000E00000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair2(L) c5t50060E80000000000000004E600000003Bd0s2 20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair2(R) c6t500060E80000000000000000E00000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair3(L) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -
VG01 pair3(R) c6t500060E80000000000000000E00000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -

```

Los ejemplos muestran que están en uso los discos siguientes:

- En el nodo 1:
 - c6t500060E8000000000000EEBA0000001Dd0s2
 - c6t500060E8000000000000EEBA0000001Ed0s2
 - c6t500060E8000000000000EEBA0000001Fd0s
- En el nodo 2:
 - c5t500060E8000000000000EEBA0000001Dd0s2
 - c5t500060E8000000000000EEBA0000001Ed0s2
 - c5t500060E8000000000000EEBA0000001Fd0s2
- En el nodo 3:
 - c5t50060E800000000000004E600000003Ad0s2
 - c5t50060E800000000000004E600000003Bd0s2
 - c5t50060E800000000000004E600000003Cd0s2

Para ver los dispositivos de DID que corresponden a estos discos, use el comando `cldevice list` como se muestra en los ejemplos siguientes.

EJEMPLO 5-10 Visualización de los DID correspondientes a los discos usados

cldevice list -v

```
DID Device  Full Device Path
-----
1          node-1:/dev/rdisk/c0t0d0  /dev/did/rdsk/d1
2          node-1:/dev/rdisk/c0t6d0  /dev/did/rdsk/d2
11         node-1:/dev/rdisk/c6t500060E8000000000000EEBA00000020d0  /dev/did/rdsk/d11
11         node-2:/dev/rdisk/c5t500060E8000000000000EEBA00000020d0  /dev/did/rdsk/d11
12         node-1:/dev/rdisk/c6t500060E8000000000000EEBA0000001Fd0  /dev/did/rdsk/d12
12         node-2:/dev/rdisk/c5t500060E8000000000000EEBA0000001Fd0  /dev/did/rdsk/d12
13         node-1:/dev/rdisk/c6t500060E8000000000000EEBA0000001Ed0  /dev/did/rdsk/d13
13         node-2:/dev/rdisk/c5t500060E8000000000000EEBA0000001Ed0  /dev/did/rdsk/d13
14         node-1:/dev/rdisk/c6t500060E8000000000000EEBA0000001Dd0  /dev/did/rdsk/d14
14         node-2:/dev/rdisk/c5t500060E8000000000000EEBA0000001Dd0  /dev/did/rdsk/d14
18         node-3:/dev/rdisk/c0t0d0  /dev/did/rdsk/d18
19         node-3:/dev/rdisk/c0t6d0  /dev/did/rdsk/d19
20         node-3:/dev/rdisk/c5t50060E800000000000004E6000000013d0  /dev/did/rdsk/d20
21         node-3:/dev/rdisk/c5t50060E800000000000004E600000003Dd0  /dev/did/rdsk/d21
22         node-3:/dev/rdisk/c5t50060E800000000000004E600000003Cd0  /dev/did/rdsk/d2223
23         node-3:/dev/rdisk/c5t50060E800000000000004E600000003Bd0  /dev/did/rdsk/d23
24         node-3:/dev/rdisk/c5t50060E800000000000004E600000003Ad0  /dev/did/rdsk/d24
```

Al combinar las instancias de DID de cada par de dispositivos replicados, `cldevice list` debe combinar la instancia de DID 12 con la 22, la instancia 13 con la 23 y la instancia 14 con la 24. Como el nodo 3 tiene la réplica primaria, ejecute el comando `cldevice -T`, desde el nodo 1 o el nodo 2. Combine siempre las instancias desde un nodo que tenga la réplica secundaria. Ejecute este comando sólo desde un nodo, no en ambos nodos.

El ejemplo siguiente muestra la salida obtenida al combinar instancias de DID ejecutando el comando en el nodo 1.

EJEMPLO 5-11 Combinación de instancias de DID

```
# cldevice replicate -D node-3
Remapping instances for devices replicated with node-3...
VG01 pair1 L node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Dd0
VG01 pair1 R node-3:/dev/rdisk/c5t50060E80000000000004E600000003Ad0
Combining instance 14 with 24
VG01 pair2 L node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Ed0
VG01 pair2 R node-3:/dev/rdisk/c5t50060E80000000000004E600000003Bd0
Combining instance 13 with 23
VG01 pair3 L node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Fd0
VG01 pair3 R node-3:/dev/rdisk/c5t50060E80000000000004E600000003Cd0
Combining instance 12 with 22
```

Al comprobar la salida de `cldevice list`, los LUN de ambos sitios tienen la misma instancia de DID. Al tener la misma instancia de DID, cada par de réplicas parece un único dispositivo de DID, como se muestra en el ejemplo siguiente.

EJEMPLO 5-12 Visualización de DID combinados

```
# cldevice list -v
DID Device Full Device Path
-----
1 node-1:/dev/rdisk/c0t0d0 /dev/did/rdisk/d1
2 node-1:/dev/rdisk/c0t6d0 /dev/did/rdisk/d2
11 node-1:/dev/rdisk/c6t500060E8000000000000E6BA00000020d0 /dev/did/rdisk/d11
11 node-2:/dev/rdisk/c5t50060E8000000000000E6BA00000020d0 /dev/did/rdisk/d11
18 node-3:/dev/rdisk/c0t0d0 /dev/did/rdisk/d18
19 node-3:/dev/rdisk/c0t6d0 /dev/did/rdisk/d19
20 node-3:/dev/rdisk/c5t50060E80000000000004E6000000013d0 /dev/did/rdisk/d20
21 node-3:/dev/rdisk/c5t50060E80000000000004E600000003Dd0 /dev/did/rdisk/d21
22 node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Fd0 /dev/did/rdisk/d1222
22 node-2:/dev/rdisk/c5t50060E8000000000000E6BA0000001Fd0 /dev/did/rdisk/d12
22 node-3:/dev/rdisk/c5t50060E80000000000004E600000003Cd0 /dev/did/rdisk/d22
23 node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Ed0 /dev/did/rdisk/d13
23 node-2:/dev/rdisk/c5t50060E8000000000000E6BA0000001Ed0 /dev/did/rdisk/d13
23 node-3:/dev/rdisk/c5t50060E80000000000004E600000003Bd0 /dev/did/rdisk/d23
24 node-1:/dev/rdisk/c6t500060E8000000000000E6BA0000001Dd0 /dev/did/rdisk/d24
24 node-2:/dev/rdisk/c5t50060E8000000000000E6BA0000001Dd0 /dev/did/rdisk/d24
24 node-3:/dev/rdisk/c5t50060E80000000000004E600000003Ad0 /dev/did/rdisk/d24
```

El paso siguiente es crear el grupo de dispositivos de administrador de volúmenes. Ejecute este comando desde el nodo que tenga la réplica primaria, en este ejemplo el nodo 3. Asigne al grupo de dispositivos el mismo nombre que al grupo de réplica, como se muestra en el ejemplo siguiente.

EJEMPLO 5-13 Creación del grupo de dispositivos de Solaris Volume Manager

```
# metaset -s VG01 -ah phys-deneb-3
# metaset -s VG01 -ah phys-deneb-1
# metaset -s VG01 -ah phys-deneb-2
# metaset -s VG01 -a /dev/did/rdisk/d22
# metaset -s VG01 -a /dev/did/rdisk/d23
# metaset -s VG01 -a /dev/did/rdisk/d24
# metaset
Set name = VG01, Set number = 1
```

Host	Owner
phys-deneb-3	Yes
phys-deneb-1	
phys-deneb-2	

Drive	Dbase
d22	Yes
d23	Yes
d24	Yes

En este punto, el grupo de dispositivos está disponible para su uso, se pueden crear metadispositivos y el grupo de dispositivos se puede desplazar a cualquiera de los tres nodos. Sin embargo, para incrementar la eficiencia de las conmutaciones y las migraciones tras error, ejecute `cldevicegroup set` para marcar el grupo de dispositivos como replicado en la configuración del clúster.

EJEMPLO 5-14 Transformación de las conmutaciones y las migraciones tras error en procesos eficientes

```
# cldevicegroup sync VG01
# cldevicegroup show VG01
=== Device Groups===
```

Device Group Name	VG01
Type:	SVM
failback:	no
Node List:	phys-deneb-3, phys-deneb-1, phys-deneb-2
preferenced:	yes
numsecondaries:	1
device names:	VG01
Replication type:	truecopy

La configuración del grupo de replicaciones se completa con este paso. Para comprobar que la configuración haya sido correcta, siga los pasos descritos en [“Comprobación de la configuración de un grupo de dispositivos globales replicado con Hitachi TrueCopy”](#) en la página 102.

Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility

En la tabla siguiente se enumeran las tareas necesarias para configurar y administrar un dispositivo replicado basado en almacenamiento de EMC Symmetrix Remote Data Facility (SRDF).

TABLA 5-3 Mapa de tareas: administrar un dispositivo replicado basado en almacenamiento de EMC SRDF

Tarea	Instrucciones
Instalar el software de SRDF en el dispositivo de almacenamiento y los nodos	La documentación incluida con el dispositivo de almacenamiento EMC.
Configurar el grupo de replications de EMC	“Configuración de un grupo de replications de EMC SRDF” en la página 109
Configurar el dispositivo de DID	“Configuración de dispositivos de DID para la replicación con EMC SRDF” en la página 111
Registrar el grupo replicado	“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132 o “Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)” en la página 145
Comprobar la configuración	“Comprobación de la configuración de un grupo de dispositivos globales replicado con EMC SRDF” en la página 113
Recuperar los datos de forma manual después de que la sala primaria de un clúster de campus falle por completo	“Recuperación de datos de EMC SRDF tras un error completo de la sala primaria” en la página 118

▼ Configuración de un grupo de replications de EMC SRDF

Antes de empezar

Antes de configurar un grupo de replications de EMC SRDF, el software EMC Solutions Enabler debe estar instalado en todos los nodos del clúster . En primer lugar, configure los grupos de dispositivos EMC SRDF en los discos compartidos del clúster. Si desea más información sobre cómo configurar los grupos de dispositivos EMC SRDF, consulte la documentación de EMC SRDF.

Con EMC SRDF, emplee dispositivos dinámicos, no estáticos. Los dispositivos estáticos necesitan varios minutos para cambiar el nodo primario de replicación y pueden afectar al tiempo de migración tras error.



Precaución – El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos conectados a la matriz de almacenamiento.**
- 2 **En cada nodo configurado con los datos replicados, descubra la configuración de dispositivo de Symmetrix.**

Este proceso puede tardar unos minutos.

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

- 3 **Si aún no ha creado los pares de réplicas, hágalo ahora.**

Para crear los pares de réplicas, use el comando `symrdf`. Si desea obtener instrucciones sobre cómo crear pares de réplicas, consulte la documentación de SRDF.

- 4 **Compruebe que la replicación de datos esté correctamente configurada en cada nodo configurado con dispositivos replicados.**

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

- 5 **Realice un intercambio del grupo de dispositivos.**

- a. **Compruebe que las réplicas primaria y secundaria estén sincronizadas.**

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

- b. **Determine los nodos que contienen la réplica primaria y la secundaria. Para ello, ejecute el comando `symdg show`.**

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

El nodo con el estado de dispositivo de RDF1 contiene la réplica primaria; el nodo con el estado de dispositivo de RDF2 contiene la réplica secundaria.

- c. **Habilite la réplica secundaria.**

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

- d. **Intercambie los dispositivos de RDF1 y RDF2.**

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

- e. **Habilite el par de réplicas.**

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

- f. **Compruebe que el nodo primario y las réplicas secundarias estén sincronizados.**

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

- 6 Repita todo el paso 5 en el nodo que contenía originalmente la réplica primaria.

Pasos siguientes Tras haber configurado un grupo de dispositivos para el dispositivo replicado EMC SRDF, debe configurar el controlador de DID utilizado por el dispositivo replicado.

▼ **Configuración de dispositivos de DID para la replicación con EMC SRDF**
Este procedimiento configura el controlador de DID que emplea el dispositivo replicado.

Antes de empezar `phys -s chost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.
- 2 Determine los dispositivos de DID que se corresponden con los dispositivos de RDF1 y de RDF2 configurados.

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

Nota – Si el sistema no muestra todos los parches del dispositivo de Oracle Solaris, establezca en 1 la variable de entorno `SYMCLI_FULL_PDEVNAME` y vuelva a escribir el comando `symdg -show`.

- 3 Determine los dispositivos de DID que se corresponden con los dispositivos de Oracle Solaris.
- ```
cldevice list -v
```
- 4 Para cada par de dispositivos de DID emparejados, combine las instancias de forma que obtenga un único dispositivo de DID replicado. Ejecute el comando siguiente desde la parte secundaria o RDF2.

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

**Nota** – Los dispositivos de replicación de datos de SRDF no admiten la opción `-T`.

|                                                |                                                                                                        |
|------------------------------------------------|--------------------------------------------------------------------------------------------------------|
| <code>-t tipo_replicación</code>               | Especifica el tipo de replicación. Para EMC SRDF, escriba <b>SRDF</b> .                                |
| <code>-g grupo_dispositivos_replicación</code> | Especifica el nombre del grupo de dispositivos como se muestra en el comando <code>symdg show</code> . |

|                             |                                                                               |
|-----------------------------|-------------------------------------------------------------------------------|
| <i>-d instancia_destino</i> | Especifica la instancia de DID que se corresponde con el dispositivo de RDF1. |
| <i>instancia_origen</i>     | Especifica la instancia de DID que se corresponde con el dispositivo de RDF2. |

**Nota** – Si comete una equivocación al combinar un dispositivo de DID, use la opción *-b* con el comando *sddidadm* para deshacer la combinación de dos dispositivos de DID.

|                             |                                                                                                                 |
|-----------------------------|-----------------------------------------------------------------------------------------------------------------|
| # <i>sddidadm -b device</i> |                                                                                                                 |
| <i>-b dispositivo</i>       | La instancia de DID que se correspondía con el dispositivo de destino cuando las instancias estaban combinadas. |

**5 Si se modifica el nombre de un grupo de dispositivos de replicación, en Hitachi TrueCopy y SRDF habrá que efectuar más pasos. Tras haber completado los pasos del 1 al 4, realice el paso adicional correspondiente.**

| Elemento | Descripción                                                                                                                                                                                                                                                                                                                                                                  |
|----------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| TrueCopy | Si se cambia el nombre del grupo de dispositivos de replicación (y del grupo de dispositivos globales correspondiente), debe volver a ejecutar el comando <i>cldevice replicate</i> para actualizar la información sobre el dispositivo replicado.                                                                                                                           |
| SRDF     | Si se cambia el nombre del grupo de dispositivos de replicación (y el del grupo de dispositivos globales correspondiente); debe actualizar la información del dispositivo replicado con el comando <i>sddidadm -b</i> para eliminar la información ya presente. El último paso consiste en emplear el comando <i>cldevice combine</i> para crear un dispositivo actualizado. |

**6 Compruebe que se hayan combinado las instancias de DID.**

# *cldevice list -v device*

**7 Compruebe que esté configurada la replicación de SRDF.**

# *cldevice show device*

**8 Compruebe que en todos los nodos se pueda tener acceso a los dispositivos de DID de todas las instancias de DID combinadas.**

# *cldevice list -v*

**Pasos siguientes** Tras configurar el controlador de identificador de dispositivos (DID) empleado por el dispositivo replicado, compruebe la configuración del grupo de dispositivos globales replicado con EMC SRDF.



## ▼ Comprobación de la configuración de un grupo de dispositivos globales replicado con EMC SRDF

### Antes de empezar

Antes de poder comprobarse es preciso crear el grupo de dispositivos globales. Puede usar grupos de dispositivos de Solaris Volume Manager, Veritas Volume Manager, ZFS o disco básico. Si desea obtener más información, consulte:

- “Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132
- “Adición y registro de un grupo de dispositivos (disco básico)” en la página 134
- “Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 135
- “Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager)” en la página 137



**Precaución** – El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Compruebe que el grupo de dispositivos primario se corresponda con el mismo nodo que contiene la réplica primaria.**

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

- 2 **Realice una conmutación de prueba para asegurarse de que los grupos de dispositivos estén bien configurados y que las réplicas puedan desplazarse de unos nodos a otros.**

Si el grupo de dispositivos está fuera de línea, póngalo en línea.

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

`-n nombre_nodo` El nodo al que se conmuta el grupo de dispositivos. Este nodo se convierte en el nuevo nodo primario.

- 3 **Compruebe que la conmutación haya sido correcta. Para ello, compare la salida de los comandos siguientes.**

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

## Ejemplo: configurar un grupo de replications de SRDF para Oracle Solaris Cluster

Este ejemplo muestra cómo efectuar los pasos de Oracle Solaris Cluster que necesarios para configurar la replicación con SRDF en el clúster. En el ejemplo se supone que ya se han realizado las tareas siguientes:

- Se ha completado el emparejamiento de los LUN para la replicación entre matrices.
- Se ha instalado el software de SRDF en el dispositivo de almacenamiento y los nodos del clúster.

Este ejemplo aparece un clúster de cuatro nodos: dos de ellos están conectados a un dispositivo de Symmetrix y los otros dos están conectados a un segundo dispositivo de Symmetrix. El grupo de dispositivos de SRDF se denomina dg1.

### EJEMPLO 5-15 Creación de pares de réplicas

Ejecute los comandos siguientes en todos los nodos.

```
symcfg discover
! This operation might take up to a few minutes.
symdev list pd
```

Symmetrix ID: 000187990182

| Device Name |                        | Directors |                  | Device    |         |          |      |
|-------------|------------------------|-----------|------------------|-----------|---------|----------|------|
| Sym         | Physical               | SA        | :P DA :IT Config | Attribute | Sts     | Cap (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 |
| ...         |                        |           |                  |           |         |          |      |

En todos los nodos de la parte de RDF1, escriba:

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

En todos los nodos de la parte de RDF2, escriba:

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

### EJEMPLO 5-16 Comprobación de la configuración de la replicación de datos

Desde un nodo del clúster, escriba:

## EJEMPLO 5-16 Comprobación de la configuración de la replicación de datos (Continuación)

```
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/rdisk/c5t600604800018790002353594D303637d0s2 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
```

EJEMPLO 5-16 Comprobación de la configuración de la replicación de datos (Continuación)

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    |        |        |          |              |

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

#

EJEMPLO 5-17 Visualización de los DID correspondientes a los discos usados

El mismo procedimiento es válido para las partes de RDF1 y de RDF2.

Observe el campo PdevName de la salida del comando dymdg show dg.

En la parte de RDF1, escriba:

```
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
...
```

**EJEMPLO 5-17** Visualización de los DID correspondientes a los discos usados (Continuación)

Para obtener el DID correspondiente, escriba:

```
sccidadm -L | grep c5t6006048000018790002353594D303637d0
217 pmoney1:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdsk/d217
217 pmoney2:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdsk/d217
#
```

Para que se muestre el DID correspondiente, escriba:

```
cldevice show d217

=== DID Device Instances ===

DID Device Name: /dev/did/rdsk/d217
Full Device Path: pmoney2:/dev/rdsk/c5t6006048000018790002353594D303637d0
Full Device Path: pmoney1:/dev/rdsk/c5t6006048000018790002353594D303637d0
Replication: none
default_fencing: global

#
```

En la parte de RDF2, escriba:

Observe el campo PdevName de la salida del comando `dymdg show dg`.

```
symdg show dg1

Group Name: dg1

Group Type : RDF2 (RDFA)

...
Standard (STD) Devices (1):
{

LdevName PdevName Sym Dev Att. Sts Cap

DEV001 /dev/rdsk/c5t6006048000018799018253594D303637d0s2 0067 WD 4315
}

Device Group RDF Information
...
```

Para obtener el DID correspondiente, escriba:

```
sccidadm -L | grep c5t6006048000018799018253594D303637d0
108 pmoney4:/dev/rdsk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108
108 pmoney3:/dev/rdsk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108
#
```

Para que se muestre el DID correspondiente, escriba:

**EJEMPLO 5-17** Visualización de los DID correspondientes a los discos usados *(Continuación)*

```
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

#
```

**EJEMPLO 5-18** Combinación de instancias de DID

Desde la parte de RDF2, escriba:

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

**EJEMPLO 5-19** Visualización de DID combinados

Desde cualquier nodo del clúster, escriba:

```
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

#
```

## ▼ Recuperación de datos de EMC SRDF tras un error completo de la sala primaria

Este procedimiento realiza una recuperación de los datos cuando la sala primaria de un clúster de campus experimenta un error completo; la sala primaria migra tras error a una sala secundaria; posteriormente, la sala primaria vuelve a estar en línea. La sala primaria del clúster de campus es el nodo primario y el sitio de almacenamiento. El error completo de una sala incluye el del host y el del almacenamiento de dicha sala. Si hay un error en la sala primaria, automáticamente Oracle Solaris Cluster migra tras error a la sala secundaria, posibilita la lectura y escritura en el almacenamiento de la sala secundaria, y habilita la migración tras error de los grupos de dispositivos y de recursos correspondientes.

Cuando la sala primaria vuelve a estar en línea, los datos del grupo de dispositivos de SRDF escritos en la sala secundaria se pueden recuperar manualmente y volver a sincronizarse. Este procedimiento recupera el grupo de dispositivos de SRDF sincronizando la sala secundaria original (este procedimiento emplea *phys-campus-2* para la sala secundaria) con la sala primaria original (*phys-campus-1*). El procedimiento también cambia el tipo de grupo de dispositivos de SRDF a RDF1 en *phys-campus-2* y a RDF2 en *phys-campus-1*.

#### Antes de empezar

Antes de poder realizar manualmente una migración tras error, debe configurar el grupo de replications de EMC y los dispositivos de DID, así como registrar el grupo de replications de EMC. Si desea obtener información sobre cómo crear un grupo de dispositivos de Solaris Volume Manager, consulte [“Adición y registro de un grupo de dispositivos \(Solaris Volume Manager\)” en la página 132](#). Si desea obtener información sobre cómo crear un grupo de dispositivos de Veritas Volume Manager, consulte [“Creación de un nuevo grupo de discos al encapsular discos \(Veritas Volume Manager\)” en la página 141](#).

---

**Nota** – Estas instrucciones presentan un método válido para recuperar datos de SRDF manualmente tras error completo de la sala primaria y volver a estar en línea. Consulte la documentación de EMC para ver otros métodos.

---

Inicie sesión en la sala primaria del clúster de campus para seguir estos pasos. En el procedimiento que se describe a continuación, *dg1* es el nombre del grupo de dispositivos de SRDF. En el momento del error, la sala primaria de este procedimiento es *phys-campus-1* y la sala secundaria es *phys-campus-2*.

- 1 Inicie sesión en la sala primaria del clúster de campus y conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 En la sala primaria, use el comando `symrdf` para consultar el estado de replicación de los dispositivos de RDF y visualizar información acerca de ellos.**

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

---

**Consejo** – Los grupos de dispositivos con estado `split` (dividido) no están sincronizados.

---

- 3 Si el estado del par de RDF es `split` (dividido) y el tipo de grupo de dispositivos es RDF1, fuerce una migración tras error del grupo de dispositivos de SRDF.**

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

- 4 Visualice el estado de los dispositivos de RDF.**

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

- 5 Después de la migración tras error, puede intercambiar los datos de los dispositivos de RDF que efectuaron la migración.

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

- 6 Compruebe el estado y demás información relativa a los dispositivos de RDF.

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

- 7 Establezca el grupo de dispositivos de SRDF en la sala primaria.

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

- 8 Confirme que el grupo de dispositivos esté sincronizado y que sea del tipo RDF2.

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

### Ejemplo 5–20 Recuperación manual de datos de EMC SRDF después de una migración tras error habida en un sitio primario

Este ejemplo muestra los pasos necesarios en Oracle Solaris Cluster para recuperar manualmente los datos de EMC SRDF después de que la sala primaria de un clúster de campus migre tras error, una sala secundaria controle y registre los datos y, posteriormente, la sala primaria vuelva a estar en línea. En el ejemplo, el grupo de dispositivos de SRDF se denomina *dg1* y el dispositivo lógico estándar es DEV001. En el momento del error, la sala primaria es *phys-campus-1* y la sala secundaria es *phys-campus-2*. Efectúe los pasos desde la sala primaria del clúster de campus, *phys-campus-1*.

```
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 00187990182 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
```

## Información general sobre la administración de sistemas de archivos de clústers

No es necesario utilizar ningún comando especial de Oracle Solaris Cluster para la administración del sistema de archivos de clúster. Un sistema de archivos de clúster se administra igual que cualquier otro sistema de archivos de Oracle Solaris, es decir, mediante comandos estándar de sistema de archivos de Oracle Solaris como `mount` y `newfs`. Puede montar sistemas de archivos de clúster si especifica la opción `-g` en el comando `mount`. Los sistemas de archivos de clúster también se pueden montar automáticamente durante el arranque. Los sistemas de archivos de clúster sólo son visibles desde el nodo de votación de un clúster global. Si necesita poder tener acceso a los datos del sistema de archivos de clúster desde un nodo sin votación, asigne los datos a dicho nodo con [zoneadm\(1M\)](#) o `HASStoragePlus`.

---

**Nota** – Cuando el sistema de archivos de clúster lee archivos, no actualiza la hora de acceso a tales archivos.

---

## Restricciones del sistema de archivos de clúster

La administración del sistema de archivos de clúster presenta las restricciones siguientes:

- El comando [unlink\(1M\)](#) no se admite en directorios que no estén vacíos.
- No se admite el comando `lockfs -d`. Como solución alternativa, utilice el comando `lockfs -n`.
- No puede volver a montar un sistema de archivos de clúster con la opción de montaje `directio` agregada en el momento en que el montaje se realiza de nuevo.
- Se admite ZFS para sistemas de archivos root, con una excepción importante. Si emplea una partición dedicada del disco de arranque para el sistema de archivos de los dispositivos globales, como sistema de archivos se debe usar sólo UFS. El espacio de nombre de dispositivos globales necesita que el sistema de archivos proxy (PxFS) se ejecute en un sistema de archivos UFS. Sin embargo, un sistema de archivos UFS para el espacio de nombre de dispositivos globales puede coexistir con un sistema de archivos ZFS para el sistema de archivos root (`/`) y otros sistemas de archivos root, por ejemplo `/var` o `/home`.

Asimismo, si se utiliza un dispositivo de lofi para alojar el espacio de nombre de dispositivos globales, no hay límites en el uso de ZFS para sistemas de archivos root.

## Directrices para admitir VxFS

Las siguientes funciones de VxFS no son compatibles con un sistema de archivos de clúster de Oracle Solaris Cluster. Sin embargo, sí son compatibles en un sistema de archivos local.

- E/S rápida
- Instantáneas
- Puntos de control de almacenamiento
- Opciones de montaje específicas de VxFS:
  - convosync (Convertir O\_SYNC)
  - mincache
  - qlog, delaylog, tmplog
- Sistema de archivos de clúster de Veritas (requiere VxVM, función de clúster y Veritas Cluster Server). La función de clúster de VxVM no es compatible con los sistemas basados en x86.

Se pueden usar notificaciones de la memoria caché, pero sólo surten efecto en el nodo concreto.

Todas las demás funciones y opciones de VxFS compatibles con un sistema de archivos de clúster, también son compatibles con el software Oracle Solaris Cluster. Consulte la documentación de VxFS para obtener más información sobre las opciones de VxFS que se admiten en una configuración de clúster.

Las directrices siguientes sobre el uso de VxFS para crear sistemas de archivos de clúster de alta disponibilidad se aplican a una configuración de Oracle Solaris 3.3.

- Siga los procedimientos que figuran a continuación, descritos en la documentación de VxFS, para crear un sistema de archivos de VxFS.
- Monte y desmonte un sistema de archivos de VxFS desde el nodo primario. El nodo primario controla el disco en el que reside el sistema de archivos de VxFS. Una operación de montaje o desmontaje de sistema de archivos de VxFS hecha desde un nodo secundario puede experimentar un error.
- Ejecute todos los comandos de administración de VxFS desde el nodo primario del sistema de archivos de clúster de VxFS.

Las directrices siguientes para administrar sistemas de archivos de clúster de UFS; no son exclusivas del software Oracle Solaris 3.3. Ahora bien, las directrices difieren de la manera de administrar sistemas de archivos de clúster UFS.

- Los archivos de un sistema de archivos de clúster de VxFS pueden administrarse desde cualquier nodo del clúster. La excepción es `ioctls`, que debe ejecutarse sólo desde el nodo primario. Si ignora si `un` comando de administración afecta a `ioctls`, ejecute el comando desde el nodo primario.
- Si se da un error en un sistema de archivos de clúster de VxFS y migra a un nodo secundario, todas las operaciones de llamada del sistema estándar que estaban en curso durante la migración tras error vuelven a ejecutarse de forma transparente en el nuevo nodo primario. Sin embargo, las operaciones relacionadas con `ioctl` que estuviesen en curso durante el proceso de migración tras error no se podrán realizar. Después de una migración tras error del sistema de archivos de clúster de VxFS, compruebe el estado del sistema de archivos de clúster. Los comandos administrativos ejecutados en el anterior primario antiguo antes de la migración tras error podrían necesitar correcciones. Si desea obtener más información, consulte la documentación de VxFS.

## Administración de grupos de dispositivos

A medida que cambien las necesidades del clúster, tal vez necesite agregar, quitar o modificar los grupos de dispositivos de dicho clúster. Oracle Solaris Cluster ofrece una interfaz interactiva, denominada `clsetup`, apta para efectuar estas modificaciones. La utilidad `clsetup` genera comandos `cluster`. Los comandos generados se muestran en los ejemplos que figuran al final de algunos procedimientos. En la tabla siguiente se enumeran tareas para administrar grupos de dispositivos, además de vínculos a los correspondientes procedimientos de esta sección.



---

**Precaución** – No ejecute `metaset -s nombre_conjunto -f -t` en un nodo del clúster que se arranque fuera del clúster si hay otros nodos que son miembros activos del clúster y al menos uno de ellos es propietario del conjunto de discos.

---

---

**Nota** – El software Oracle Solaris Cluster crea automáticamente un grupo de dispositivos de discos básicos para cada dispositivo de disco y de cinta del clúster. Sin embargo, los grupos de dispositivos del clúster permanecen fuera de línea hasta que se acceda a dichos grupos como dispositivos globales.

---

TABLA 5-4 Mapa de tareas: administrar grupos de dispositivos

| Tarea                                                                                                                                          | Instrucciones                                                                                                                                                                                                                                                                                                                            |
|------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Actualizar el espacio de nombre los dispositivos globales sin un rearranque de reconfiguración con el comando <code>cldevice populate</code>   | <a href="#">“Actualización del espacio de nombre de dispositivos globales” en la página 126</a>                                                                                                                                                                                                                                          |
| Cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales                       | <a href="#">“Cómo cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales” en la página 127</a>                                                                                                                                                                         |
| Desplazar un espacio de nombre de dispositivos globales                                                                                        | <a href="#">“Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de <code>lofi</code>” en la página 129</a><br><a href="#">“Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada” en la página 130</a> |
| Agregar conjuntos de discos de Solaris Volume Manager y registrarlos como grupos de dispositivos mediante el comando <code>metaset</code>      | <a href="#">“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 132</a>                                                                                                                                                                                                                               |
| Agregar y registrar un grupo de dispositivos de discos básicos con el comando <code>cldevicegroup</code>                                       | <a href="#">“Adición y registro de un grupo de dispositivos (disco básico)” en la página 134</a>                                                                                                                                                                                                                                         |
| Agregar un grupo de dispositivos con nombre para ZFS con el comando <code>cldevicegroup</code>                                                 | <a href="#">“Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 135</a>                                                                                                                                                                                                                                        |
| Agregar y registrar un nuevo grupo de discos como grupo de dispositivos con el método que se prefiera                                          | <a href="#">“Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager)” en la página 137</a>                                                                                                                                                                                                                   |
| Eliminar grupos de dispositivos de Solaris Volume Manager de la configuración con los comandos <code>metaset</code> y <code>metaclear</code>   | <a href="#">“Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 138</a>                                                                                                                                                                                                             |
| Eliminar un nodo de todos los grupos de dispositivos con los comandos <code>cldevicegroup</code> , <code>metaset</code> y <code>clsetup</code> | <a href="#">“Eliminación de un nodo de todos los grupos de dispositivos” en la página 138</a>                                                                                                                                                                                                                                            |
| Eliminar un nodo de un grupo de dispositivos de Solaris Volume Manager con el comando <code>metaset</code>                                     | <a href="#">“Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)” en la página 139</a>                                                                                                                                                                                                                           |

TABLA 5-4 Mapa de tareas: administrar grupos de dispositivos (Continuación)

| Tarea                                                                                                                                           | Instrucciones                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|-------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Agregar grupos de discos de Veritas Volume Manager como grupos de dispositivos con los comandos VxVM y clsetup                                  | <p>“Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager)” en la página 137</p> <p>“Creación de un nuevo grupo de discos al encapsular discos (Veritas Volume Manager)” en la página 141</p> <p>“Adición de un volumen nuevo a un grupo de dispositivos ya creado (Veritas Volume Manager)” en la página 143</p> <p>“Conversión de un grupo de discos en un grupo de dispositivos (Veritas Volume Manager)” en la página 144</p> <p>“Asignación de un número menor nuevo a un grupo de dispositivos (Veritas Volume Manager)” en la página 144</p> <p>“Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)” en la página 145</p> <p>“Conversión de un grupo de discos local en un grupo de dispositivos (VxVM)” en la página 149</p> <p>“Conversión de un grupo de dispositivos en un grupo de discos local (VxVM)” en la página 150</p> <p>“Registro de modificaciones en la configuración de grupos de discos (Veritas Volume Manager)” en la página 148</p> |
| Eliminar los grupos de dispositivos de Veritas Volume Manager de la configuración mediante los comandos de clsetup (para generar cldevicegroup) | <p>“Eliminación de un volumen de un grupo de dispositivos (Veritas Volume Manager)” en la página 151</p> <p>“Eliminación y anulación del registro de un grupo de dispositivos (Veritas Volume Manager)” en la página 152</p>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| Agregar un nodo a un grupo de dispositivos de Veritas Volume Manager con clsetup para generar cldevicegroup                                     | “Adición de un nodo a un grupo de dispositivos (Veritas Volume Manager)” en la página 153                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| Eliminar un nodo de un grupo de dispositivos de Veritas Volume Manager con clsetup para generar cldevicegroup                                   | “Eliminación de un nodo de un grupo de dispositivos (Veritas Volume Manager)” en la página 155                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
| Eliminar un nodo de un grupo de dispositivos de discos básicos con el comando cldevicegroup                                                     | “Eliminación de un nodo de un grupo de dispositivos de discos básicos” en la página 157                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| Modificar las propiedades del grupo de dispositivos con clsetup para generar cldevicegroup                                                      | “Cambio de propiedades de los grupos de dispositivos” en la página 158                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      |
| Mostrar las propiedades y los grupos de dispositivos con el comando cldevicegroup show                                                          | “Enumeración de la configuración de grupos de dispositivos” en la página 163                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |

TABLA 5-4 Mapa de tareas: administrar grupos de dispositivos (Continuación)

| Tarea                                                                                                                                 | Instrucciones                                                                              |
|---------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------|
| Cambiar el número deseado de secundarios de un grupo de dispositivos con <code>clsetup</code> para generar <code>cldevicegroup</code> | “Establecimiento del número de secundarios para un grupo de dispositivos” en la página 160 |
| Conmutar el primario de un grupo de dispositivos con el comando <code>cldevicegroup switch</code>                                     | “Conmutación al nodo primario de un grupo de dispositivos” en la página 165                |
| Poner un grupo de dispositivos en estado de mantenimiento con el comando <code>metaset</code> o <code>vxdg</code>                     | “Colocación de un grupo de dispositivos en estado de mantenimiento” en la página 166       |

## ▼ Actualización del espacio de nombre de dispositivos globales

Al agregar un nuevo dispositivo global, actualice manualmente el espacio de nombre de dispositivos globales con el comando `cldevice populate`.

**Nota** – El comando `cldevice populate` no tiene efecto alguno si el nodo que lo está ejecutando no es miembro del clúster. Tampoco tiene ningún efecto si el sistema de archivos `/global/.devices/node@ID_nodo` no está montado.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Ejecute el comando `devfsadm(1M)` en todos los nodos del clúster.**

Este comando puede ejecutarse simultáneamente en todos los nodos del clúster.

- 3 **Reconfigure el espacio de nombre.**

```
cldevice populate
```

- 4 **Antes de intentar crear conjuntos de discos, compruebe que el comando `cldevice populate` haya finalizado su ejecución en cada uno de los nodos.**

El comando `cldevice` se llama a sí mismo de forma remota en todos los nodos, incluso cuando el comando se ejecuta desde un solo nodo. Para determinar si ha concluido el procesamiento del comando `cldevice populate`, ejecute el comando siguiente en todos los nodos del clúster.

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

**Ejemplo 5–21 Actualización del espacio de nombre de dispositivos globales**

El ejemplo siguiente muestra la salida generada al ejecutarse correctamente el comando `cldevice populate`.

```
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
```

## ▼ **Cómo cambiar el tamaño de un dispositivo `lofi` que se utiliza para el espacio de nombres de dispositivos globales**

Si utiliza un dispositivo `lofi` para el espacio de nombres de dispositivos globales en uno o más nodos del clúster global, lleve a cabo este procedimiento para cambiar el tamaño del dispositivo.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en un nodo donde quiera cambiar el tamaño del dispositivo `lofi` para el espacio de nombres de dispositivos globales.**

- 2 **Evacuar los servicios fuera del nodo y reiniciar el nodo en un modo que no sea de clúster**

Lleve a cabo esta tarea para garantizar que este nodo no preste servicio a los dispositivos globales mientras se realice este procedimiento. Para obtener instrucciones, consulte [“Rearranque de un nodo en un modo que no sea de clúster” en la página 79](#).

- 3 **Desmonte el sistema de archivos del dispositivo global y desconecte el dispositivo `lofi`.**

El sistema de archivos de dispositivos globales se monta de forma local.

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

*Ensure that the lofi device is detached*

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

*The command returns no output if the device is detached*

---

**Nota** – Si el sistema de archivos se monta mediante la opción `-m` opción, no se agregará ninguna entrada archivo `mnttab`. El comando `umount` puede mostrar una advertencia similar a la siguiente:

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

Puede ignorar esta advertencia.

---

**4 Elimine y vuelva a crear el archivo `/globaldevices` con el tamaño necesario.**

En el siguiente ejemplo se muestra la creación de un archivo `/globaldevices` cuyo tamaño es de 200 MB.

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

**5 Cree un sistema de archivos para el espacio de nombres de dispositivos globales.**

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

**6 Inicie el nodo en modo de clúster.**

El nuevo sistema de archivos se rellena con los dispositivos globales.

```
phys-schost# reboot
```

**7 Migre al nodo todos los servicios que desee ejecutar en él.**

## Migración del espacio de nombre de dispositivos globales

Puede crear un espacio de nombre en un dispositivo de archivo de bucle invertido (lofi), en lugar de crear uno para dispositivos globales en una partición dedicada. Esta función es útil si se instala software de Oracle Solaris Cluster en sistemas que ya tienen instalado el sistema operativo Solaris 10.



---

**Nota** – Se admite ZFS para sistemas de archivos root, con una excepción importante. Si emplea una partición dedicada del disco de arranque para el sistema de archivos de los dispositivos globales, como sistema de archivos se debe usar sólo UFS. El espacio de nombre de dispositivos globales necesita que el sistema de archivos proxy (PxFS) se ejecute en un sistema de archivos UFS. Sin embargo, un sistema de archivos UFS para el espacio de nombre de dispositivos globales puede coexistir con un sistema de archivos ZFS para el sistema de archivos root (/) y otros sistemas de archivos root, por ejemplo /var o /home. Asimismo, si se utiliza un dispositivo de lofi para alojar el espacio de nombre de dispositivos globales, no hay límites en el uso de ZFS para sistemas de archivos root.

---

Los procedimientos siguientes describen cómo desplazar un espacio de nombre de dispositivos globales ya existente desde una partición dedicada hasta un dispositivo de lofi y viceversa:

- “Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de lofi” en la página 129
- “Migración del espacio de nombre de dispositivos globales desde un dispositivo de lofi hasta una partición dedicada” en la página 130

## ▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de lofi

- 1 Conviértase en superusuario del nodo de votación del clúster global cuya ubicación de espacio de nombre desee modificar.
- 2 Evacue los servicios fuera del nodo y reinicie el nodo en un modo que no sea de clúster.  
Haga esto para garantizar que este nodo no preste servicio a los dispositivos globales mientras se realice este procedimiento. Para obtener instrucciones, consulte “Rearranque de un nodo en un modo que no sea de clúster” en la página 79.
- 3 Compruebe que en el nodo no exista ningún archivo denominado /.globaldevices. Si lo hay, elimínelo.
- 4 Cree el dispositivo de lofi.  

```
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 En el archivo /etc/vfstab, comente la entrada del espacio de nombre de dispositivos globales. Esta entrada presenta una ruta de montaje que comienza con  
/global/.devices/node@nodeID.

6 Desmonte la partición de dispositivos globales `/global/.devices/node@nodeID`.

7 Inhabilite y vuelva a habilitar los servicios de SMF `globaldevices` y `scmountdev`.

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

A continuación, se crea un dispositivo de `lofi` en `/globaldevices` y se monta en el sistema de archivos de dispositivos globales.

8 Repita estos pasos en otros nodos cuyo espacio de nombre de dispositivos globales quiera migrar desde una partición hasta un dispositivo de `lofi`.

9 Desde un nodo, complete los espacios de nombre de dispositivos globales.

```
/usr/cluster/bin/cldevice populate
```

Antes de llevar a cabo ninguna otra acción en el clúster, compruebe que el procesamiento del comando haya concluido en cada uno de los nodos.

```
ps -ef \ grep cldevice populate
```

El espacio de nombre de dispositivos globales reside ahora en el dispositivo de `lofi`.

10 Migre al nodo todos los servicios que desee ejecutar en él.

## ▼ Migración del espacio de nombre de dispositivos globales desde un dispositivo de `lofi` hasta una partición dedicada

1 Conviértase en superusuario del nodo de votación del clúster global cuya ubicación de espacio de nombre desee modificar.

2 Evacuar los servicios fuera del nodo y reiniciar el nodo en un modo que no sea de clúster

Lleve a cabo esta tarea para garantizar que este nodo no preste servicio a los dispositivos globales mientras se realice este procedimiento. Para obtener instrucciones, consulte [“Rearranque de un nodo en un modo que no sea de clúster” en la página 79](#).

3 En un disco local del nodo, cree una nueva partición que cumpla con los requisitos siguientes:

- Un tamaño mínimo de 512 MB
- Usa el sistema de archivos UFS

- 4 Agregue una entrada al archivo `/etc/vfstab` para que la nueva partición se monte como sistema de archivos de dispositivos globales.

- Determine el ID de nodo del nodo actual.

```
/usr/sbin/clinfo -nnode ID
```

- Cree la nueva entrada en el archivo `/etc/vfstab` con el formato siguiente:

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

Por ejemplo, si la partición que elige es `/dev/did/rdisk/d5s3`, la entrada nueva que se agrega al archivo `/etc/vfstab` debe ser: `/dev/did/dsk/d5s3 /dev/did/rdisk/d5s3 /global/.devices/node@3 ufs 2 no global`

- 5 Desmonte la partición de los dispositivos globales `/global/.devices/node@ID_nodo`.

- 6 Quite el dispositivo de `lofi` asociado con el archivo `/.globaldevices`.

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

- 7 Elimine el archivo `/.globaldevices`.

```
rm /.globaldevices
```

- 8 Inhabilite y vuelva a habilitar los servicios de SMF `globaldevices` y `scmountdev`.

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

Ahora la partición está montada como sistema de archivos del espacio de nombre de dispositivos globales.

- 9 Repita estos pasos en otros nodos cuyo espacio de nombre de dispositivos globales desee migrar desde un dispositivo de `lofi` hasta una partición.

- 10 Desde un nodo del clúster, ejecute el comando `cldevice populate` para llenar el espacio de nombre de dispositivos globales.

```
/usr/cluster/bin/cldevice populate
```

Antes de pasar a realizar ninguna otra acción en cualquiera de los nodos, compruebe que el proceso concluya en todos los nodos del clúster.

```
ps -ef | grep cldevie populate
```

El espacio de nombre de dispositivos globales ahora reside en la partición dedicada.

- 11 Migre al nodo todos los servicios que desee ejecutar en él.

## Adición y registro de grupos de dispositivos

Puede agregar y registrar grupos de dispositivos para Solaris Volume Manager, ZFS, Veritas Volume Manager o disco básico.

### ▼ Adición y registro de un grupo de dispositivos (Solaris Volume Manager)

Use el comando `metaset` para crear un conjunto de discos de Solaris Volume Manager y registrarlo como grupo de dispositivos de Oracle Solaris Cluster. Al registrar el conjunto de discos, el nombre que le haya asignado se asigna automáticamente al grupo de dispositivos.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.



---

**Precaución** – El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

---

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en uno de los nodos conectado a los discos en los que esté creando el conjunto de discos.**
- 2 **Agregue el conjunto de discos de Solaris Volume Manager y regístrelo como un grupo de dispositivos con Oracle Solaris Cluster. Para crear un grupo de discos de múltiples propietarios, use la opción `-M`.**

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

`-s conjunto_discos` Especifica el conjunto de discos que se va a crear.

`-a -h lista_nodos` Agregue la lista de nodos que pueden controlar el conjunto de discos.

`-M` Designa el grupo de discos como grupo de múltiples propietarios.

---

**Nota** – Al ejecutar el comando `metaset` para configurar un grupo de dispositivos de Solaris Volume Manager en un clúster, de forma predeterminada se obtiene un nodo secundario, sea cual sea el número de nodos incluidos en dicho grupo de dispositivos. La cantidad de nodos secundarios puede cambiarse mediante la utilidad `clsetup` tras haber creado el grupo de dispositivos. Consulte [“Establecimiento del número de secundarios para un grupo de dispositivos” en la página 160](#) si desea obtener más información sobre la migración de disco tras error.

---

**3 Si configura un grupo de dispositivos replicado, establezca la propiedad de replicación para el grupo de dispositivos.**

```
cldevicegroup sync devicegroup
```

**4 Compruebe que se haya agregado el grupo de dispositivos.**

El nombre del grupo de dispositivos coincide con el del conjunto de discos que se especifica con `metaset`.

```
cldevicegroup list
```

**5 Enumere las asignaciones de DID.**

```
cldevice show | grep Device
```

- Elija las unidades que comparten los nodos del clúster que vayan a controlar el conjunto de discos o que tengan la posibilidad de hacerlo.
- Use el nombre de dispositivo de DID completo, que tiene el formato `/dev/did/rdisk/d N`, al agregar una unidad a un conjunto de discos.

En el ejemplo siguiente, las entradas del dispositivo de DID `/dev/did/rdisk/d3` indican que `phys-schost-1` y `phys-schost-2` comparten la unidad.

```
=== 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/c1t1d0
Full Device Path: phys-schost-2:/dev/rdisk/c1t1d0
...
```

**6 Agregue las unidades al conjunto de discos.**

Utilice el nombre completo de la ruta de DID.

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

`-s nombre_conjunto`      Especifique el nombre del conjunto de discos, idéntico al del grupo de dispositivos.

`-a`                              Agrega la unidad al conjunto de discos.

---

**Nota** – No use el nombre de dispositivo de nivel inferior (cNtX dY) al agregar una unidad a un conjunto de discos. Ya que el nombre de dispositivo de nivel inferior es un nombre local y no único para todo el clúster, si se utiliza es posible que se prive al metaconjunto de la capacidad de conmutar a otra ubicación.

---

## 7 Compruebe el estado del conjunto de discos y de las unidades.

```
metaset -s setname
```

### Ejemplo 5–22 Adición a un grupo de dispositivos de Solaris Volume Manager

En el ejemplo siguiente se muestra la creación del conjunto de discos y el grupo de dispositivos con las unidades de disco /dev/did/rdisk/d1 y /dev/did/rdisk/d2; asimismo, se comprueba que se haya creado el grupo de dispositivos.

```
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
```

## ▼ Adición y registro de un grupo de dispositivos (disco básico)

El software Oracle Solaris Cluster admite el uso de grupos de dispositivos de discos básicos y otros administradores de volúmenes. Al configurar Oracle Solaris Cluster, los grupos de dispositivos se configuran automáticamente para cada dispositivo básico del clúster. Utilice este procedimiento para reconfigurar automáticamente estos grupos de dispositivos creados a fin de usarlos con el software Sun Oracle Solaris Cluster.

Puede crear un grupo de dispositivos de disco básico por los siguientes motivos:

- Desea agregar más de un DID al grupo de dispositivos.
- Necesita cambiar el nombre del grupo de dispositivos.
- Desea crear una lista de grupos de dispositivos sin utilizar la opción -v del comando cldg.



---

**Precaución** – Si crea un grupo de dispositivos en dispositivos replicados, el nombre del grupo de dispositivos que crea (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.

---

- 1 **Identifique los dispositivos que desee usar y anule la configuración de cualquier grupo de dispositivos predefinido.**

Los comandos siguientes sirven para eliminar los grupos de dispositivos predefinidos de d7 y d8.

```
paris-1# cldevicegroup disable dsk/d7 dsk/d8
paris-1# cldevicegroup offline dsk/d7 dsk/d8
paris-1# cldevicegroup delete dsk/d7 dsk/d8
```

- 2 **Cree el grupo de dispositivos de disco básico con los dispositivos deseados.**

El comando siguiente crea un grupo de dispositivos globales, rawdg, que contiene d7 y d8.

```
paris-1# cldevicegroup create -n phys-paris-1,phys-paris-2 -t rawdisk
-d d7,d8 rawdg
paris-1# /usr/cluster/lib/dcs/cldg show rawdg -d d7 rawdg
paris-1# /usr/cluster/lib/dcs/cldg show rawdg -d d8 rawdg
```

## ▼ Adición y registro de un grupo de dispositivos replicado (ZFS)

Para replicar ZFS, debe crear un grupo de dispositivos con un nombre y hacer que figuren en una lista los discos que pertenecen a zpool. Un dispositivo sólo puede pertenecer a un grupo de dispositivos a la vez; por eso, si ya tiene un grupo de dispositivos de Oracle Solaris Cluster que contiene el dispositivo, debe eliminar este grupo antes de agregar el dispositivo al nuevo grupo de dispositivos de ZFS.

El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager, Veritas Volume Manager o disco básico) debe ser idéntico al del grupo de dispositivos replicado.



**Precaución** – Hoy por hoy, ZFS no es compatible con las tecnologías de replicación de datos de otros proveedores. Para obtener información sobre las actualizaciones de la compatibilidad de ZFS, consulte las notas sobre la versión más actuales de Oracle Solaris Cluster.

- 1 **Elimine los grupos de dispositivos predeterminados que se correspondan con los dispositivos de zpool.**

Por ejemplo, si dispone de un zpool denominado mypool que contiene dos dispositivos, /dev/did/dsk/d2 y /dev/did/dsk/d13, debe eliminar los dos grupos de dispositivos predeterminados llamados d2 y d13.

```
cldevicegroup offline dsk/d2 dsk/d13
cldevicegroup remove dsk/d2 dsk/d13
```

- 2 Cree un grupo de dispositivos cuyos DID se correspondan con los del grupo de dispositivos eliminado en el Paso 1.

```
cldevicegroup create -d d2,d13 -t rawdisk mypool
```

Esta acción crea un grupo de dispositivos denominado mypool (con el mismo nombre que zpool), que administra los dispositivos básicos /dev/did/dsk/d2 y /dev/did/dsk/d13.

- 3 Cree un zpool que contenga estos dispositivos.

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

- 4 Cree un grupo de recursos para administrar la migración de los dispositivos replicados (en el grupo de dispositivos) que sólo cuenten con zonas globales en su lista de nodos.

```
clrg create -n pnode1,pnode2 migrate_truecopydg-rg
```

- 5 Cree un recurso hasp-rs en el grupo de recursos creado en el Paso 4; establezca la propiedad globaldevicepaths en un grupo de dispositivos de disco básico. Este grupo de dispositivos se creó en el Paso 2.

```
clrs create -t HASToragePlus -x globaldevicepaths=mypool -g \
migrate_truecopydg-rg hasp2migrate_mypool
```

- 6 Si el grupo de recursos de aplicación se va a ejecutar en zonas locales, cree un grupo de recursos cuya lista de nodos contenga las pertinentes zonas locales. Las zonas globales que se correspondan con las zonas locales deben figurar en la lista de nodos del grupo de recursos creado en el Paso 4. Establezca el valor +++ en la propiedad rg\_affinities desde este grupo de recursos en el grupo de recursos creado en el Paso 4.

```
clrg create -n pnode1:zone-1,pnode2:zone-2 -p \
RG_affinities=+++migrate_truecopydg-rg sybase-rg
```

- 7 Cree un recurso HASToragePlus (hasp-rs) para el zpool creado en el Paso 3 en el grupo de recursos creado en el Paso 4 o en el Paso 6. Configure la propiedad resource\_dependencies para el recurso hasp-rs creado en el Paso 5.

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

- 8 Use el nuevo nombre del grupo de recursos cuando se precise un nombre de grupo de dispositivos.



## ▼ Creación de un nuevo grupo de discos al inicializar discos (Veritas Volume Manager)

---

**Nota** – Este procedimiento sólo es válido para inicializar discos. Si está encapsulando discos, use el procedimiento “[Creación de un nuevo grupo de discos al encapsular discos \(Veritas Volume Manager\)](#)” en la página 141.

---

Tras agregar el grupo de discos VxVM, debe registrar el grupo de dispositivos.

Si usa VxVM para configurar grupos de discos compartidos para Oracle RAC, use la función de clúster de VxVM como se describe en la *Guía de referencia del administrador de Veritas Volume Manager*.

- 1 **Conviértase en superusuario de cualquier nodo del clúster que esté *conectado físicamente* a los discos que componen el grupo de discos que se agrega.**
- 2 **Cree el grupo de discos y el volumen de VxVM.**

Use el método que prefiera para crear el grupo de discos y el volumen.

---

**Nota** – Si está configurando un volumen duplicado o reflejado, use Dirty Region Logging (DRL) para reducir el tiempo de recuperación transcurrido tras un error de nodo. Sin embargo, es posible que DRL reduzca el rendimiento de E/S.

---

Consulte la documentación de Veritas Volume Manager obtener información sobre los procedimientos para completar este paso.

- 3 **Registre el grupo de discos de VxVM como grupo de dispositivos de Oracle Solaris Cluster.**  
Consulte “[Registro de un grupo de discos como grupo de dispositivos \(Veritas Volume Manager\)](#)” en la página 145.

No registre los grupos de discos compartidos de Oracle RAC con la estructura de clúster.

## Mantenimiento de grupos de dispositivos

Puede llevar a cabo una serie de tareas administrativas para los grupos de dispositivos.

## Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)

Los grupos de dispositivos son conjuntos de discos de Solaris Volume Manager que se han registrado con Oracle Solaris Cluster. Para eliminar un grupo de dispositivos de Solaris Volume Manager, use los comandos `metaclear` y `metaset`. Dichos comandos eliminan el grupo de dispositivos que tenga el mismo nombre y anulan el registro del grupo de discos como grupo de dispositivos de Oracle Solaris Cluster.

Consulte la documentación de Solaris Volume Manager para saber los pasos que deben seguirse al eliminar un conjunto de discos.

### ▼ Eliminación de un nodo de todos los grupos de dispositivos

Siga este procedimiento para eliminar un nodo del clúster de todos los grupos de dispositivos que incluyen el nodo en sus listas de posibles nodos primarios.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo que va a eliminar de todos los grupos de dispositivos en calidad de nodo primario potencial.**
- 2 **Determine el grupo o los grupos de dispositivos donde el nodo que se va a eliminar figure como miembro.**

Busque el nombre del nodo en la Lista de nodos del grupo de dispositivos en cada uno de los dispositivos.

```
cldevicegroup list -v
```
- 3 **Si alguno de los grupos de dispositivos identificados en el [Paso 2](#) pertenece al tipo de grupo de dispositivos SVM, siga los pasos descritos en “[Eliminación de un nodo de un grupo de dispositivos \(Solaris Volume Manager\)](#)” en la [página 139](#) para todos los grupos de dispositivos de dicho tipo.**

- 4 Si alguno de los grupos de dispositivos identificados en el [Paso 2](#) pertenece al tipo de grupo de dispositivos VxVM, siga los pasos descritos en [“Eliminación de un nodo de un grupo de dispositivos \(Veritas Volume Manager\)” en la página 155](#) para todos los grupos de dispositivos de dicho tipo.
- 5 Determine los grupos de discos de dispositivos básicos de los cuales el nodo que se va a eliminar forma parte como miembro.  

```
cldevicegroup list -v
```
- 6 Si alguno de los grupos de dispositivos que figuran en la lista del [Paso 5](#) pertenece a los tipos de grupos Disco o Local\_Disk, siga los pasos descritos en [“Eliminación de un nodo de un grupo de dispositivos de discos básicos” en la página 157](#) para todos estos grupos de dispositivos.
- 7 Compruebe que el nodo se haya eliminado de la lista de posibles nodos primarios en todos los grupos de dispositivos.  

El comando no devuelve ningún resultado si el nodo ya no figura en la lista como nodo primario potencial en ninguno de los grupos de dispositivos.

```
cldevicegroup list -v nodename
```

## ▼ Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)

Siga este procedimiento para eliminar un nodo de clúster de la lista de nodos primarios potenciales de un grupo de dispositivos de Solaris Volume Manager. Repita el comando `metaset` en cada uno de los grupos de dispositivos donde desee eliminar el nodo.



**Precaución** – No ejecute `metaset -s nombre_conjunto -f -t` en un nodo del clúster que se arranque fuera del clúster si hay otros nodos activos como miembros del clúster y al menos uno de ellos es propietario del conjunto de discos.

`phys -schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que el nodo continúe como miembro del grupo de dispositivos y que este último sea un grupo de dispositivos de Solaris Volume Manager.

El tipo de grupo de dispositivos SDS/SVM indica que se trata de un grupo de dispositivos de Solaris Volume Manager.

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

- 2 Determine qué nodo es el primario del grupo de dispositivos.

```
cldevicegroup status devicegroup
```

- 3 Conviértase en superusuario del nodo propietario del grupo de dispositivos que desee modificar.

- 4 Elimine del grupo de dispositivos el nombre de host del nodo.

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

-s                      Especifica el nombre del grupo de dispositivos.  
*nombre\_conjunto*

-d                      Elimina del grupo de dispositivos los nodos identificados con -h.

-h *lista\_nodos*        Especifica el nombre del nodo o los nodos que se van a eliminar.

---

**Nota** – La actualización puede tardar varios minutos.

---

Si el comando falla, agregue la opción -f (forzar).

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

- 5 Repita el [Paso 4](#) con cada uno de los grupos de dispositivos donde desee eliminar el nodo como nodo primario potencial.

- 6 Compruebe que el nodo se haya eliminado del grupo de dispositivos.

El nombre del grupo de dispositivos coincide con el del conjunto de discos que se especifica con metaset.

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

### **Ejemplo 5–23** Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)

En el ejemplo siguiente se muestra cómo se elimina el nombre de host `phys-schost-2` de la configuración de un grupo de dispositivos. En este ejemplo se elimina `phys-schost-2` como nodo primario potencial del grupo de dispositivos designado. Compruebe la eliminación del nodo; para ello, ejecute el comando `cldevicegroup show`. Compruebe que el nodo eliminado ya no aparezca en el texto de la pantalla.

```
[Determine the Solaris Volume Manager
device group for the node:]
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
[Determine which node is the current primary for the device group:]
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
[Become superuser on the node that currently owns the device group.]
[Remove the host name from the device group:]
metaset -s dg-schost-1 -d -h phys-schost-2
[Verify removal of the node:]
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
```

## ▼ Creación de un nuevo grupo de discos al encapsular discos (Veritas Volume Manager)

**Nota** – Este procedimiento sólo es válido para encapsular discos. Si está inicializando discos, siga el procedimiento [“Creación de un nuevo grupo de discos al inicializar discos \(Veritas Volume Manager\)” en la página 137.](#)

Puede convertir discos que no sean root en grupos de dispositivos de Oracle Solaris Cluster si los encapsula como grupos de discos de VxVM y, a continuación, registra los grupos de discos como grupos de dispositivos de Oracle Solaris Cluster.

La encapsulación de discos sólo se admite durante la creación de un grupo de discos de VxVM. Tras crear un grupo de discos de VxVM y registrarlo como grupo de dispositivos de Oracle Solaris Cluster, sólo se deben agregar al grupo de discos los discos que puedan inicializarse.

Si usa VxVM para configurar grupos de discos compartidos para Oracle RAC, use la función de clúster de VxVM como se describe en la *Guía de referencia del administrador de Veritas Volume Manager*.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Si el disco que se está encapsulando tiene entradas del sistema de archivos en el archivo `/etc/vfstab`, compruebe que la opción `mount at boot` esté establecida en `no`.**

Vuélvala a establecer en `yes` una vez que el disco esté encapsulado y registrado como grupo de dispositivos de Oracle Solaris Cluster.

- 3 **Encapsule los discos.**

Use los menús `vxdiskadm` o la interfaz gráfica de usuario para encapsular los discos. VxVM requiere dos particiones libres y cilindros sin asignación al inicio o al final del disco. Asimismo, el segmento dos debe establecerse para todo el disco. Consulte la página de comando `man vxdiskadm` si desea obtener más información.

- 4 **Cierre el nodo y reinícielo.**

El comando `clnode evacuate` conmuta todos los grupos de recursos y de dispositivos, incluidos los nodos de votación no globales de un clúster global, del nodo especificado al siguiente nodo por orden de preferencia. Use el comando `shutdown` para cerrar y reiniciar el nodo.

```
clnode evacuate node[...]
shutdown -g0 -y -i6
```

- 5 **Si es necesario, conmute todos los grupos de recursos y de dispositivos para devolverlos al nodo original.**

Si los grupos de recursos y de dispositivos estaban configurados para conmutar al nodo primario en caso de error, este paso es innecesario.

```
cldevicegroup switch -n node devicegroup
clresourcegroup switch -z zone -n node resourcegroup
```

`nodo` El nombre del nodo.

`zona` El nombre del nodo sin votación, *nodo*, que puede controlar el grupo de recursos. Indique la *zona* sólo si ha especificado un nodo sin votación al crear el grupo de

recursos.

- 6 **Registre el grupo de discos de VxVM como grupo de dispositivos de Oracle Solaris Cluster.**  
Consulte “[Registro de un grupo de discos como grupo de dispositivos \(Veritas Volume Manager\)](#)” en la página 145.  
No registre los grupos de discos compartidos de Oracle RAC con la estructura de clúster.
- 7 Si la opción `mount at boot` se establece en `no` en el [Paso 2](#), vuélvala a establecer en `yes`.

## ▼ Adición de un volumen nuevo a un grupo de dispositivos ya creado (Veritas Volume Manager)

Al agregar un volumen nuevo a un grupo de dispositivos ya existente de VxVM, realice el procedimiento correspondiente desde el nodo primario del grupo de dispositivos en línea.

---

**Nota** – Después de agregar el volumen, debe registrar la modificación de la configuración mediante el procedimiento “[Registro de modificaciones en la configuración de grupos de discos \(Veritas Volume Manager\)](#)” en la página 148.

---

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.administer` en todos los nodos del clúster.**
- 2 **Determine el nodo primario para el grupo de dispositivos al que va a agregar el nuevo volumen.**  
`# cldevicegroup status`
- 3 **Si el grupo de dispositivos está fuera de línea, póngalo en línea.**

`# cldevicegroup switch -n nodename devicegroup`

*nombre\_nodo* Especifica el nombre del nodo al cual se conmuta el grupo de dispositivos. Este nodo se convierte en el nuevo nodo primario.

*grupo\_dispositivos* Especifica el grupo de dispositivos que se conmuta.

- 4 **En el nodo primario (el que controla el grupo de dispositivos), cree el volumen de VxVM en el grupo de discos.**

Consulte la documentación de Veritas Volume Manager para ver el procedimiento utilizado con el fin de crear el volumen de VxVM.

- 5 **Sincronice las modificaciones del grupo de discos de VxVM para actualizar el espacio de nombre global.**

```
cldevicegroup sync
```

[“Registro de modificaciones en la configuración de grupos de discos \(Veritas Volume Manager\)” en la página 148.](#)

## ▼ **Conversión de un grupo de discos en un grupo de dispositivos (Veritas Volume Manager)**

Un grupo de discos de VxVM ya creado puede convertirse en un grupo de dispositivos de Oracle Solaris Cluster si el grupo de discos se importa al nodo actual y, a continuación, se registra como grupo de dispositivos de Oracle Solaris Cluster.

- 1 **Conviértase en superusuario en un nodo de clúster.**
- 2 **Importe el grupo de discos de VxVM al nodo actual.**
- 3 **Registre el grupo de discos de VxVM como grupo de dispositivos de Oracle Solaris Cluster.**

Consulte [“Registro de un grupo de discos como grupo de dispositivos \(Veritas Volume Manager\)” en la página 145.](#)

## ▼ **Asignación de un número menor nuevo a un grupo de dispositivos (Veritas Volume Manager)**

Si el registro de un grupo de dispositivos no se completa correctamente porque el número menor entra en conflicto con otro grupo de discos, asigne al nuevo grupo de discos un número menor nuevo. Tras asignar el número menor nuevo, vuelva a ejecutar el procedimiento para registrar el grupo de discos como grupo de dispositivos de Oracle Solaris Cluster.

- 1 **Conviértase en superusuario en un nodo de clúster.**
- 2 **Determine los números menores que están en uso.**

```
ls -l /global/.devices/node@nodeid/dev/vx/dsk/*
```



- 3 Elija otro múltiplo de 1.000 que no se utilice como número menor de base para el nuevo grupo de discos.
- 4 Asigne el número menor nuevo al grupo de discos.  
`# vxdg remminor diskgroup base-minor-number`
- 5 Registre el grupo de discos de VxVM como grupo de dispositivos de Oracle Solaris Cluster.  
 Consulte “Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)” en la página 145.

#### Ejemplo 5–24 Asignación de un número menor nuevo a un grupo de dispositivos

En este ejemplo se usan los números menores 16000-16002 y 4000-4001. El comando `vxdg remminor` se utiliza para asignar el número menor de base 5000 al nuevo grupo de dispositivos.

```
ls -l /global/.devices/node@nodeid/dev/vx/dsk/*

/global/.devices/node@nodeid/dev/vx/dsk/dg1
brw----- 1 root root 56,16000 Oct 7 11:32 dg1v1
brw----- 1 root root 56,16001 Oct 7 11:32 dg1v2
brw----- 1 root root 56,16002 Oct 7 11:32 dg1v3

/global/.devices/node@nodeid/dev/vx/dsk/dg2
brw----- 1 root root 56,4000 Oct 7 11:32 dg2v1
brw----- 1 root root 56,4001 Oct 7 11:32 dg2v2
vxdg remminor dg3 5000
```

## ▼ Registro de un grupo de discos como grupo de dispositivos (Veritas Volume Manager)

Este procedimiento emplea la utilidad `clsetup` para registrar el grupo de discos de VxVM asociado como grupo de dispositivos de Oracle Solaris Cluster.

---

**Nota** – Tras haber registrado un grupo de dispositivos con el clúster, no importe ni exporte jamás un grupo de discos de VxVM mediante comandos de VxVM. Si realiza cambios en el grupo de discos de VxVM o en el volumen, siga el procedimiento “Registro de modificaciones en la configuración de grupos de discos (Veritas Volume Manager)” en la página 148 para registrar las modificaciones en la configuración del grupo de dispositivos. Este procedimiento garantiza que el espacio de nombre global se mantenga en estado correcto.

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

**Antes de empezar**

Antes de registrar un grupo de dispositivos de VxVM, compruebe que se hayan cumplido los requisitos siguientes:

- Se cuenta con privilegios de superusuario en un nodo del clúster.
- El nombre del grupo de discos de VxVM se registrará como grupo de dispositivos.
- Hay un orden de preferencia de los nodos para controlar el grupo de dispositivos.
- Hay un número de nodos secundarios para el grupo de dispositivos.

Al definir el orden de preferencia, también se indica si el grupo de dispositivos debe conmutarse al primer nodo por orden de preferencia si se da un error en el nodo y posteriormente se vuelve al clúster.

Consulte `cldevicegroup(1CL)` si desea obtener más información sobre las opciones de preferencias y conmutación de procesos de nodos.

Los nodos no primarios del clúster (de repuesto) realizan una transición para convertirse en nodos secundarios de según el orden de preferencias de los nodos. El número predeterminado de nodos secundarios de un grupo de dispositivos suele establecerse en uno. Esta configuración predeterminada reduce al mínimo la disminución del rendimiento debido a los puntos de control de distintos nodos secundarios durante el funcionamiento normal. Por ejemplo, en un clúster de cuatro nodos, el comportamiento predeterminado configura un nodo primario, uno secundario y dos de repuesto. Consulte también “[Establecimiento del número de secundarios para un grupo de dispositivos](#)” en la página 160.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Inicie la utilidad `clsetup`.**

```
clsetup
```

Aparece el menú principal.

- 3 **Para trabajar con grupos de dispositivos de VxVM, escriba el número correspondiente a la opción de los grupos de dispositivos y los volúmenes.**

Se muestra el menú Grupos de dispositivos.

- 4 **Para registrar un grupo de dispositivos de VxVM, escriba el número correspondiente a la opción de registrar un grupo de discos de VxVM como grupo de dispositivos.**

Siga las instrucciones y escriba el nombre del grupo de discos de VxVM que se vaya a registrar como grupo de dispositivos de Oracle Solaris Cluster.

Si este grupo de dispositivos se replica mediante replicación basada en almacenamiento, dicho nombre debe coincidir con el del nombre del grupo de replicaciones.

Si usa VxVM para configurar grupos de discos compartidos para Oracle Parallel Server u Oracle RAC, no registre los grupos de discos compartidos en la estructura del clúster. Use la función de clúster de VxVM como se describe en *Veritas Volume Manager Administrator's Reference Guide*.

- 5 **Si aparece el error siguiente al intentar registrar el grupo de dispositivos, cambie el número menor del grupo de dispositivos.**

```
cldevicegroup: Failed to add device group - in use
```

Para cambiar el número menor del grupo de dispositivos, siga el procedimiento [“Asignación de un número menor nuevo a un grupo de dispositivos \(Veritas Volume Manager\)” en la página 144](#). Este procedimiento permite asignar un número menor que no entre en conflicto con un número menor utilizado por otro grupo de dispositivos.

- 6 **Si configura un grupo de dispositivos replicado, establezca la propiedad de replicación para el grupo de dispositivos.**

```
cldevicegroup sync devicegroup
```

- 7 **Compruebe que el grupo de dispositivos esté registrado y en línea.**

Si el grupo de dispositivos está registrado correctamente, la información del nuevo grupo de dispositivos se muestra al usar el comando siguiente.

```
cldevicegroup status devicegroup
```

---

**Nota** – Si modifica información de configuración de un grupo de discos o un volumen de VxVM registrado en el clúster, debe sincronizar el grupo de dispositivos con `clsetup`. Entre dichas modificaciones de la configuración están agregar o eliminar volúmenes, así como modificar el grupo, el propietario o los permisos de los volúmenes. Volver a registrar registro después de modificar la configuración asegura que el espacio de nombre global se mantenga en el estado correcto. Consulte [“Actualización del espacio de nombre de dispositivos globales” en la página 126](#).

---

### Ejemplo 5-25 Registro de un grupo de dispositivos de Veritas Volume Manager

El ejemplo siguiente muestra el comando `cldevicegroup` generado por `clsetup` cuando registra un grupo de dispositivos de VxVM (`dg1`) y el paso de comprobación. En este ejemplo se supone que el grupo de discos y el volumen de VxVM ya se habían creado.

```
clsetup

cldevicegroup create -t vxvm -n phys-schost-1,phys-schost-2 -p failback=true dg1

cldevicegroup status dg1

=== Cluster Device Groups ===

--- Device Group Status ---

Device Group Name Primary Secondary Status

dg1 phys-schost-1 phys-schost-2 Online
```

**Véase también** Para crear un sistema de archivos de clúster en el grupo de dispositivos de VxVM, consulte [“Adición de un sistema de archivos de clúster” en la página 174](#).

Si hay problemas con el número menor, consulte [“Asignación de un número menor nuevo a un grupo de dispositivos \(Veritas Volume Manager\)” en la página 144](#).

## ▼ Registro de modificaciones en la configuración de grupos de discos (Veritas Volume Manager)

Si modifica algún elemento de la información de configuración de un grupo de discos o volumen de VxVM, es necesario registrar las modificaciones en la configuración del grupo de dispositivos de Oracle Solaris Cluster. La operación de registro garantiza que el espacio de nombre se mantenga en estado correcto.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

**2 Inicie la utilidad `clsetup`.**

```
clsetup
```

Aparece el menú principal.

**3 Para trabajar con grupos de dispositivos de VxVM, escriba el número correspondiente a la opción de los grupos de dispositivos y los volúmenes.**

Se muestra el menú Grupos de dispositivos.

**4 Para registrar las modificaciones en la configuración, escriba el número correspondiente a la opción de sincronizar la información del volumen para un grupo de dispositivos de VxVM.**

Siga las instrucciones y escriba el nombre del grupo de discos de VxVM cuya configuración se haya modificado.

**Ejemplo 5–26 Registro de las modificaciones en la configuración de grupos de discos de Veritas Volume Manager**

El ejemplo siguiente muestra el comando `cldevicegroup` generado por `clsetup` al registrarse un grupo de dispositivos de VxVM modificado ( `dg1`). En este ejemplo se supone que el grupo de discos y el volumen de VxVM ya se habían creado.

```
clsetup
```

```
cldevicegroup sync dg1
```

## ▼ Conversión de un grupo de discos local en un grupo de dispositivos (VxVM)

Siga este procedimiento para cambiar un grupo de discos de VxVM local a un grupo de dispositivos de VxVM al que se pueda acceder globalmente.

**1 Conviértase en superusuario en un nodo de clúster.****2 Inicie la utilidad `clsetup`.**

```
clsetup
```

**3 Anule la configuración de la propiedad `localonly`.**

a. Elija la opción del menú Grupos de dispositivos y volúmenes.

b. Elija la opción del menú Restablecer un grupo de discos local a un grupo de dispositivos de VxVM.

- c. Siga las instrucciones para anular la configuración de la propiedad `localonly`.
- 4 Especifique los nodos que pueden actuar controlar el grupo de discos.
  - a. Vuelva al menú principal de la utilidad `clsetup`.
  - b. Elija la opción del menú Grupos de dispositivos y volúmenes.
  - c. Elija la opción del menú para registrar un grupo de discos.
  - d. Siga las instrucciones para especificar los nodos que pueden controlar el grupo de discos.
  - e. Cuando haya terminado, salga de la utilidad `clsetup`.
- 5 Compruebe que se haya configurado el grupo de dispositivos.  
`phys-schost# cldevicegroup show`

## ▼ Conversión de un grupo de dispositivos en un grupo de discos local (VxVM)

Siga este procedimiento para convertir un grupo de dispositivos de VxVM en un grupo de discos de VxVM local no administrado por el software Oracle Solaris Cluster. El grupo de discos local puede tener más de un nodo en su lista de nodos, pero sólo lo puede controlar un único nodo a la vez.

- 1 Conviértase en superusuario en un nodo del clúster.
- 2 Ponga fuera de línea el grupo de dispositivos.  
`phys-schost# cldevicegroup offline devicegroup`
- 3 Anule el registro del grupo de dispositivos.
  - a. Inicie la utilidad `clsetup`.  
`phys-schost# clsetup`
  - b. Elija la opción del menú Grupos de dispositivos y volúmenes.
  - c. Elija la opción del menú Anular el registro de un grupo de discos de VxVM.
  - d. Siga las instrucciones para especificar el grupo de discos de VxVM cuyo registro va a anular en el software Oracle Solaris Cluster.

e. Cierre la utilidad `clsetup`.

**4 Compruebe que el grupo de discos ya no esté registrado en el software Oracle Solaris Cluster.**

```
phys-schost# cldevicegroup status
```

En la salida del comando ya no debería aparecer el grupo de dispositivos cuyo registro ha anulado.

**5 Importe el grupo de discos.**

```
phys-schost# vxdg import diskgroup
```

**6 Establezca la propiedad `localonly` del grupo de discos.**

a. Inicie la utilidad `clsetup`.

```
phys-schost# clsetup
```

b. Elija la opción del menú Grupos de dispositivos y volúmenes.

c. Elija la opción del menú Definir un grupo de discos VxVM como grupo de discos local.

d. Siga las instrucciones para establecer la propiedad `localonly` y especificar el único nodo que controlará el grupo de discos.

e. Cuando haya terminado, salga de la utilidad `clsetup`.

**7 Compruebe que el grupo de discos esté configurado correctamente como grupo de discos local.**

```
phys-schost# vxdg list diskgroup
```

## ▼ Eliminación de un volumen de un grupo de dispositivos (Veritas Volume Manager)

---

**Nota** – Tras eliminar el volumen del grupo de dispositivos, debe registrar las modificaciones en la configuración hechas en el grupo de dispositivos mediante el procedimiento “[Registro de modificaciones en la configuración de grupos de discos \(Veritas Volume Manager\)](#)” en la página 148.

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en cualquier nodo del clúster.**
- 2 **Determine el nodo primario y el estado del grupo de dispositivos.**  
`# cldevicegroup status devicegroup`
- 3 **Si el grupo de dispositivos está fuera de línea, póngalo en línea.**  
`# cldevicegroup online devicegroup`
- 4 **En el nodo primario (el que controla el grupo de dispositivos), elimine el volumen de VxVM en el grupo de discos.**  
`# vxedit -g diskgroup -rf rm volume`  
`-g grupo_discos`      Especifica el grupo de discos de VxVM que contiene el volumen.  
`-rf rm volumen`      Elimina el volumen indicado. La opción `-r` convierte la operación en recursiva. La opción `-f` es necesaria para eliminar un volumen habilitado.
- 5 **Con la utilidad `clsetup`, registre las modificaciones en la configuración del grupo de dispositivos para actualizar el espacio de nombre global.**  
Consulte [“Registro de modificaciones en la configuración de grupos de discos \(Veritas Volume Manager\)” en la página 148.](#)

## ▼ Eliminación y anulación del registro de un grupo de dispositivos (Veritas Volume Manager)

Eliminar un grupo de dispositivos de Oracle Solaris Cluster hace que el grupo de discos de VxVM correspondiente se exporte en lugar de destruirse. Sin embargo, aunque el grupo de discos de VxVM siga existiendo, no se puede utilizar en el clúster a menos que se vuelva a registrar.

Este procedimiento emplea la utilidad `clsetup` para eliminar un grupo de discos de VxVM y anular su registro como grupo de dispositivos de Oracle Solaris Cluster.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.



Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Ponga fuera de línea el grupo de dispositivos.**

```
cldevicegroup offline devicegroup
```

- 3 **Inicie la utilidad `clsetup`.**

```
clsetup
```

Aparece el menú principal.

- 4 **Para trabajar con grupos de dispositivos de VxVM, escriba el número correspondiente a la opción de los grupos de dispositivos y los volúmenes.**

Se muestra el menú Grupos de dispositivos.

- 5 **Para anular el registro de un grupo de discos de VxVM, escriba el número correspondiente a la opción de anular el registro de un grupo de dispositivos de VxVM.**

Siga las instrucciones y escriba el nombre del grupo de discos de VxVM cuyo registro se vaya a anular.

### **Ejemplo 5–27 Eliminación y anulación del registro de un grupo de dispositivos de Veritas Volume Manager**

El ejemplo siguiente muestra el grupo de dispositivos de VxVM `dg1` puesto fuera de línea y el comando `cldevicegroup` generado por `clsetup` cuando elimina y anula el registro del grupo de dispositivos.

```
cldevicegroup offline dg1
clsetup
cldevicegroup delete dg1
```

## **▼ Adición de un nodo a un grupo de dispositivos (Veritas Volume Manager)**

Este procedimiento agrega un nodo a un grupo de dispositivos mediante la utilidad `clsetup`.

Los requisitos para agregar un nodo a un grupo de dispositivos de VxVM son:

- Se cuenta con privilegios de superusuario en un nodo del clúster.

- El nombre del grupo de dispositivos de VxVM al que se vaya a agregar el nodo.
- El nombre o ID de nodo de los nodos que se vayan a agregar.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en cualquier nodo del clúster.**
- 2 Inicie la utilidad `clsetup`.**  
`# clsetup`  
Aparece el menú principal.
- 3 Para trabajar con grupos de dispositivos de VxVM, escriba el número correspondiente a la opción de los grupos de dispositivos y los volúmenes.**  
Se muestra el menú Grupos de dispositivos.
- 4 Para agregar un nodo a un grupo de dispositivos de VxVM, escriba el número correspondiente a la opción de agregar un nodo a un grupo de dispositivos de VxVM.**  
Siga las instrucciones y escriba los nombres de nodo y de grupo de dispositivos.
- 5 Compruebe que se haya agregado el nodo.**  
Busque la información del grupo de dispositivos para el nuevo disco, que se muestra con el comando siguiente.  
`# cldevicegroup show devicegroup`

### **Ejemplo 5–28** Adición de un nodo a un grupo de dispositivos de Veritas Volume Manager

En el ejemplo siguiente se muestra el comando `cldevicegroup` generado por `clsetup` cuando agrega un nodo (`phys-schost-3`) a un grupo de dispositivos de VxVM (`dg1`) y el paso de comprobación.

```
clsetup
cldevicegroup add-node -n phys-schost-3 dg1
cldevicegroup show dg1

=== Device Groups ===

Device Group Name: dg1
Type: VxVM
failback: yes
```

```

Node List: phys-schost-1, phys-schost-3
preferred: no
numsecondaries: 1
diskgroup names: dg1

```

## ▼ Eliminación de un nodo de un grupo de dispositivos (Veritas Volume Manager)

Siga este procedimiento para eliminar un nodo del clúster de la lista de nodos primarios potenciales de un grupo de dispositivos (o grupo de discos) de Veritas Volume Manager (o VxVM).

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### 1 Compruebe que el nodo continúe como miembro del grupo y que este último sea un grupo de dispositivos de VxVM.

El tipo de grupo de dispositivos VxVM indica que se trata de un grupo de dispositivos de VxVM.

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

### 2 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en un nodo del clúster actual.

### 3 Inicie la utilidad `clsetup`.

```
clsetup
```

Aparece el menú principal.

### 4 Para reconfigurar un grupo de dispositivos, escriba el número correspondiente a la opción de los grupos de dispositivos y volúmenes.

### 5 Para eliminar el nodo del grupo de dispositivos de VxVM, escriba el número correspondiente a la opción de eliminar un nodo de un grupo de dispositivos de VxVM.

Siga los indicadores para eliminar un nodo del clúster del grupo de dispositivos. Se le solicitará que facilite la información siguiente:

- Grupo de dispositivos de VxVM
- Nombre del nodo

**6 Compruebe que el nodo se haya eliminado del grupo o los grupos de dispositivos de VxVM.**

```
cldevicegroup show devicegroup
```

**Ejemplo 5–29 Eliminación de un nodo de un grupo de dispositivos (VxVM)**

En este ejemplo se muestra la eliminación del nodo denominado `phys-schost-1` del grupo de dispositivos `dg1` de VxVM.

[Determine the VxVM device group for the node:]

```
cldevicegroup show dg1
```

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

```
Device Group Name: dg1
Type: VxVM
failback: no
Node List: phys-schost-1, phys-schost-2
preferenced: no
numsecondaries: 1
diskgroup names: dg1
```

[Become superuser and start the `clsetup` utility:]

```
clsetup
```

Select Device groups and volumes>Remove a node from a VxVM device group.

Answer the questions when prompted.

You will need the following information.

| Name:                  | Example:      |
|------------------------|---------------|
| VxVM device group name | dg1           |
| node names             | phys-schost-1 |

[Verify that the `cldevicegroup` command executed properly:]

```
cldevicegroup remove-node -n phys-schost-1 dg1
```

Command completed successfully.

Dismiss the `clsetup` Device Groups Menu and Main Menu.

[Verify that the node was removed:]

```
cldevicegroup show dg1
```

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

```
Device Group Name: dg1
Type: VxVM
failback: no
Node List: phys-schost-2
preferenced: no
numsecondaries: 1
device names: dg1
```

## ▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos

Siga este procedimiento para eliminar un nodo de clúster de la lista de nodos primarios potenciales de un grupo de dispositivos de discos básicos.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en uno de los nodos del clúster *distinto del nodo que se vaya a eliminar*.**
- 2 **Identifique los grupos de dispositivos conectados al nodo que se vaya a eliminar y determine cuáles son grupos de dispositivos de discos básicos.**  

```
cldevicegroup show -n nodename -t rawdisk +
```
- 3 **Inhabilite la propiedad `localonly` de cada grupo de dispositivos de discos básicos `Local_Disk`.**  

```
cldevicegroup set -p localonly=false devicegroup
```

Consulte la página de comando `man cldevicegroup(1CL)` si desea obtener más información sobre la propiedad `localonly`.
- 4 **Compruebe que haya inhabilitado la propiedad `localonly` en todos los grupos de dispositivos de discos básicos conectados al nodo que vaya a eliminar.**  

El tipo de grupo de dispositivos `Disk` indica que la propiedad `localonly` está inhabilitada para ese grupo de dispositivos de discos básicos.

```
cldevicegroup show -n nodename -t rawdisk -v +
```
- 5 **Elimine el nodo de todos los grupos de dispositivos de discos básicos identificados en el [Paso 2](#).**  

Complete este paso en cada grupo de dispositivos de discos básicos conectado con el nodo que se va a eliminar.

```
cldevicegroup remove-node -n nodename devicegroup
```

### Ejemplo 5–30 Eliminación de un nodo de un grupo de dispositivos básicos

En este ejemplo se muestra cómo eliminar un nodo (`phys - schost - 2`) de un grupo de dispositivos de discos básicos. Todos los comandos se ejecutan desde otro nodo del clúster (`phys - schost - 1`).

[Identify the device groups connected to the node being removed, and determine which are raw-disk device groups:]

```
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/d2
Type: VxVM
failback: true
Node List: pbrave2
preferenced: false
localonly: false
autogen: true
numsecondaries: 1
diskgroup name: vxdg1
```

```
Device Group Name: dsk/d1
Type: SVM
failback: false
Node List: pbrave1, 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
```

[Disable the localonly flag for each local disk on the node:]

```
phys-schost-1# cldevicegroup set -p localonly=false dsk/d4
```

[Verify that the localonly flag is disabled:]

```
phys-schost-1# cldevicegroup show -n phys-schost-2 -t rawdisk +
```

```
(dsk/d4) Device group type: Disk
(dsk/d8) Device group type: Local_Disk
```

[Remove the node from all raw-disk device groups:]

```
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
```

## ▼ Cambio de propiedades de los grupos de dispositivos

El método para establecer la propiedad primaria de un grupo de dispositivos se basa en el establecimiento de un atributo de preferencia de propiedad denominado `preferenced`. Si el atributo no está definido, el propietario primario de un grupo de dispositivos que no está sujeto

a ninguna otra relación de propiedad es el primer nodo que intente acceder a un disco de dicho grupo. Sin embargo, si este atributo está definido, debe especificar el orden de preferencia en el cual los nodos intentan establecer la propiedad.

Si inhabilita el atributo `preferred`, el atributo `failback` también se inhabilita automáticamente. Sin embargo, si intenta habilitar o volver a habilitar el atributo `preferred`, debe elegir entre habilitar o inhabilitar el atributo `failback`.

Si el atributo `preferred` se habilita o inhabilita, se solicitará que reestablecer el orden de los nodos en la lista de preferencias de propiedades primarias.

Este procedimiento emplea el comando `clsetup` para establecer o anular la definición de los atributos `preferred` y `failback` para los grupos de dispositivos de Solaris Volume Manager o VxVM.

**Antes de empezar**

Para llevar a cabo este procedimiento, necesita el nombre del grupo de dispositivos para el cual está cambiando valores de atributos.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en cualquier nodo del clúster.**
- 2 Inicie la utilidad `clsetup`.**  
`# clsetup`  
 Aparece el menú principal.
- 3 Para trabajar con grupos de dispositivos, escriba el número correspondiente a la opción de los grupos de dispositivos y los volúmenes.**  
 Se muestra el menú Grupos de dispositivos.
- 4 Para modificar las propiedades esenciales de un grupo de dispositivos, escriba el número correspondiente a la opción de modificar las propiedades esenciales de un grupo de dispositivos de VxVM o Solaris Volume Manager.**  
 Se muestra el menú Cambiar las propiedades esenciales.

- 5 **Para modificar una propiedad de un grupo de dispositivos, escriba el número correspondiente a la opción de cambiar las propiedades de preferencia o de conmutación de procesos.**

Siga las instrucciones para definir las opciones de `preferenced` y `failback` de un grupo de dispositivos.

- 6 **Compruebe que se hayan modificado los atributos del grupo de dispositivos.**

Busque la información del grupo de dispositivos mostrada por el comando siguiente.

```
cldevicegroup show -v devicegroup
```

### Ejemplo 5–31 Modificación de las propiedades de grupos de dispositivos

En el ejemplo siguiente se muestra el comando `cldevicegroup` generado por `clsetup` cuando define los valores de los atributos para un grupo de dispositivos (`dg-schost-1`).

```
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 |
| preferenced:       | yes                          |
| numsecondaries:    | 1                            |
| diskset names:     | dg-schost-1                  |

## ▼ Establecimiento del número de secundarios para un grupo de dispositivos

La propiedad `numsecondaries` especifica el número de nodos dentro de un grupo de dispositivos que pueden controlar el grupo si falla el nodo primario. El número predeterminado de nodos secundarios para servicios de dispositivos es de uno. El valor puede establecerse como cualquier número entero entre uno y la cantidad de nodos proveedores no primarios operativos del grupo de dispositivos.

Este valor de configuración es importante para equilibrar el rendimiento y la disponibilidad del clúster. Por ejemplo, incrementar el número de nodos secundarios aumenta las oportunidades de que un grupo de dispositivos sobreviva a múltiples errores simultáneos dentro de un clúster. Incrementar el número de nodos secundarios también reduce el rendimiento habitualmente durante el funcionamiento normal. Un número menor de nodos secundarios suele mejorar el rendimiento, pero reduce la disponibilidad. Ahora bien, un número superior de nodos secundarios no siempre implica una mayor disponibilidad del sistema de archivos o del grupo



de dispositivos. Consulte el [Capítulo 3, “Key Concepts for System Administrators and Application Developers”](#) de *Oracle Solaris Cluster Concepts Guide* si desea obtener más información.

Si modifica la propiedad `numsecondaries`, se agregan o eliminan nodos secundarios en el grupo de dispositivos si la modificación causa una falta de coincidencia entre el número real de nodos secundarios y el número deseado.

Este procedimiento emplea la utilidad `clsetup` para establecer la propiedad `numsecondaries` en todos los tipos de grupos de dispositivos. Consulte `cldevicegroup(1CL)` si desea obtener información sobre las opciones de los grupos de dispositivos al configurar cualquier grupo de dispositivos.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en cualquier nodo del clúster.**
- 2 **Inicie la utilidad `clsetup`.**  
`# clsetup`  
 Aparece el menú principal.
- 3 **Para trabajar con grupos de dispositivos, seleccione la opción del menú Grupos de dispositivos y volúmenes.**  
 Se muestra el menú Grupos de dispositivos.
- 4 **Para modificar las propiedades esenciales de un grupo de dispositivos, seleccione la opción Cambiar las propiedades esenciales de un grupo de dispositivos.**  
 Se muestra el menú Cambiar las propiedades esenciales.
- 5 **Para modificar el número de nodos secundarios, escriba el número correspondiente la opción de cambiar la propiedad `numsecondaries`.**  
 Siga las instrucciones y escriba el número de nodos secundarios que se van a configurar para el grupo de dispositivos. El comando `cldevicegroup` correspondiente se ejecuta a continuación, se imprime un registro y la utilidad vuelve al menú anterior.

**6 Valide la configuración del grupo de dispositivos.**

```
cldevicegroup show dg-schost-1
=== Device Groups ===

Device Group Name: dg-schost-1
Type: VxVm This might also be SDS or Local_Disk.
failback: yes
Node List: phys-schost-1, phys-schost-2 phys-schost-3
preferenced: yes
numsecondaries: 1
diskgroup names: dg-schost-1
```

---

**Nota** – Si modifica información de configuración de un grupo de discos o un volumen de VxVM registrado en el clúster, debe volver a registrar el grupo de dispositivos con `clsetup`. Entre dichas modificaciones de la configuración están agregar o eliminar volúmenes, así como modificar el grupo, el propietario o los permisos de los volúmenes. Volver a registrar registro después de modificar la configuración asegura que el espacio de nombre global se mantenga en el estado correcto. Consulte [“Actualización del espacio de nombre de dispositivos globales” en la página 126](#).

---

**7 Compruebe que se haya modificado el atributo del grupo de dispositivos.**

Busque la información del grupo de dispositivos mostrada por el comando siguiente.

```
cldevicegroup show -v devicegroup
```

**Ejemplo 5–32 Modificación del número de nodos secundarios (Solaris Volume Manager)**

En el ejemplo siguiente se muestra el comando `cldevicegroup` que genera `clsetup` al configurar el número de nodos secundarios para un grupo de dispositivos (`dg-schost-1`). En este ejemplo se supone que el grupo de discos y el volumen ya se habían creado.

```
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
```

**Ejemplo 5–33** Establecimiento del número de nodos secundarios (Veritas Volume Manager)

El ejemplo siguiente muestra el comando `cldevicegroup` generado por `clsetup` al definir en dos el número de nodos secundarios para un grupo de dispositivos (`dg-schost-1`). Consulte [“Establecimiento del número de secundarios para un grupo de dispositivos” en la página 160](#) si desea obtener información sobre cómo modificar el número de secundarios tras haber creado un grupo de dispositivos.

```
cldevicegroup set -p numsecondaries=2 dg-schost-1

cldevicegroup show dg-schost-1
=== Device Groups ===

Device Group Name: dg-schost-1
Type: VxVM
failback: yes
Node List: phys-schost-1, phys-schost-2
preferenced: yes
numsecondaries: 1
diskgroup names: dg-schost-1
```

**Ejemplo 5–34** Definición del número de nodos secundarios con el valor predeterminado

En el ejemplo siguiente se muestra el uso de un valor nulo de secuencia de comandos para configurar el número predeterminado de nodos secundarios. El grupo de dispositivos se configura para usar el valor predeterminado, aunque dicho valor sufra modificaciones.

```
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
```

## ▼ Enumeración de la configuración de grupos de dispositivos

No es necesario ser un superusuario para enumerar un listado con la configuración. Sin embargo, se debe disponer de autorización `solaris.cluster.read`.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

● **Utilice uno de los métodos de la lista siguiente.**

|                                                      |                                                                                                                         |
|------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------|
| Oracle Solaris Cluster Manager GUI                   | Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.                              |
| <code>cldevicegroup show</code>                      | Use <code>cldevicegroup show</code> para enumerar la configuración de todos los grupos de dispositivos del clúster.     |
| <code>cldevicegroup show grupo_dispositivos</code>   | Use <code>cldevicegroup show grupo_dispositivos</code> para enumerar la configuración de un solo grupo de dispositivos. |
| <code>cldevicegroup status grupo_dispositivos</code> | Use <code>cldevicegroup status grupo_dispositivos</code> para determinar el estado de un solo grupo de dispositivos.    |
| <code>cldevicegroup status +</code>                  | Use <code>cldevicegroup status +</code> para determinar el estado de todos los grupos de dispositivos del clúster.      |

Use la opción `-v` con cualquiera de estos comandos para obtener información más detallada.

**Ejemplo 5-35** Enumeración del estado de todos los grupos de dispositivos

```
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
```

**Ejemplo 5-36** Enumeración de la configuración de un determinado grupo de dispositivos

```
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                  |

## ▼ Conmutación al nodo primario de un grupo de dispositivos

Este procedimiento también es válido para iniciar (poner en línea) un grupo de dispositivos inactivo.

También puede poner en línea un grupo de dispositivos inactivo o conmutar el nodo primario de un grupo de dispositivos mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma un perfil que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Use `cldevicegroup switch` para conmutar el nodo primario del grupo de dispositivos.**

```
cldevicegroup switch -n nodename devicegroup
```

`-n nombre_nodo` Especifica el nombre del nodo al que conmutar. Este nodo se convierte en el nuevo nodo primario.

`grupo_dispositivos` Especifica el grupo de dispositivos que se conmuta.

- 3 **Compruebe que el grupo de dispositivos se haya conmutado al nuevo nodo primario.**

Si el grupo de dispositivos está registrado correctamente, la información del nuevo grupo de dispositivos se muestra al usar el comando siguiente.

```
cldevice status devicegroup
```

### Ejemplo 5-37 Conmutación del nodo primario de un grupo de dispositivos

En el ejemplo siguiente se muestra cómo conmutar el nodo primario de un grupo de dispositivos y cómo comprobar el cambio.

```
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
```

▼ **Colocación de un grupo de dispositivos en estado de mantenimiento**

Poner un grupo de dispositivos en estado de mantenimiento impide que dicho grupo se ponga automáticamente en línea cada vez que se obtenga acceso a uno de sus dispositivos. Conviene poner un grupo de dispositivos en estado de mantenimiento al finalizar procedimientos de reparación que requieran consentimiento para todas las actividades E/S hasta que terminen las operaciones de reparación. Asimismo, poner un grupo de dispositivos en estado de mantenimiento ayuda a impedir la pérdida de datos, al asegurar que el grupo de dispositivos no se pone en línea en un nodo mientras el conjunto o el grupo de discos se esté reparando en otro nodo.

Para obtener instrucciones sobre cómo restaurar un conjunto de discos dañado, consulte [“Restauración de un conjunto de discos dañado” en la página 299](#).

---

**Nota** – Antes de poder colocar un grupo de dispositivos en estado de mantenimiento, deben detenerse todos los accesos a sus dispositivos y desmontarse todos los sistemas de archivos dependientes.

---

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

**1 Ponga el grupo de dispositivos en estado de mantenimiento.**

**a. Si el grupo de dispositivos está habilitado, inhabilitelo.**

```
cldevicegroup disable devicegroup
```

**b. Ponga fuera de línea el grupo de dispositivos.**

```
cldevicegroup offline devicegroup
```

**2 Si el procedimiento de reparación que se efectúa requiere la propiedad de un conjunto o un grupo de discos, importe manualmente ese conjunto o grupo de discos.**

En Solaris Volume Manager:

```
metaset -C take -f -s diskset
```



**Precaución** – Si va a asumir la propiedad de un conjunto de discos de Solaris Volume Manager, *debe* usar el comando `metaset -C take` cuando el grupo de dispositivos esté en estado de mantenimiento. Al usar `metaset -t`, el grupo de dispositivos se pone en línea como parte del proceso de pasar a ser propietario. Si va a importar un grupo de discos de VxVM, debe usar el indicador `-t` al importar el grupo de discos. Al usar el indicador `-t`, se impide que el grupo de discos se importe automáticamente si se reanuda este nodo.

En Veritas Volume Manager:

```
vxdg -t import disk-group-name
```

**3 Complete el procedimiento de reparación que debe realizar.****4 Deje libre la propiedad del conjunto o del grupo de discos.**

**Precaución** – Antes de sacar el grupo de dispositivos fuera del estado de mantenimiento, debe liberar la propiedad del conjunto o grupo de discos. Si hay un error al liberar la propiedad, pueden perderse datos.

- En Solaris Volume Manager:

```
metaset -C release -s diskset
```

- En Veritas Volume Manager:

```
vxdg deport diskgroupname
```

**5 Ponga en línea el grupo de dispositivos.**

```
cldevicegroup online devicegroup
cldevicegroup enable devicegroup
```

**Ejemplo 5–38 Colocación de un grupo de dispositivos en estado de mantenimiento**

Este ejemplo muestra cómo poner el grupo de dispositivos `dg-schost-1` en estado de mantenimiento y cómo sacarlo de dicho estado.

```
[Place the device group in maintenance state.]
cldevicegroup disable dg-schost-1
cldevicegroup offline dg-schost-1
[If needed, manually import the disk set or disk group.]
For Solaris Volume Manager:
metaset -C take -f -s dg-schost-1
For Veritas Volume Manager:
vxdg -t import dg1

[Complete all necessary repair procedures.]

[Release ownership.]
For Solaris Volume Manager:
metaset -C release -s dg-schost-1
For Veritas Volume Manager:
vxdg deport dg1

[Bring the device group online.]
cldevicegroup online dg-schost-1
cldevicegroup enable dg-schost-1
```

## Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento

La instalación del software Oracle Solaris Cluster asigna reservas de SCSI automáticamente a todos los dispositivos de almacenamiento. Siga los procedimientos descritos a continuación para comprobar la configuración de los dispositivos y, si es necesario, para anular la configuración de un dispositivo.

- “Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento” en la página 168
- “Visualización del protocolo SCSI de un solo dispositivo de almacenamiento” en la página 169
- “Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento” en la página 170
- “Modificación del protocolo de protección en un solo dispositivo de almacenamiento” en la página 172

### ▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.



Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione autorización de RBAC `solaris.cluster.read`.**
- 2 **Desde cualquier nodo, visualice la configuración del protocolo SCSI predeterminado global.**

```
cluster show -t global
```

Para obtener más información, consulte la página de comando man [cluster\(1CL\)](#).

### **Ejemplo 5–39 Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento**

En el ejemplo siguiente se muestra la configuración del protocolo SCSI de todos los dispositivos de almacenamiento del clúster.

```
cluster show -t global
```

```
=== Cluster ===
```

|                    |                                |
|--------------------|--------------------------------|
| Cluster Name:      | racerxx                        |
| installmode:       | disabled                       |
| heartbeat_timeout: | 10000                          |
| heartbeat_quantum: | 1000                           |
| private_netaddr:   | 172.16.0.0                     |
| private_netmask:   | 255.255.248.0                  |
| max_nodes:         | 64                             |
| max_privatenets:   | 10                             |
| global_fencing:    | pathcount                      |
| Node List:         | phys-racerxx-1, phys-racerxx-2 |

## **▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento**

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione autorización de RBAC `solaris.cluster.read`.**

## 2 Desde cualquier nodo, visualice la configuración del protocolo SCSI del dispositivo de almacenamiento.

```
cldevice show device
```

*dispositivo*      Nombre de la ruta del dispositivo o un nombre de dispositivo.

Para obtener más información, consulte la página de comando man [cldevice\(1CL\)](#).

### Ejemplo 5–40 Visualización del protocolo SCSI de un solo dispositivo

En el ejemplo siguiente se muestra el protocolo SCSI del dispositivo `/dev/rdisk/c4t8d0`.

```
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                   |

## ▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento

La protección se puede activar o desactivar globalmente para todos los dispositivos de almacenamiento conectados a un clúster. La configuración de protección predeterminada de un solo dispositivo de almacenamiento anula la configuración global al establecer la protección predeterminada del dispositivo en `pathcount`, `prefer3` o en `nofencing`. Si la configuración de protección de un dispositivo de almacenamiento está establecida en `global`, el dispositivo de almacenamiento usa la configuración global. Por ejemplo, si un dispositivo de almacenamiento presenta la configuración predeterminada `pathcount`, la configuración no se modificará si utiliza este procedimiento para cambiar la configuración del protocolo SCSI global a `prefer3`. Complete el procedimiento [“Modificación del protocolo de protección en un solo dispositivo de almacenamiento” en la página 172](#) para modificar la configuración predeterminada de un solo dispositivo.



**Precaución** – Si la protección se desactiva en un momento inadecuado, sus datos pueden ser vulnerables y sufrir daños durante la migración tras error de una aplicación. Si se plantea desactivar la protección, tenga muy en cuenta la posibilidad de sufrir daños en los datos. La protección puede desactivarse si el dispositivo de almacenamiento compartido no admite el protocolo SCSI o si desea permitir el acceso al almacenamiento del clúster desde hosts situados fuera de dicho clúster.

Para cambiar la configuración de protección predeterminada de un dispositivo de quórum, anule la configuración del dispositivo, modifique la configuración de protección y vuelva a configurar el dispositivo de quórum. Si piensa desactivar y activar regularmente la protección de dispositivos entre los que figuran los de quórum, considere la posibilidad de configurar el quórum a través de un servicio de servidor de quórum y así eliminar las interrupciones en el funcionamiento del quórum.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que cuente con autorización RBAC `solaris.cluster.modify`.**
- 2 **Establezca el protocolo de protección para todos los dispositivos de almacenamiento que no sean de quórum.**

```
cluster set -p global_fencing={pathcount | prefer3 | nofencing | nofencing-noscrub}
```

|                                |                                                                                                                                                                                                     |
|--------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>-p global_fencing</code> | Establece el algoritmo de protección predeterminado global para todos los dispositivos compartidos.                                                                                                 |
| <code>prefer3</code>           | Emplea el protocolo SCSI-3 para los dispositivos que cuenten con más de dos rutas.                                                                                                                  |
| <code>pathcount</code>         | Determina el protocolo de protección mediante el número de rutas de DID conectadas al dispositivo compartido. La configuración de <code>pathcount</code> se utiliza con los dispositivos de quórum. |
| <code>nofencing</code>         | Desactiva la protección con la configuración del estado de todos los dispositivos de almacenamiento.                                                                                                |
| <code>nofencing-noscrub</code> | Al limpiar y el dispositivo, se garantiza que éste esté libre de todo tipo de información de reserva de SCSI persistente y se permite que sistemas situados fuera del clúster tengan acceso         |

al almacenamiento. Use la opción `nofencing-noscrub` sólo con dispositivos de almacenamiento con problemas graves con las reservas de SCSI.

#### **Ejemplo 5-41** Establecimiento de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento

En el ejemplo siguiente se establece en el protocolo SCSI-3 el protocolo de protección para todos los dispositivos de almacenamiento del clúster.

```
cluster set -p global_fencing=prefer3
```

### ▼ **Modificación del protocolo de protección en un solo dispositivo de almacenamiento**

El protocolo de protección puede configurarse también para un solo dispositivo de almacenamiento.

---

**Nota** – Para cambiar la configuración de protección predeterminada de un dispositivo de quórum, anule la configuración del dispositivo, modifique la configuración de protección y vuelva a configurar el dispositivo de quórum. Si piensa desactivar y activar regularmente la protección de dispositivos entre los que figuran los de quórum, considere la posibilidad de configurar el quórum a través de un servicio de servidor de quórum y así eliminar las interrupciones en el funcionamiento del quórum.

---

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.



---

**Precaución** – Si la protección se desactiva en un momento inadecuado, sus datos pueden ser vulnerables y sufrir daños durante la migración tras error de una aplicación. Si se plantea desactivar la protección, tenga muy en cuenta la posibilidad de sufrir daños en los datos. La protección puede desactivarse si el dispositivo de almacenamiento compartido no admite el protocolo SCSI o si desea permitir el acceso al almacenamiento del clúster desde hosts situados fuera de dicho clúster.

---

1 Conviértase en superusuario o asuma una función que cuente con autorización RBAC `solaris.cluster.modify`.

2 Establezca el protocolo de protección del dispositivo de almacenamiento.

```
cldevice set -p default_fencing ={pathcount | \
scsi3 | global | nofencing | nofencing-noscrub} device
```

`-p default_fencing` Modifica la propiedad `default_fencing` del dispositivo.

`pathcount` Determina el protocolo de protección mediante el número de rutas de DID conectadas al dispositivo compartido.

`scsi3` Usa el protocolo SCSI-3.

`global` Usa la configuración de protección predeterminada global. La configuración global se utiliza con los dispositivos que no son de quórum.

Desactiva la protección con la configuración del estado de protección para la instancia de DID especificada.

`nofencing-noscrub` Al limpiar y el dispositivo, se garantiza que éste esté libre de todo tipo de información de reserva de SCSI persistente y se permite que sistemas situados fuera del clúster tengan acceso al almacenamiento. Use la opción `nofencing-noscrub` sólo con dispositivos de almacenamiento con problemas graves con las reservas de SCSI.

*dispositivo* Especifica el nombre de la ruta de dispositivo o el nombre del dispositivo.

Si desea obtener más información, consulte la página de comando `man cluster(1CL)`.

### Ejemplo 5-42 Configuración del protocolo de protección de un solo dispositivo

En el ejemplo siguiente se configura el dispositivo `d5`, especificado por su número de dispositivo, con el protocolo SCSI-3.

```
cldevice set -p default_fencing=prefer3 d5
```

En el ejemplo siguiente se desactiva la protección predeterminada del dispositivo `d11`.

```
#cldevice set -p default_fencing=nofencing d11
```

# Administración de sistemas de archivos de clúster

El sistema de archivos de clúster está disponible globalmente. Se puede leer y tener acceso a él desde cualquier nodo del clúster.

TABLA 5-5 Mapa de tareas: administrar sistemas de archivos de clúster

| Tarea                                                                                               | Instrucciones                                                                       |
|-----------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Agregar sistemas de archivos de clúster tras la instalación inicial de Oracle Solaris Cluster       | <a href="#">“Adición de un sistema de archivos de clúster” en la página 174</a>     |
| Eliminar un sistema de archivos de clúster                                                          | <a href="#">“Eliminación de un sistema de archivos de clúster” en la página 178</a> |
| Comprobar los puntos de montaje globales de un clúster para verificar la coherencia entre los nodos | <a href="#">“Comprobación de montajes globales en un clúster” en la página 180</a>  |

## ▼ Adición de un sistema de archivos de clúster

Efectúe esta tarea en cada uno de los sistemas de archivos de clúster creados tras la instalación inicial de Oracle Solaris Cluster.



**Precaución** – Compruebe que haya especificado el nombre del dispositivo de disco correcto. Al crear un sistema de archivos de clúster se destruyen todos los datos de los discos. Si especifica un nombre de dispositivo equivocado, podría borrar datos que no tuviera previsto eliminar.

Antes de agregar un sistema de archivos de clúster adicional, compruebe que se cumplan los requisitos siguientes:

- Se cuenta con privilegios de superusuario en un nodo del clúster.
- Software de administración de volúmenes instalado y configurado en el clúster.
- Grupo de dispositivos (Solaris Volume Manager o VxVM) o segmento de disco de bloques donde poder crear el sistema de archivos de clúster.

Si ha usado Oracle Solaris Cluster Manager para instalar servicios de datos, ya hay uno o más sistemas de archivos de clúster si los discos compartidos en los que crear los sistemas de archivos de clúster eran suficientes.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

1 **Conviértase en superusuario en un nodo de clúster.**

Lleve a cabo este procedimiento desde la zona global si no hay zonas globales configuradas en el clúster.

**Consejo** – Para crear sistemas de archivos con mayor rapidez, conviértase en superusuario en el nodo principal del dispositivo global para el que desea crear un sistema de archivos.

2 **Cree un sistema de archivos.**



**Caution** – Todos los datos de los discos se destruyen al crear un sistema de archivos. Compruebe que haya especificado el nombre del dispositivo de disco correcto. Si se especifica un nombre equivocado, podría borrar datos que no tuviera previsto eliminar.

■ **Para crear un sistema de archivos UFS, utilice el comando `newfs(1M)`**

`phys-schost# newfs raw-disk-device`

La tabla siguiente muestra ejemplos de nombres para el argumento *dispositivo\_disco\_básico*. Cada administrador de volúmenes aplica sus propias convenciones de asignación de nombres.

| Administrador de volúmenes | Nombre de dispositivo de disco de ejemplo | Descripción                                                        |
|----------------------------|-------------------------------------------|--------------------------------------------------------------------|
| Solaris Volume Manager     | <code>/dev/md/nfs/rdisk/d1</code>         | Dispositivo de disco básico d1 dentro del conjunto de discos nfs   |
| Veritas Volume Manager     | <code>/dev/vx/rdisk/oradg/vol01</code>    | Dispositivo de disco básico vol01 dentro del grupo de discos oradg |
| Ninguno                    | <code>/dev/global/rdisk/d1s3</code>       | Dispositivo de disco básico d1s3                                   |

■ **Para crear un sistema de archivos Sistema de archivos de Veritas (VxFS), siga los procedimientos que se indican en la documentación de VxFS.**

3 **Cree un directorio de puntos de montaje en cada nodo del clúster para el sistema de archivos de dicho clúster.**

*Todos los nodos* deben tener un punto de montaje, aunque no se acceda al sistema de archivos de clúster en un nodo concreto.

---

**Consejo** – Para facilitar la administración, cree el punto de montaje en el directorio `/global/grupo_dispositivos/`. Esta ubicación permite distinguir fácilmente los sistemas de archivos de clúster disponibles de forma global de los sistemas de archivos locales.

---

`phys-schost# mkdir -p /global/device-group/mountpoint/`

*grupo\_dispositivos*      Nombre del directorio correspondiente al nombre del grupo de dispositivos que contiene el dispositivo.

*punto\_montaje*          Nombre del directorio en el que se monta el sistema de archivos de clúster.

**4 En cada uno de los nodos del clúster, agregue una entrada en el archivo `/etc/vfstab` para el punto de montaje.**

Consulte la página de comando `man vfstab(4)` para obtener información detallada.

---

**Nota** – Si hay zonas no globales configuradas en el clúster, monte los sistemas de archivos de clúster de la zona global en una ruta del directorio root de la zona global.

---

- a. Especifique en cada entrada las opciones de montaje requeridas para el tipo de sistema de archivos que utilice.
- b. Para montar de forma automática el sistema de archivos de clúster, establezca el campo `mount at boot` en `yes`.
- c. Compruebe que la información de la entrada `/etc/vfstab` de cada sistema de archivos de clúster sea idéntica en todos los nodos.
- d. Compruebe que las entradas del archivo `/etc/vfstab` de cada nodo muestren los dispositivos en el mismo orden.
- e. Compruebe las dependencias de orden de arranque de los sistemas de archivos.

Por ejemplo, fíjese en la situación hipotética siguiente: `phys-schost-1` monta el dispositivo de disco `d0` en `/global/oracle/` y `phys-schost-2` monta el dispositivo de disco `d1` en `/global/oracle/logs/`. Con esta configuración, `phys-schost-2` sólo puede arrancar y montar `/global/oracle/logs/` cuando `phys-schost-1` arranque y monte `/global/oracle/`.

**5 Ejecute la utilidad de comprobación de la configuración en cualquier nodo del clúster.**

`phys-schost# cluster check -k vfstab`



La utilidad de comprobación de la configuración verifica la existencia de los puntos de montaje. Además, comprueba que las entradas del archivo `/etc/vfstab` sean correctas en todos los nodos del clúster. Si no hay ningún error, el comando no devuelve nada.

Para obtener más información, consulte la página de comando `man cluster(1CL)`.

### 6 Monte el sistema de archivos de clúster.

```
phys-schost# mount /global/device-group/mountpoint/
```

- Para sistemas de archivos UFS, monte el sistema de archivos de clúster desde cualquier nodo del clúster.
- Para VxFS, monte el sistema de archivos de clúster desde el elemento maestro del grupo *grupo\_dispositivos*, para asegurarse de que el montaje del sistema de archivos se efectúe correctamente.

Asimismo, desmonte un sistema de archivos de VxFS desde el elemento maestro del grupo *grupo\_dispositivos*, para asegurarse de que el sistema de archivos se desmonte correctamente.

---

**Nota** – Para administrar un sistema de archivos de clúster VxFS en un entorno de Oracle Solaris Cluster, ejecute los comandos de administración sólo desde el nodo principal en que está montado el sistema de archivos de clúster VxFS.

---

### 7 Compruebe que el sistema de archivos de clúster esté montado en todos los nodos de dicho clúster.

Puede utilizar los comandos `df` o `mount` para enumerar los sistemas de archivos montados. Para obtener más información, consulte las páginas de comando `man df(1M)` o `mount(1M)`.

Se puede obtener acceso a los sistemas de archivos del clúster desde la zona global y desde la zona no global.

## Ejemplo 5–43 Creación de un sistema de archivos de clúster UFS

En el ejemplo siguiente, se crea un sistema de archivos de clúster UFS en el volumen de Solaris Volume Manager `/dev/md/oracle/rdisk/d1`. Se agrega una entrada para el sistema de archivos de clúster en el archivo `vfstab` de cada nodo. A continuación, se ejecuta el comando `cluster check` desde un nodo. Tras comprobar que configuración se haya efectuado correctamente, se monta el sistema de archivos de clúster desde un nodo y se verifica en todos los nodos.

```
phys-schost# newfs /dev/md/oracle/rdisk/d1
...
phys-schost# mkdir -p /global/oracle/d1
phys-schost# vi /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
...
phys-schost# cluster check -k vfstab
phys-schost# mount /global/oracle/d1
phys-schost# mount
...
/global/oracle/d1 on /dev/md/oracle/dsk/d1 read/write/setuid/global/logging/largefiles
on Sun Oct 3 08:56:16 2005
```

## ▼ Eliminación de un sistema de archivos de clúster

Para *quitar* un sistema de archivos de clúster, no tiene más que desmontarlo. Para quitar o eliminar también los datos, quite el dispositivo de disco subyacente (o metadispositivo o metavolumen) del sistema.

---

**Nota** – Los sistemas de archivos de clúster se desmontan automáticamente como parte del cierre del sistema sucedido al ejecutar `cluster shutdown` para detener todo el clúster. Un sistema de archivos de clúster no se desmonta al ejecutar `shutdown` para detener un solo nodo. Sin embargo, si el nodo que se cierra es el único que tiene una conexión con el disco, cualquier intento de tener acceso al sistema de archivos de clúster en ese disco da como resultado un error.

---

Compruebe que se cumplan los requisitos siguientes antes de desmontar los sistemas de archivos de clúster:

- Se cuenta con privilegios de superusuario en un nodo del clúster.
- Sistema de archivos no ocupado. Un sistema de archivos está ocupado si un usuario está trabajando en un directorio del sistema de archivos o si un programa tiene un archivo abierto en dicho sistema de archivos. El usuario o el programa pueden estar trabajando en cualquier nodo del clúster.

### 1 Conviértase en superusuario en un nodo de clúster.

### 2 Determine los sistemas de archivos de clúster que están montados.

```
mount -v
```

### 3 En cada nodo, enumere todos los procesos que utilicen el sistema de archivos de clúster para saber los procesos que va a detener.

```
fuser -c [-u] mountpoint
```

-c                      Informa sobre los archivos que son puntos de montaje de sistemas de archivos y sobre cualquier archivo dentro de sistemas de archivos montados.

- u (Opcional) Muestra el nombre de inicio de sesión del usuario para cada ID de proceso.
- punto\_montaje* Especifica el nombre del sistema de archivos del clúster para el que desea detener procesos.

#### 4 Detenga todos los procesos del sistema de archivos de clúster en cada uno de los nodos.

Use el método que prefiera para detener los procesos. Si conviene, utilice el comando siguiente para forzar la conclusión de los procesos asociados con el sistema de archivos de clúster.

```
fuser -c -k mountpoint
```

Se envía un SIGKILL a cada proceso que usa el sistema de archivos de clúster.

#### 5 Compruebe en todos los nodos que no haya ningún proceso que esté usando el sistema de archivos.

```
fuser -c mountpoint
```

#### 6 Desmonte el sistema de archivos desde un solo nodo.

```
umount mountpoint
```

*punto\_montaje* Especifica el nombre del sistema de archivos del clúster que desea desmontar. Puede ser el nombre del directorio en el que está montado el sistema de archivos de clúster o la ruta del nombre del dispositivo del sistema de archivos.

#### 7 (Opcional) Edite el archivo `/etc/vfstab` para eliminar la entrada del sistema de archivos de clúster que se va a eliminar.

Realice este paso en cada nodo del clúster que tenga una entrada de este sistema de archivos de clúster en el archivo `/etc/vfstab`.

#### 8 (Opcional) Quite el dispositivo de disco `group/metadevice/volume/plex`.

Para obtener más información, consulte la documentación del administrador de volúmenes.

### Ejemplo 5–44 Eliminación de un sistema de archivos de clúster

En el ejemplo siguiente se quita un sistema de archivos de clúster UFS montado en el metadispositivo o el volumen de Solaris Volume Manager `/dev/md/oracle/rdisk/d1`.

```
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

(On each node, remove the highlighted entry:)
vi /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

[Save and exit.]

```

Para borrar los datos del sistema de archivos de clúster, elimine el dispositivo subyacente. Para obtener más información, consulte la documentación del administrador de volúmenes.

## ▼ Comprobación de montajes globales en un clúster

La utilidad `cluster(1CL)` comprueba la sintaxis de las entradas de sistemas de archivos de clúster en el archivo `/etc/vfstab`. Si no hay ningún error, el comando no devuelve nada.

---

**Nota** – Ejecute el comando `cluster check` tras realizar modificaciones en la configuración, por ejemplo (como quitar un sistema de archivos de clúster, que puedan haber afectado a los dispositivos o a los componentes de administración de volúmenes).

---

- 1 Conviértase en superusuario en un nodo de clúster.
- 2 Compruebe los montajes globales del clúster.

```
cluster check -k vfstab
```

## Administración de la supervisión de rutas de disco

Los comandos de administración de supervisión de rutas de disco permiten recibir notificaciones sobre errores en rutas de disco secundarias. Siga los procedimientos descritos en esta sección para realizar tareas administrativas asociadas con la supervisión de las rutas de disco. Consulte el [Capítulo 3, “Key Concepts for System Administrators and Application Developers” de Oracle Solaris Cluster Concepts Guide](#) si desea obtener información conceptual sobre el daemon de supervisión de las rutas de disco. Consulte la página de comando `man cldevice(1CL)` si desea ver una descripción de las opciones del comando y otros comandos relacionados. Si desea obtener más información sobre cómo ajustar el daemon `scdpmd`, consulte la página de comando `man scdpmd.conf(4)`. Consulte también la página de comando `man syslogd(1M)` para ver los errores registrados sobre los que informa el daemon.

**Nota** – Las rutas de disco se agregan automáticamente a la lista de supervisión al incorporar dispositivos de E/S a un nodo mediante el comando `cldevice`. Asimismo, la supervisión de rutas de disco se detiene automáticamente al quitar los dispositivos de un nodo con los comandos de Oracle Solaris Cluster.

TABLA 5–6 Mapa de tareas: administrar la supervisión de rutas de disco

| Tarea                                                                                                                                                                                                                                                        | Instrucciones                                                                                                                                                                                                                                                            |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Supervisar una ruta de disco                                                                                                                                                                                                                                 | “Supervisión de una ruta de disco” en la página 181                                                                                                                                                                                                                      |
| Anular la supervisión de una ruta de disco                                                                                                                                                                                                                   | “Anulación de la supervisión de una ruta de disco” en la página 183                                                                                                                                                                                                      |
| Imprimir el estado de rutas de disco erróneas correspondientes a un nodo                                                                                                                                                                                     | “Impresión de rutas de disco erróneas” en la página 184                                                                                                                                                                                                                  |
| Supervisar rutas de disco desde un archivo                                                                                                                                                                                                                   | “Supervisión de rutas de disco desde un archivo” en la página 185                                                                                                                                                                                                        |
| Habilitar o inhabilitar el rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas                                                                                                                                       | “Habilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 187<br><br>“Inhabilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 188 |
| Resolver un estado incorrecto de una ruta de disco. Puede aparecer un informe sobre un estado de ruta de disco incorrecta cuando el dispositivo de DID supervisado no está disponible al arrancar y la instancia de DID no se carga en el controlador de DID | “Resolución de un error de estado de ruta de disco” en la página 184                                                                                                                                                                                                     |

Entre los procedimientos de la sección siguiente que ejecutan el comando `cldevice` está el argumento de ruta de disco. El argumento de ruta de disco consta del nombre de un nodo y el de un disco. El nombre del nodo no es imprescindible y está predefinido para convertirse en `all` si no se especifica.

## ▼ Supervisión de una ruta de disco

Efectúe esta tarea para supervisar las rutas de disco del clúster.



**Precaución** – Los nodos que ejecuten versiones anteriores a Sun Cluster 3.1 10/03 no admiten la supervisión de rutas de disco. No utilice los comandos de supervisión de rutas de disco en el transcurso de una actualización. Tras actualizarse todos los nodos, deben estar en línea para emplear los comandos de supervisión de rutas de disco.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**
- 2 **Supervise una ruta de disco.**  
`# cldevice monitor -n node disk`
- 3 **Compruebe que se supervise la ruta de disco.**  
`# cldevice status device`

**Ejemplo 5–45** Supervisión de una ruta de disco en un solo nodo

En el ejemplo siguiente se supervisa la ruta de disco `schost-1:/dev/did/rdisk/d1` desde un solo nodo. El daemon de DPM del nodo `schost-1` es el único que supervisa la ruta al disco `/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     |

**Ejemplo 5–46** Supervisión de una ruta de disco en todos los nodos

En el ejemplo siguiente se supervisa la ruta de disco `schost-1:/dev/did/dsk/d1` desde todos los nodos. La supervisión de rutas de disco se inicia en todos los nodos en los que `/dev/did/dsk/d1` sea una ruta válida.

```
cldevice monitor /dev/did/dsk/d1
cldevice status /dev/did/dsk/d1
```

| Device Instance   | Node          | Status |
|-------------------|---------------|--------|
| -----             |               |        |
| /dev/did/rdisk/d1 | phys-schost-1 | Ok     |

Ejemplo 5-47    Relectura de la configuración de disco desde CCR

En el ejemplo siguiente se fuerza al daemon a que relea la configuración de disco desde CCR e imprima las rutas de disco supervisadas con su estado.

```
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
```

▼    **Anulación de la supervisión de una ruta de disco**

Siga este procedimiento para anular la supervisión de una ruta de disco.



**Precaución** – Los nodos que ejecuten versiones anteriores a Sun Cluster 3.1 10/03 no admiten la supervisión de rutas de disco. No utilice los comandos de supervisión de rutas de disco en el transcurso de una actualización. Tras actualizarse todos los nodos, deben estar en línea para emplear los comandos de supervisión de rutas de disco.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1    **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC solaris.cluster.modify en todos los nodos del clúster.**
- 2    **Determine el estado de la ruta de disco cuya supervisión se vaya a anular.**  
      # cldevice status device
- 3    **Anule la supervisión de las correspondientes rutas de disco en cada uno de los nodos.**  
      # cldevice unmonitor -n node disk

### Ejemplo 5-48 Anulación de la supervisión de una ruta de disco

En el ejemplo siguiente se anula la supervisión de la ruta de disco `schost-2:/dev/did/rdisk/d1` y se imprimen las rutas de disco de todo el clúster, con sus estados.

```
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 |

## ▼ Impresión de rutas de disco erróneas

Siga este procedimiento para imprimir las rutas de disco erróneas correspondientes a un clúster.



**Precaución** – Los nodos que ejecuten versiones anteriores a Sun Cluster 3.1 10/03 no admiten la supervisión de rutas de disco. No utilice los comandos de supervisión de rutas de disco en el transcurso de una actualización. Tras actualizarse todos los nodos, deben estar en línea para emplear los comandos de supervisión de rutas de disco.

- 1 **Conviértase en superusuario en un nodo de clúster.**
- 2 **Imprima las rutas de disco erróneas de todo el clúster.**

```
cldevice status -s fail
```

### Ejemplo 5-49 Impresión de rutas de disco erróneas

En el ejemplo siguiente se imprimen las rutas de disco erróneas de todo el clúster.

```
cldevice status -s fail
```

| Device Instance | Node          | Status |
|-----------------|---------------|--------|
| -----           | ----          | -----  |
| dev/did/dsk/d4  | phys-schost-1 | fail   |

## ▼ Resolución de un error de estado de ruta de disco

Si ocurre alguno de los siguientes eventos, es posible que la supervisión de rutas de disco no pueda actualizar el estado de una ruta con errores al volver a conectarse en línea:

- Un error de la ruta supervisada hace que se re arranque un nodo.
- El dispositivo que está en la ruta de DID supervisada sólo vuelve a ponerse en línea después de que también lo haya hecho el nodo re arrancado.



Se informa sobre un estado de ruta de disco incorrecta porque el dispositivo de DID supervisado no está disponible al arrancar y la instancia de DID no se carga en el controlador de DID. Cuando suceda esto, actualice manualmente la información de DID.

- 1 Desde un nodo, actualice el espacio de nombre de dispositivos globales.  
`# cldevice populate`
- 2 Compruebe en todos los nodos que haya finalizado el procesamiento del comando antes de continuar con el paso siguiente.  
El comando se ejecuta de forma remota en todos los nodos, incluso al ejecutarse en un solo nodo. Para determinar si ha concluido el procesamiento del comando, ejecute el comando siguiente en todos los nodos del clúster.  
`# ps -ef | grep cldevice populate`
- 3 Compruebe que, dentro del intervalo de tiempo de consulta de DPM, la ruta de disco errónea tenga ahora un estado correcto.

`# cldevice status disk-device`

| Device Instance | Node          | Status |
|-----------------|---------------|--------|
| -----           | ----          | -----  |
| dev/did/dsk/dN  | phys-schost-1 | Ok     |

## ▼ Supervisión de rutas de disco desde un archivo

Complete el procedimiento descrito a continuación para supervisar o anular la supervisión de rutas de disco desde un archivo.

Para modificar la configuración del clúster mediante un archivo, primero se debe exportar la configuración actual. Esta operación de exportación crea un archivo XML que luego puede modificarse para establecer los elementos de la configuración que vaya a cambiar. Las instrucciones de este procedimiento describen el proceso completo.



**Precaución** – Los nodos que ejecuten versiones anteriores a Sun Cluster 3.1 10/03 no admiten la supervisión de rutas de disco. No utilice los comandos de supervisión de rutas de disco en el transcurso de una actualización. Tras actualizarse todos los nodos, deben estar en línea para emplear los comandos de supervisión de rutas de disco.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC solaris.cluster.modify en todos los nodos del clúster.**
- 2 **Exporte la configuración del dispositivo a un archivo XML.**  

```
cldevice export -o configurationfile
```

-o *archivo\_configuración* Especifique el nombre de archivo del archivo XML.
- 3 **Modifique el archivo de configuración de manera que se supervisen las rutas de dispositivos.**  
Busque las rutas de dispositivos que desee supervisar y establezca el atributo `monitored` en `true`.
- 4 **Supervise las rutas de dispositivos.**  

```
cldevice monitor -i configurationfile
```

-i *archivo\_configuración* Especifique el nombre de archivo del archivo XML modificado.
- 5 **Compruebe que la ruta de dispositivo ya se supervise.**  

```
cldevice status
```

### **Ejemplo 5–50** Supervisar rutas de disco desde un archivo

En el ejemplo siguiente, la ruta de dispositivo entre el nodo `phys-schost-2` y el dispositivo `d3` se supervisa con archivo XML.

El primer paso es exportar la configuración del clúster actual.

```
cldevice export -o deviceconfig
```

El archivo XML `deviceconfig` muestra que la ruta entre `phys-schost-2` y `d3` actualmente no se supervisa.

```
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
.
.
.
 <deviceList readonly="true">
 <device name="d3" ctd="clt8d0">
 <devicePath nodeRef="phys-schost-1" monitored="true"/>
 <devicePath nodeRef="phys-schost-2" monitored="false"/>
 </device>
 </deviceList>
</cluster>
```

Para supervisar esa ruta, establezca el atributo `monitored` en `true`, tal y como se explica a continuación.

```
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
.
.
.
 <deviceList readonly="true">
 <device name="d3" ctd="clt8d0">
 <devicePath nodeRef="phys-schost-1" monitored="true"/>
 <devicePath nodeRef="phys-schost-2" monitored="true"/>
 </device>
 </deviceList>
</cluster>
```

Use el comando `cldevice` para leer el archivo y activar la supervisión.

```
cldevice monitor -i deviceconfig
```

Use el comando `cldevice` para comprobar que el archivo ya se esté supervisando.

```
cldevice status
```

**Véase también** Si desea obtener más información sobre cómo exportar la configuración del clúster y usar el archivo XML resultante para establecer la configuración del clúster, consulte las páginas de comando [man cluster\(1CL\)](#) y [clconfiguration\(5CL\)](#).

## ▼ **Habilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas**

Al habilitar esta función, un nodo se rearranca automáticamente siempre que se cumplan las condiciones siguientes:

- Fallan todas las rutas de disco compartido supervisadas del nodo.
- Se puede acceder a uno de los discos como mínimo desde un nodo diferente del clúster.

Al rearrancar el nodo se rearrancan todos los grupos de recursos y grupos de dispositivos que se controlan en ese nodo en otro nodo.

Si no se tiene acceso a todas las rutas de disco compartido supervisadas de un nodo tras rearrancar automáticamente el nodo, éste no rearranca automáticamente otra vez. Sin embargo, si alguna de las rutas de disco está disponible tras rearrancar el nodo pero falla posteriormente, el nodo rearranca automáticamente otra vez.

Al habilitar la propiedad `reboot_on_path_failure`, los estados de las rutas de disco local no se tienen en cuenta al determinar si es necesario rearrancar un nodo. Esto afecta sólo a discos compartidos supervisados.

- 1 **Conviértase en superusuario en uno de los nodos del clúster o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Habilite el re arranque automático de un nodo para *todos* los nodos del clúster si fallan todas las rutas de disco compartido supervisadas de dicho nodo.**

```
clnode set -p reboot_on_path_failure=enabled +
```

## ▼ **Inhabilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas**

Si se inhabilita esta función y fallan todas las rutas de disco compartido supervisadas de un nodo, dicho nodo *no* re arranca automáticamente.

- 1 **Conviértase en superusuario en uno de los nodos del clúster o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Inhabilite el re arranque automático de un nodo para *todos* los nodos del clúster si fallan todas las rutas de disco compartido supervisadas de dicho nodo.**

```
clnode set -p reboot_on_path_failure=disabled +
```

# Administración de quórum

---

Este capítulo presenta los procedimientos para administrar dispositivos de quórum dentro de los servidores de quórum de Oracle Solaris Cluster y Oracle Solaris Cluster. Si desea obtener información sobre conceptos de quórum, consulte *“Quorum and Quorum Devices” de Oracle Solaris Cluster Concepts Guide*.

- [“Administración de dispositivos de quórum” en la página 189](#)
- [“Administración de servidores de quórum de Oracle Solaris Cluster” en la página 215](#)

## Administración de dispositivos de quórum

Un dispositivo de quórum es un dispositivo de almacenamiento o servidor de quórum compartido por dos o más nodos y que aporta votos usados para establecer un quórum. Esta sección explica los procedimientos para administrar dispositivos de quórum.

Puede usar el comando `clquorum(1CL)` para realizar todos los procedimientos administrativos de los dispositivos de quórum. También puede efectuar algunos procedimientos mediante el uso de la utilidad interactiva `clsetup(1CL)` o la GUI de Oracle Solaris Cluster Manager. Siempre que es posible, los procedimientos de quórum de esta sección se describen con la utilidad `clsetup`. La ayuda en línea de Oracle Solaris Cluster Manager describe cómo realizar los procedimientos relativos al quórum mediante la GUI. Al trabajar con dispositivos de quórum, tenga en cuenta las directrices siguientes:

- Todos los comandos relacionados con el quórum deben ejecutarse en el nodo de votación del clúster global.
- Si el comando `clquorum` se ve interrumpido o falla, la información de la configuración de quórum podría volverse incoherente en la base de datos de configuración del clúster. De darse esta incoherencia, vuelva a ejecutar el comando o ejecute el comando `clquorum reset` para restablecer la configuración de quórum.

- Para que el clúster tenga la máxima disponibilidad, compruebe que el número total de votos aportado por los dispositivos de quórum sea menor que el aportado por los nodos. De lo contrario, los nodos no pueden formar un clúster ninguno de los dispositivos de quórum está disponible aunque todos los nodos estén funcionando.
- No agregue ningún disco configurado como dispositivo de quórum a una agrupación de almacenamiento ZFS de Solaris. Si se agrega un dispositivo de quórum configurado a una agrupación de almacenamiento ZFS de Solaris, el disco se vuelve a etiquetar como disco EFI, se pierde la información de configuración de quórum y el disco ya no proporciona un voto de quórum al clúster. Una vez que un disco esté en una agrupación de almacenamiento, ya se puede configurar como dispositivo de quórum. También se puede anular la configuración del disco, agregarlo a la agrupación de almacenamiento y luego volverlo a configurar como dispositivo de quórum.

**Nota** – El comando `clsetup` es una interfaz interactiva para los demás comandos de Oracle Solaris Cluster. Cuando se ejecuta `clsetup`, el comando genera los pertinentes comandos, en este caso se trata de comandos `clquorum`. Los comandos generados se muestran en los ejemplos que figuran al final de los procedimientos.

Para ver la configuración de quórum, use `clquorum show`. El comando `clquorum list` muestra los nombres de los dispositivos de quórum del clúster. El comando `clquorum status` ofrece información sobre el estado y el número de votos.

La mayoría de los ejemplos de esta sección proceden de un clúster de tres nodos.

**TABLA 6-1** Lista de tareas: administrar el quórum

Tarea	Instrucciones
Agregar un dispositivo de quórum a un clúster mediante <code>clsetup(1CL)</code>	<a href="#">“Adición de un dispositivo de quórum” en la página 192</a>
Eliminar un dispositivo de quórum de un clúster mediante <code>clsetup</code> (para generar <code>clquorum</code> )	<a href="#">“Eliminación de un dispositivo de quórum” en la página 204</a>
Eliminar el último dispositivo de quórum de un clúster mediante <code>clsetup</code> (para generar <code>clquorum</code> )	<a href="#">“Eliminación del último dispositivo de quórum de un clúster” en la página 205</a>
Reemplazar un dispositivo de quórum de un clúster mediante los procedimientos de agregar y quitar	<a href="#">“Sustitución de un dispositivo de quórum” en la página 206</a>
Modificar una lista de dispositivos de quórum mediante los procedimientos de agregar y quitar	<a href="#">“Modificación de una lista de nodos de dispositivo de quórum” en la página 207</a>

TABLA 6-1 Lista de tareas: administrar el quórum (Continuación)

Tarea	Instrucciones
Poner un dispositivo de quórum en estado de mantenimiento mediante <code>clsetup</code> (para generar <code>clquorum</code> )  (Mientras se encuentra en estado de mantenimiento, el dispositivo de quórum no participa en las votaciones para establecer el quórum.)	<a href="#">“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 210</a>
Restablecer la configuración de quórum a su estado predeterminado mediante <code>clsetup</code> (para generar <code>clquorum</code> )	<a href="#">“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 211</a>
Enumerar en una lista los dispositivos de quórum y los números de votos mediante el comando <code>clquorum(1CL)</code>	<a href="#">“Enumeración de una lista con la configuración de quórum” en la página 213</a>

## Reconfiguración dinámica con dispositivos de quórum

Debe tener en cuenta diversos aspectos al desarrollar operaciones de reconfiguración dinámica o DR (Dynamic Reconfiguration) en los dispositivos de quórum de un clúster.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la reconfiguración dinámica de Oracle Solaris también son válidos para la reconfiguración dinámica en Oracle Solaris Cluster excepto la inactividad del sistema operativo. Por lo tanto, consulte la documentación sobre la reconfiguración dinámica de Oracle Solaris *antes* de utilizarla con el software Oracle Solaris Cluster. Preste especial atención a los problemas que afecten a los dispositivos de E/S no conectados a red durante las operaciones de desconexión de reconfiguración dinámica.
- Oracle Solaris Cluster rechaza operaciones de eliminación de tarjetas de DR efectuadas cuando hay una interfaz configurada para un dispositivo de quórum.
- Si la operación de DR pertenece a un dispositivo activo, Oracle Solaris Cluster rechaza la operación e identifica a los dispositivos que se verían afectados por ella.

Para eliminar un dispositivo de quórum, complete los pasos siguientes en el orden que se indica.

TABLA 6-2 Mapa de tareas: reconfiguración dinámica con dispositivos de quórum

Tarea	Instrucciones
1. Habilitar un nuevo dispositivo de quórum para sustituir el que se va a eliminar	<a href="#">“Adición de un dispositivo de quórum” en la página 192</a>

TABLA 6-2 Mapa de tareas: reconfiguración dinámica con dispositivos de quórum (Continuación)

Tarea	Instrucciones
2. Inhabilitar el dispositivo de quórum que se va a eliminar	<a href="#">“Eliminación de un dispositivo de quórum” en la página 204</a>
3. Efectuar la operación de eliminación de reconfiguración dinámica en el dispositivo que se va a eliminar	<i>Sun Enterprise 10000 Dynamic Reconfiguration Reference Manual</i> (de la recopilación <i>Solaris 10 on Sun Hardware</i> ).

## Adición de un dispositivo de quórum

Esta sección presenta procedimientos para agregar un dispositivo de quórum. Compruebe que todos los nodos del clúster estén en línea antes de agregar un nuevo dispositivo de quórum. Si desea obtener información sobre cómo determinar el número de números de votos de quórum necesario para el clúster, configuraciones de quórum recomendadas y protección de errores, consulte [“Quorum and Quorum Devices” de Oracle Solaris Cluster Concepts Guide](#).



**Precaución** – No agregue ningún disco configurado como dispositivo de quórum a una agrupación de almacenamiento ZFS de Solaris. Si se agrega un dispositivo de quórum configurado a una agrupación de almacenamiento ZFS de Solaris, el disco se vuelve a etiquetar como disco EFI, se pierde la información de configuración de quórum y el disco ya no proporciona un voto de quórum al clúster. Una vez que un disco esté en una agrupación de almacenamiento, ya se puede configurar como dispositivo de quórum. También se puede anular la configuración del disco, agregarlo a la agrupación de almacenamiento y luego volverlo a configurar como dispositivo de quórum.

El software Oracle Solaris Cluster admite los tipos de dispositivos de quórum siguientes:

- Disco compartido con conexión directa- para dispositivos con tecnología SCSI o Serial Attached Technology Attachment (SATA)
- NAS de Sun
- Sun Storage 7000 Unified Storage Systems de Oracle
- NAS de Network Appliance (NetApp)
- Oracle Solaris Cluster Quorum Server

En las secciones siguientes se presentan procedimientos para agregar estos dispositivos:

- [“Adición de un dispositivo de quórum de disco compartido” en la página 193](#)
- [“Adición de un dispositivo de quórum de almacenamiento conectado a NAS de Network Appliance” en la página 197](#)
- [“Adición de un dispositivo de quórum de servidor de quórum” en la página 200](#)



---

**Nota** – Los discos replicados no se pueden configurar como dispositivos de quórum. Si se intenta agregar un disco replicado como dispositivo de quórum, se recibe el mensaje de error siguiente, el comando detiene su ejecución y genera un código de error.

*Disk-name is a replicated device. Replicated devices cannot be configured as quorum devices.*

---

Un dispositivo de quórum de disco compartido es cualquier dispositivo de almacenamiento conectado que sea compatible con el software de Oracle Solaris Cluster. El disco compartido se conecta a dos o más nodos del clúster. Si se activa la protección, un disco con doble puerto puede configurarse como dispositivo de quórum que utilice SCSI-2 o SCSI-3 (la opción predeterminada es SCSI-2). Si se activa la protección y el dispositivo compartido está conectado a más de dos nodos, el disco compartido puede configurarse como dispositivo de quórum que use el protocolo SCSI-3 (es el predeterminado si hay más de dos nodos). Puede emplear el identificador de anulación de SCSI para que el software de Oracle Solaris Cluster deba usar el protocolo SCSI-3 con los discos compartidos de doble puerto.

Si desactiva la protección en un disco compartido, a continuación puede configurarlo como dispositivo de quórum que use el protocolo de quórum de software. Esto sería cierto al margen de que el disco fuese compatible con los protocolos SCSI-2 o SCSI-3. El quórum del software es un protocolo de Oracle que emula un formato de Reservas de grupo persistente (PGR) SCSI.




---

**Precaución** – Si utiliza discos que no son compatibles con SCSI (como SATA), debe desactivarse la protección de SCSI.

---

Para dispositivos de quórum, puede usar un disco que contenga datos de usuario o que sea miembro de un grupo de dispositivos. El protocolo que utiliza el subsistema de quórum con un disco compartido puede verse si mira el valor de `access-mode` para el disco compartido en la salida del comando `cluster show`.

Estos procedimientos también pueden realizarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Consulte las páginas de comando `man clsetup(1CL)` y `clquorum(1CL)` si desea obtener información sobre los comandos que se usan en los procedimientos siguientes.

## ▼ Adición de un dispositivo de quórum de disco compartido

El software de Oracle Solaris Cluster admite los dispositivos de disco compartido (SCSI y SATA) como dispositivos de quórum. Un dispositivo de SATA no es compatible con una reserva de SCSI; para configurar estos discos como dispositivos de quórum, debe inhabilitar el indicador de protección de la reserva de SCSI y utilizar el protocolo de quórum de software.

Para completar este procedimiento, identifique una unidad de disco por su ID de dispositivo (DID), que comparten los nodos. Use el comando `cldevice show` para ver la lista de nombres de DID. Consulte la página de comando `man cldevice(1CL)` si desea obtener información adicional. Compruebe que todos los nodos del clúster estén en línea antes de agregar un nuevo dispositivo de quórum.

Use este procedimiento para configurar dispositivos SATA o SCSI

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 Inicie la utilidad `clsetup`.**

`# clsetup`

Aparece el menú principal de `clsetup`.

- 3 Escriba el número correspondiente a la opción de quórum.**

Aparece el menú Quórum.

- 4 Escriba el número correspondiente a la opción de agregar un dispositivo de quórum; a continuación escriba `yes` cuando la utilidad `clsetup` solicite que confirme el dispositivo de quórum que va a agregar.**

La utilidad `clsetup` pregunta qué tipo de dispositivo de quórum se desea agregar.

- 5 Escriba el número correspondiente a la opción de un dispositivo de quórum de disco compartido.**

La utilidad `clsetup` pregunta qué dispositivo global quiere utilizar.

- 6 Escriba el dispositivo global que va a usar.**

La utilidad `clsetup` solicita que confirme que el nuevo dispositivo de quórum debe agregarse al dispositivo global especificado.

- 7 Escriba `yes` para seguir agregando el nuevo dispositivo de quórum.**

Si se agrega correctamente el nuevo dispositivo de quórum, la utilidad `clsetup` muestra el correspondiente mensaje.

**8 Compruebe que se haya agregado el dispositivo de quórum.**

```
clquorum list -v
```

**Ejemplo 6-1 Adición de un dispositivo de quórum de disco compartido**

En el ejemplo siguiente se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum de disco compartido y el paso de comprobación.

```
Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on any
cluster node.
```

```
[Start the clsetup utility:]
```

```
clsetup
```

```
[Select Quorum>Add a quorum device]
```

```
[Answer the questions when prompted.]
```

```
[You will need the following information.]
```

```
[Information: Example:]
[Directly attached shared disk shared_disk]
[Global device d20]
```

```
[Verify that the clquorum command was completed successfully:]
```

```
clquorum add d20
```

```
Command completed successfully.
```

```
[Quit the clsetup Quorum Menu and Main Menu.]
```

```
[Verify that the quorum device is added:]
```

```
clquorum list -v
```

```
Quorum Type

d20 shared_disk
scphyshost-1 node
scphyshost-2 node
```

## ▼ **Cómo agregar un dispositivo de quórum NAS de Sun NAS o de Sun Storage 7000 Unified Storage Systems**

Compruebe que todos los nodos del clúster estén en línea antes de agregar un nuevo dispositivo de quórum.

`phys -schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Use la GUI de NAS de Sun o la GUI de Sun Unified Storage para configurar un dispositivo iSCSI en el archivador de NAS de Sun.

Si tiene un dispositivo de Sun Unified Storage, utilice la GUI para realizar los pasos siguientes.  
Si tiene un dispositivo NAS de Sun, utilice los comandos siguientes:

- a. Cree un volumen de archivos con un tamaño aproximado de 50 MB.
- b. En cada uno de los nodos, cree una lista de acceso de iSCSI.
  - i. Emplee el nombre del clúster como nombre de la lista de acceso de iSCSI.
  - ii. Agregue a la lista de acceso el nombre del nodo iniciador de cada nodo del clúster. No son necesarios CHAP ni IQN.

- c. Configure el LUN de iSCSI.

Puede usar el nombre del volumen de archivos de apoyo como nombre para el LUN.  
Agregue al LUN la lista de acceso de cada nodo.

- 2 Detecte el LUN de iSCSI en cada uno de los nodos y establezca la lista de acceso de iSCSI con configuración estática.

```
iscsiadm modify discovery -s enable

iscsiadm list discovery
Discovery:
 Static: enabled
 Send Targets: disabled
 iSNS: disabled

iscsiadm add static-config iqn.LUNName,IPAddress_of_NASDevice
devfsadm -i iscsi
cldevice refresh
```

- 3 Desde un nodo del clúster, configure los DID correspondientes al LUN de iSCSI.

```
/usr/cluster/bin/cldevice populate
```

- 4 Identifique el dispositivo DID que representa el LUN del dispositivo NAS que ya se ha configurado en el clúster mediante iSCSI. Use el comando `cldevice show` para ver la lista de nombres de DID. Consulte la página de comando `man cldevice(1CL)` si desea obtener información adicional.

- 5 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.

- 6 Utilice el comando `clquorum` para agregar el dispositivo NAS como dispositivo del quórum mediante el dispositivo DID identificado en el [Paso 4](#).

```
clquorum add d20
```

El clúster tiene reglas predeterminadas para decidir si se deben usar scsi-2, scsi-3 o los protocolos de quórum del software. Consulte la [clquorum\(1CL\)](#) para más información.

### **Ejemplo 6–2 Adición de un dispositivo de quórum NAS de Sun NAS o de Sun Storage 7000 Unified Storage Systems**

En el ejemplo siguiente se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum NAS y un paso de comprobación.

```
Add an iSCSI device on the Sun NAS filer.
Use the Sun NAS GUI to create a file volume that is approximately 50mb in size.
File Volume Operations -> Create File Volume
For each node, create an iSCSI access list.
iSCSI Configuration -> Configure Access List
Add the initiator node name of each cluster node to the access list.
*** Need GUI or command syntax for this step. ***
Configure the iSCSI LUN
iSCSI Configuration -> Configure iSCSI LUN
On each of the cluster nodes, discover the iSCSI LUN and set the iSCSI access list to static configuration.
iscsiadm modify discovery -s enable
iscsiadm list discovery
Discovery:
 Static: enabled
 Send Targets: enabled
 iSNS: disabled
iscsiadm add static-config
iqn.1986-03.com.sun0-1:000e0c66efe8.4604DE16.thinquorum,10.11.160.20
devsadm -i iscsi
From one cluster node, configure the DID devices for the iSCSI LUN.
/usr/cluster/bin/scldevice populate
/usr/cluster/bin/scldevice populate
Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on any cluster node.

[Add the NAS device as a quorum device
using the DID device:]
clquorum add d20

Command completed successfully.
```

### **▼ Adición de un dispositivo de quórum de almacenamiento conectado a NAS de Network Appliance**

Al usar un dispositivo de almacenamiento conectado a red (Network-Attached Storage, NAS) de Network Appliance (NetApp) como dispositivo de quórum, deben cumplirse los requisitos siguientes:

- Debe instalar la licencia de iSCSI de NetApp.
- Debe configurar un LUN de iSCSI LUN en el archivador integrado en el clúster para utilizarlo como dispositivo de quórum.
- Debe configurar la unidad NAS de NetApp de forma que use NTP para sincronizar la hora.

- Al menos uno de los servidores NTP seleccionados para el archivador integrado en el clúster debe ser un servidor NTP para los nodos de Oracle Solaris Cluster.
- Al arrancar el clúster, arranque siempre el dispositivo de NAS antes de iniciar los nodos del clúster.

Si arranca los dispositivos en un orden incorrecto, los nodos no pueden detectar el dispositivo de quórum. Si en esta situación se diera un error en un nodo, es posible que el clúster deje de funcionar. Si se interrumpe el servicio, debe rearrancar todo el clúster entero, o quitar el dispositivo de quórum NAS de NetApp y volverlo a agregar.

- Un clúster puede usar un dispositivo de NAS sólo para un dispositivo de quórum.  
Si necesita más dispositivos de quórum, puede configurar otro almacenamiento compartido. Los clústers adicionales que usen el mismo dispositivo de NAS pueden emplear LUN separados en ese dispositivo como dispositivos de quórum.

Consulte la documentación siguiente de Oracle Solaris Cluster para obtener información sobre cómo instalar un dispositivo de almacenamiento NAS de NetApp en un entorno de Oracle Solaris Cluster: [Oracle Solaris Cluster 3.3 With Network-Attached Storage Devices Manual](#).

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que todos los nodos de Oracle Solaris Cluster estén en línea y se puedan comunicar con el archivador integrado en el clúster de NetApp.**
- 2 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 3 Inicie la utilidad `clsetup`.**

```
clsetup
```

Aparece el menú principal de `clsetup`.

- 4 Escriba el número correspondiente a la opción de quórum.**

Aparece el menú Quórum.

- 5 Escriba el número correspondiente a la opción de agregar un dispositivo de quórum. A continuación, escriba `yes` para confirmar que va a agregar un dispositivo de quórum.**

La utilidad `clsetup` pregunta qué tipo de dispositivo de quórum se desea agregar.

- 6 Escriba el número correspondiente a la opción del dispositivo de quórum de netapp\_nas. A continuación, escriba yes para confirmar que va a agregar un dispositivo de quórum netapp\_nas.**

La utilidad `clsetup` solicita que se indique el nombre del nuevo dispositivo de quórum.

- 7 Escriba el nombre del dispositivo de quórum que va a agregar.**

Elija el nombre que quiera para el dispositivo de quórum. Este nombre usa sólo para procesar futuros comandos administrativos.

La utilidad `clsetup` solicita que indique el nombre del archivador correspondiente al nuevo dispositivo de quórum.

- 8 Escriba el nombre del archivador del nuevo dispositivo de quórum.**

Es la dirección o el nombre accesible desde la red del archivador.

La utilidad `clsetup` solicita que indique el ID de LUN correspondiente al archivador.

- 9 Escriba el ID del LUN dispositivo de quórum en el archivador.**

La utilidad `clsetup` pregunta si el nuevo dispositivo de quórum se debe agregar al archivador.

- 10 Escriba yes para seguir agregando el nuevo dispositivo de quórum.**

Si se agrega correctamente el nuevo dispositivo de quórum, la utilidad `clsetup` muestra el correspondiente mensaje.

- 11 Compruebe que se haya agregado el dispositivo de quórum.**

```
clquorum list -v
```

### Ejemplo 6-3 Adición de un dispositivo de quórum NAS de NetApp

En el ejemplo siguiente se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum NAS de NetApp. El ejemplo también muestra un paso de comprobación.

```
Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on any cluster node.
[Start the clsetup utility:]
clsetup
[Select Quorum>Add a quorum device]
[Answer the questions when prompted.]
[You will need the following information.]
[Information: Example:]
[Quorum Device Netapp_nas quorum device]
[Name: qd1]
[Filer: nas1.sun.com]
[LUN ID: 0]
[Verify that the clquorum command was completed successfully:]
clquorum add -t netapp_nas -p filer=nas1.sun.com,-p lun_id=0 qd1
Command completed successfully.
[Quit the clsetup Quorum Menu and Main Menu.]
[Verify that the quorum device is added:]
```

```
clquorum list -v
Quorum Type

qd1 netapp_nas
scphyshost-1 node
scphyshost-2 node
```

## ▼ Adición de un dispositivo de quórum de servidor de quórum

### Antes de empezar

Antes de poder agregar un servidor de quórum de Oracle Solaris Cluster como dispositivo del quórum, el software del servidor del quórum de Oracle Solaris Cluster debe estar instalado en la máquina de host y el servidor de quórum debe haberse iniciado y estar en ejecución. Si desea obtener información sobre cómo instalar el servidor de quórum, consulte [“Instalación y configuración del software Servidor de quórum” de Guía de instalación del software Oracle Solaris Cluster](#).

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**
- 2 **Compruebe que todos los nodos de Oracle Solaris Cluster estén en línea y puedan comunicarse con el servidor de quórum de Oracle Solaris Cluster.**
  - a. **Compruebe que los conmutadores de red conectados directamente con los nodos del clúster cumplan uno de los criterios siguientes:**
    - El conmutador es compatible con el protocolo RSTP.
    - El conmutador tiene habilitado el modo de puerto rápido.

Se necesita una de estas funciones para que la comunicación entre los nodos del clúster y el servidor de quórum sea inmediata. Si el conmutador ralentizada dicha comunicación se ralentizase de forma significativa, el clúster interpretaría este impedimento de la comunicación como una pérdida del dispositivo de quórum.



- b. Si la red pública utiliza subredes de longitud variable o CIDR (Classless Inter-Domain Routing), modifique los archivos siguientes en cada uno de los nodos.

Si usa subredes con clases, tal y como se define en RFC 791, no es necesario seguir estos pasos.

- i. Agregue al archivo `/etc/inet/netmasks` una entrada por cada subred pública que emplee el clúster.

La entrada siguiente es un ejemplo que contiene una dirección IP de red pública y una máscara de red:

```
10.11.30.0 255.255.255.0
```

- ii. Anexe `netmask + broadcast + al nombre de host` para cada archivo `/etc/hostname.adaptador`.

```
nodename netmask + broadcast +
```

- c. Agregue el nombre de host del servidor de quórum a todos los nodos del clúster en el archivo `/etc/inet/hosts` o `/etc/inet/ipnodes`.

Agregue al archivo una asignación entre nombre de host y dirección como la siguiente:

```
ipaddress qshost1
```

*dirección\_ip* Dirección IP del equipo donde se ejecuta el servidor de quórum.

*host1\_sq* Nombre de host del equipo donde se ejecuta el servidor de quórum.

- d. Si usa un servicio de nombres, agregue la asignación entre nombre y dirección de host del servidor de quórum a la base de datos del servicio de nombres.

### 3 Inicie la utilidad `clsetup`.

```
clsetup
```

Aparece el menú principal de `clsetup`.

### 4 Escriba el número correspondiente a la opción de quórum.

Aparece el menú Quórum.

### 5 Escriba el número correspondiente a la opción de agregar un dispositivo de quórum. A continuación, escriba *yes* para confirmar que va a agregar un dispositivo de quórum.

La utilidad `clsetup` pregunta qué tipo de dispositivo de quórum se desea agregar.

### 6 Escriba el número correspondiente a la opción de un dispositivo de quórum de servidor de quórum. A continuación, escriba *yes* para confirmar que va a agregar un dispositivo de quórum de servidor de quórum.

La utilidad `clsetup` solicita que se indique el nombre del nuevo dispositivo de quórum.

**7 Escriba el nombre del dispositivo de quórum que va a agregar.**

Elija el nombre que quiera para el dispositivo de quórum. Este nombre usa sólo para procesar futuros comandos administrativos.

La utilidad `clsetup` solicita que indique el nombre del archivador correspondiente al nuevo dispositivo de quórum.

**8 Escriba el nombre de host del servidor de quórum.**

Este nombre especifica la dirección IP del equipo en que se ejecuta el servidor de quórum o el nombre de host del equipo dentro de la red.

Según la configuración IPv4 o IPv6 del host, la dirección IP del equipo debe especificarse en el archivo `/etc/hosts`, en el archivo `/etc/inet/ipnodes` o en ambos.

---

**Nota** – Todos los nodos del clúster deben tener acceso al equipo que se especifique y el equipo debe ejecutar el servidor de quórum.

---

La utilidad `clsetup` solicita que indique el número de puerto del servidor de quórum.

**9 Escriba el número de puerto que el servidor de quórum usa para comunicarse con los nodos del clúster.**

La utilidad `clsetup` solicita que confirme que debe agregarse el nuevo dispositivo de quórum.

**10 Escriba yes para seguir agregando el nuevo dispositivo de quórum.**

Si se agrega correctamente el nuevo dispositivo de quórum, la utilidad `clsetup` muestra el correspondiente mensaje.

**11 Compruebe que se haya agregado el dispositivo de quórum.**

```
clquorum list -v
```

**Ejemplo 6–4 Adición de un dispositivo de quórum de servidor de quórum**

En el ejemplo siguiente se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum de servidor de quórum. El ejemplo también muestra un paso de comprobación.

Become superuser or assume a role that provides `solaris.cluster.modify` RBAC authorization on any cluster node.

```
[Start the clsetup utility:]
clsetup
[Select Quorum > Add a quorum device]
[Answer the questions when prompted.]
[You will need the following information.]
 [Information: Example:]
 [Quorum Device quorum_server quorum device]
```

```

[Name: qd1]
[Host Machine Name: 10.11.124.84]
[Port Number: 9001]

[Verify that the clquorum command was completed successfully:]
clquorum add -t quorum_server -p qshost=10.11.124.84,-p port=9001 qd1

 Command completed successfully.
[Quit the clsetup Quorum Menu and Main Menu.]
[Verify that the quorum device is added:]
clquorum list -v

Quorum Type

qd1 quorum_server
scphyshost-1 node
scphyshost-2 node

clquorum status

=== Cluster Quorum ===
-- Quorum Votes Summary --

 Needed Present Possible

 3 5 5

-- 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

qd1 1 1 Online
d3s2 1 1 Online
d4s2 1 1 Online

```

## Eliminación o sustitución de un dispositivo de quórum

Esta sección presenta los procedimientos siguientes para eliminar o reemplazar un dispositivo de quórum:

- “Eliminación de un dispositivo de quórum” en la página 204
- “Eliminación del último dispositivo de quórum de un clúster” en la página 205
- “Sustitución de un dispositivo de quórum” en la página 206

## ▼ Eliminación de un dispositivo de quórum

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Cuando se quita un dispositivo de quórum, éste ya no participa en la votación para establecer el quórum. Todos los clústers de dos nodos necesitan como mínimo un dispositivo de quórum configurado. Si éste es el último dispositivo de quórum de un clúster, habrá un error de `clquorum(1CL)` al tratar de quitar el dispositivo de la configuración. Si va a quitar un nodo, elimine todos los dispositivos de quórum que tenga conectados.

---

**Nota** – Si el dispositivo que va a quitar es el último dispositivo de quórum del clúster, consulte el procedimiento [“Eliminación del último dispositivo de quórum de un clúster”](#) en la página 205.

---

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**
- 2 **Determine el dispositivo de quórum que se va a eliminar.**  
`# clquorum list -v`
- 3 **Ejecute la utilidad `clsetup(1CL)`.**  
`# clsetup`  
Aparece el menú principal.
- 4 **Escriba el número correspondiente a la opción de quórum.**
- 5 **Escriba el número correspondiente a la opción de quitar un dispositivo de quórum.**  
Responda a las preguntas que aparecen durante el proceso de eliminación.
- 6 **Salga de `clsetup`.**
- 7 **Compruebe que se haya eliminado el dispositivo de quórum.**  
`# clquorum list -v`

## Ejemplo 6-5 Eliminación de un dispositivo de quórum

En este ejemplo se muestra cómo quitar un dispositivo de quórum de un clúster que tiene configurados dos o más dispositivos de quórum.

Become superuser or assume a role that provides `solaris.cluster.modify` RBAC authorization on any cluster node.

```
[Determine the quorum device to be removed:]
clquorum list -v
[Start the clsetup utility:]
clsetup
[Select Quorum>Remove a quorum device]
[Answer the questions when prompted.]
Quit the clsetup Quorum Menu and Main Menu.]
[Verify that the quorum device is removed:]
clquorum list -v
```

Quorum	Type
-----	----
scphyshost-1	node
scphyshost-2	node
scphyshost-3	node

### Errores más frecuentes

Si se pierde la comunicación entre el clúster y el host del servidor de quórum durante una operación para quitar un dispositivo de quórum de servidor de quórum, debe limpiar la información de configuración caducada acerca del host del servidor de quórum. Si desea obtener instrucciones sobre cómo realizar esta limpieza, consulte [“Limpieza de la información caducada sobre clústers del servidor de quórum” en la página 220](#).

## ▼ Eliminación del último dispositivo de quórum de un clúster

Este procedimiento elimina el último dispositivo del quórum de un clúster de dos nodos mediante la opción `clquorum force, - F`. En general, primero debe quitar el dispositivo que ha fallado y después agregar el dispositivo de quórum que lo reemplaza. Si no se trata del último dispositivo de quórum de un nodo con dos clústeres, siga los pasos descritos en [“Eliminación de un dispositivo de quórum” en la página 204](#).

Agregar un dispositivo de quórum implica reconfigurar el nodo que afecta al dispositivo de quórum que ha fallado y genera una situación de error grave en el equipo. La opción Forzar permite quitar el dispositivo de quórum que ha fallado sin generar una situación de error grave en el equipo. El comando `clquorum(1CL)` permite quitar el dispositivo de la configuración. Después de quitar el dispositivo de quórum que ha fallado, puede agregar un nuevo dispositivo con el comando `clquorum add`. Consulte [“Adición de un dispositivo de quórum” en la página 192](#).

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**
- 2 **Quite el dispositivo de quórum con el comando `clquorum`. Si el dispositivo de quórum ha fallado, use la opción Forzar -F para eliminar el dispositivo que ha fallado.**

```
clquorum remove -F qd1
```

---

**Nota** – También puede situar el nodo que se va a eliminar en estado de mantenimiento y, a continuación, quitar el dispositivo de quórum con el comando `clquorum remove quorum`. Mientras el clúster se halla en modo de instalación, las opciones del menú de administración del clúster `clsetup(1CL)` no están disponibles. Consulte [“Colocación de un nodo en estado de mantenimiento” en la página 274](#) si desea obtener más información.

---

- 3 **Compruebe que se haya quitado el dispositivo de quórum.**

```
clquorum list -v
```

### Ejemplo 6–6 Eliminación del último dispositivo de quórum

En este ejemplo se muestra cómo poner el clúster en modo de mantenimiento y quitar el último dispositivo de quórum de la configuración del clúster.

```
[Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on any
 cluster node.]
[Place the cluster in install mode:]
cluster set -p installmode=enabled
[Remove the quorum device:]
clquorum remove d3
[Verify that the quorum device has been removed:]
clquorum list -v
 Quorum Type

scphyshost-1 node
scphyshost-2 node
scphyshost-3 node
```

### ▼ Sustitución de un dispositivo de quórum

Siga este procedimiento para reemplazar un dispositivo de quórum por otro dispositivo de quórum. Puede reemplazar un dispositivo de quórum por un dispositivo de tipo similar, por ejemplo sustituir un dispositivo de NAS por otro dispositivo de NAS, o bien reemplazar el dispositivo por uno de otro tipo diferente, por ejemplo sustituir un dispositivo de NAS por un disco compartido.

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### 1 Configure un nuevo dispositivo de quórum.

Primero debe agregarse un nuevo dispositivo de quórum a la configuración para que ocupe el lugar del anterior. Consulte [“Adición de un dispositivo de quórum” en la página 192](#) para agregar un nuevo dispositivo de quórum al clúster.

### 2 Quite el dispositivo que va a sustituir como dispositivo de quórum.

Consulte [“Eliminación de un dispositivo de quórum” en la página 204](#) para quitar el antiguo dispositivo de quórum de la configuración.

### 3 Si el dispositivo de quórum es un disco que ha tenido un error, sustituya el disco.

Consulte los procedimientos de hardware del receptáculo para el disco en [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#).

## Mantenimiento de dispositivos de quórum

Esta sección explica los procedimientos para mantener dispositivos de quórum.

- [“Modificación de una lista de nodos de dispositivo de quórum” en la página 207](#)
- [“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 210](#)
- [“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 211](#)
- [“Enumeración de una lista con la configuración de quórum” en la página 213](#)
- [“Reparación de un dispositivo de quórum” en la página 214](#)

### ▼ Modificación de una lista de nodos de dispositivo de quórum

Puede emplear la utilidad `clsetup(1CL)` para agregar un nodo o para quitar un nodo de la lista de nodos de un dispositivo de quórum. Para modificar la lista de nodos de un dispositivo de quórum, debe quitar el dispositivo de quórum, modificar las conexiones físicas de los nodos con el dispositivo de quórum que ha extraído y reincorporar el dispositivo de quórum a la configuración del clúster. Al agregar un dispositivo de quórum, `clquorum(1CL)` configura automáticamente las rutas entre el nodo y el disco para todos los nodos conectados con el disco.

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**
- 2 Determine el nombre del dispositivo de quórum que va a modificar.**  
`# clquorum list -v`
- 3 Inicie la utilidad `clsetup`.**  
`# clsetup`  
Aparece el menú principal.
- 4 Escriba el número correspondiente a la opción Quórum.**  
Aparece el menú Quórum.
- 5 Escriba el número correspondiente a la opción de quitar un dispositivo de quórum.**  
Siga las instrucciones. Se preguntará el nombre del disco que se va a eliminar.
- 6 Agregue o elimine las conexiones del nodo con el dispositivo de quórum.**
- 7 Escriba el número correspondiente a la opción de agregar un dispositivo de quórum.**  
Siga las instrucciones. Se solicitará el nombre del disco que se va a usar como dispositivo de quórum.
- 8 Compruebe que se haya agregado el dispositivo de quórum.**  
`# clquorum list -v`

**Ejemplo 6-7**    Modificación de una lista de nodos de dispositivo de quórum

En el ejemplo siguiente se muestra cómo usar la utilidad `clsetup` para agregar o quitar nodos de una lista de nodos de dispositivo de quórum. En este ejemplo, el nombre del dispositivo de quórum es `d2` y como resultado final del procedimiento se agrega otro nodo a la lista de nodos del dispositivo de quórum.

[Become superuser or assume a role that provides `solaris.cluster.modify` RBAC authorization on any node in the cluster.]

[Determine the quorum device name:]

```
clquorum list -v
Quorum Type

d2 shared_disk
```



```

sc-phys-schost-1 node
sc-phys-schost-2 node
sc-phys-schost-3 node

[Start the clsetup utility:]
clsetup

[Type the number that corresponds with the quorum option.]
.
[Type the number that corresponds with the option to remove a quorum device.]
.
[Answer the questions when prompted.]
[You will need the following information:]

 Information: Example:
 Quorum Device Name: d2

[Verify that the clquorum command completed successfully:]
clquorum remove d2
 Command completed successfully.

[Verify that the quorum device was removed.]
clquorum list -v
Quorum Type

sc-phys-schost-1 node
sc-phys-schost-2 node
sc-phys-schost-3 node

[Type the number that corresponds with the Quorum option.]
.
[Type the number that corresponds with the option to add a quorum device.]
.
[Answer the questions when prompted.]
[You will need the following information:]

 Information Example:
 quorum device name d2

[Verify that the clquorum command was completed successfully:]
clquorum add d2
 Command completed successfully.

Quit the clsetup utility.

[Verify that the correct nodes have paths to the quorum device.
In this example, note that phys-schost-3 has been added to the
enabled hosts list.]
clquorum show d2 | grep Hosts
=== Quorum Devices ===

Quorum Device Name: d2
 Hosts (enabled): phys-schost-1, phys-schost-2, phys-schost-3

[Verify that the modified quorum device is online.]

clquorum status d2
=== Cluster Quorum ===

```

--- Quorum Votes by Device ---

Device Name	Present	Possible	Status
-----	-----	-----	-----
d2	1	1	Online

## ▼ Colocación de un dispositivo de quórum en estado de mantenimiento

Use el comando `clquorum(1CL)` para poner un dispositivo de quórum en estado de mantenimiento. Hoy por hoy, la utilidad `clsetup(1CL)` no tiene esta capacidad. Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Ponga el dispositivo de quórum en estado de mantenimiento cuando vaya a apartar del servicio el dispositivo de quórum durante un período de tiempo prolongado. De esta forma, el número de votos de quórum del dispositivo de quórum se establece en cero y no aporta nada al número de quórum mientras se efectúan las tareas de mantenimiento en el dispositivo. La información de configuración del dispositivo de quórum se conserva durante el estado de mantenimiento.

---

**Nota** – Todos los clústers de dos nodos deben tener configurado al menos un dispositivo de quórum. Si éste es el último dispositivo de quórum de un clúster de dos nodos, `clquorum` el dispositivo no puede ponerse en estado de mantenimiento.

---

Para poner un nodo de un clúster en estado de mantenimiento, consulte “Colocación de un nodo en estado de mantenimiento” en la página 274.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

**1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

**2 Ponga el dispositivo de quórum en estado de mantenimiento.**

`# clquorum disable device`

*dispositivo*      Especifica el nombre DID del dispositivo de disco que se va a cambiar, por ejemplo `d4`.

3 Compruebe que el dispositivo de quórum esté en estado de mantenimiento.

La salida del dispositivo puesto en estado de mantenimiento debe tener cero como valor para los Votos del dispositivo del quórum.

```
clquorum status device
```

Ejemplo 6-8 Colocación de un dispositivo de quórum en estado de mantenimiento

En el ejemplo siguiente se muestra cómo poner un dispositivo de quórum en estado de mantenimiento y cómo comprobar los resultados.

```
clquorum disable d20
clquorum status d20

=== Cluster Quorum ===

--- Quorum Votes by Device ---

Device Name Present Possible Status

d20 1 1 Offline
```

**Véase también** Para volver a habilitar el dispositivo de quórum, consulte [“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 211](#).

Para poner un nodo en estado de mantenimiento, consulte [“Colocación de un nodo en estado de mantenimiento” en la página 274](#).

▼ **Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento**

Siga este procedimiento cuando desee sacar un dispositivo de quórum del estado de mantenimiento y restablecer el valor predeterminado para el número de votos.



**Precaución** – Si no especifica ni la opción `globaldev` ni `node`, el número de quórum se restablece para todo el clúster.

Al configurar un dispositivo de quórum, el software Oracle Solaris Cluster asigna al dispositivo de quórum un número de votos de  $N-1$ , donde  $N$  es el número de votos conectados al dispositivo de quórum. Por ejemplo, un dispositivo de quórum conectado a dos nodos con números de votos cuyo valor no sea cero tiene un número de quórum de uno (dos menos uno).

- Para sacar un nodo de un clúster y sus dispositivos de quórum asociados del estado de mantenimiento, consulte [“Procedimiento para sacar un nodo del estado de mantenimiento” en la página 276](#).

- Para obtener más información sobre los números de votos de quórum, consulte [“About Quorum Vote Counts”](#) de *Oracle Solaris Cluster Concepts Guide*.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en todos los nodos del clúster.**

- 2 **Restablezca el número de quórum.**

```
clquorum enable device
```

*dispositivo*      Especifica el nombre DID del dispositivo de quórum que se va a restablecer, por ejemplo `d4`.

- 3 **Si va a restablecer el número de quórum porque un nodo estaba en estado de mantenimiento, rearranque el nodo.**

- 4 **Compruebe el número de votos de quórum.**

```
clquorum show +
```

### Ejemplo 6–9 Restablecimiento del número de votos de quórum (dispositivo de quórum)

En el ejemplo siguiente se restablece el número de quórum predeterminado en un dispositivo de quórum y se comprueba el resultado.

```
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/rdsd/d20s2
Type: shared_disk
Access Mode: scsi2
Hosts (enabled): phys-schost-2, phys-schost-3

```

## ▼ Enumeración de una lista con la configuración de quórum

Este procedimiento también puede efectuarse mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Para enumerar la configuración de quórum no hace falta ser superusuario. Se puede asumir cualquier función que proporcione la autorización de RBAC `solaris.cluster.read`.

---

**Nota** – Al incrementar o reducir el número de conexiones de nodos con un dispositivo de quórum, el número de votos de quórum no se recalcula de forma automática. Puede reestablecer el voto de quórum correcto si quita todos los dispositivos de quórum y después los vuelve a agregar a la configuración. En caso de un nodo de dos clústers, agregue temporalmente un nuevo dispositivo de quórum antes de quitar y volver a agregar el dispositivo de quórum original. A continuación, elimine el dispositivo de quórum temporal.

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### ● Use `clquorum(1CL)` para enumerar en una lista la configuración de quórum.

```
% clquorum show +
```

#### Ejemplo 6–10 Enumeración en una lista la configuración de quórum

```
% 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/rdsd/d20s2
 Type: shared_disk
 Access Mode: scsi2
 Hosts (enabled): phys-schost-2, phys-schost-3
```

## ▼ Reparación de un dispositivo de quórum

Siga este procedimiento para reemplazar un dispositivo de quórum que no funcione correctamente.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### 1 Quite el dispositivo de disco que va a sustituir como dispositivo de quórum.

---

**Nota** – Si el dispositivo que pretende quitar es el último dispositivo de quórum, se recomienda agregar primero otro disco como nuevo dispositivo de quórum. Así, se dispone de un dispositivo de quórum si hubiera un error durante el procedimiento de sustitución. Consulte [“Adición de un dispositivo de quórum” en la página 192](#) para agregar un nuevo dispositivo de quórum.

---

Consulte [“Eliminación de un dispositivo de quórum” en la página 204](#) para quitar un dispositivo de disco como dispositivo de quórum.

### 2 Sustituya el dispositivo de disco.

Para reemplazar el dispositivo de disco, consulte los procedimientos sobre el receptáculo para discos en [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#).

### 3 Agregue el disco sustituido como nuevo dispositivo de quórum.

Consulte [“Adición de un dispositivo de quórum” en la página 192](#) para agregar un disco como nuevo dispositivo de quórum.

---

**Nota** – Si ha agregado un dispositivo de quórum adicional en el [Paso 1](#), ahora puede quitarlo con seguridad. Consulte [“Eliminación de un dispositivo de quórum” en la página 204](#) para quitar el dispositivo de quórum.

---

## Administración de servidores de quórum de Oracle Solaris Cluster

El servidor de quórum de Oracle Solaris Cluster ofrece un dispositivo de quórum que no es un dispositivo de almacenamiento compartido. Esta sección presenta los procedimientos siguientes para administrar servidores de quórum de Oracle Solaris Cluster:

- [“Información general sobre el archivo de configuración del servidor de quórum” en la página 215](#)
- [“Inicio y detención del software del servidor del quórum” en la página 216](#)
- [“Inicio de un servidor de quórum” en la página 216](#)
- [“Detención de un servidor de quórum” en la página 217](#)
- [“Visualización de información sobre el servidor de quórum” en la página 218](#)
- [“Limpieza de la información caducada sobre clústers del servidor de quórum” en la página 220](#)

Si desea obtener información sobre cómo instalar y configurar servidores de quórum de Oracle Solaris Cluster, consulte [“Instalación y configuración del software Servidor de quórum” de Guía de instalación del software Oracle Solaris Cluster](#).

## Información general sobre el archivo de configuración del servidor de quórum

Al instalar el software Oracle Solaris Cluster, se crea un archivo predeterminado de configuración, `/etc/scqsd/scqsd.conf`, que contiene información sobre un único servidor de quórum predeterminado. Todas las líneas del archivo `/etc/scqsd/scqsd.conf` tienen el comando especificado a continuación:

```
/usr/cluster/lib/sc/scqsd [-d quorumdirectory] [-i instancename] -p port
```

`/usr/cluster/lib/sc/scqsd`      La ruta completa donde ha instalado el software Oracle Solaris Cluster. Este valor debe ser `/usr/cluster/lib/sc/scqsd`.

`-d directorio_quórum`      La ruta del directorio donde el servidor de quórum puede almacenar datos de quórum.

	El proceso del servidor de quórum crea un archivo por clúster en este directorio para almacenar información de quórum específica del clúster. De forma predeterminada, el valor de esta opción es <code>/var/scqsd</code> . Este directorio debe ser exclusivo para cada servidor de quórum que configure.
<code>-i nombre_instancia</code>	Nombre exclusivo que se elige para la instancia del servidor de quórum.
<code>-p puerto</code>	El número de puerto en el que el servidor de quórum escucha las solicitudes del clúster. El puerto predeterminado es 9000.

Los nombres de instancias son opcionales. Si especifica un nombre para el servidor de quórum, dicho nombre debe ser exclusivo entre todos los servidores de quórum del sistema. Si elige omitir la opción del nombre de la instancia, debe referirse al servidor de quórum por el puerto en el que escucha.

## Inicio y detención del software del servidor del quórum

Estos procedimientos describen cómo iniciar y detener el software Oracle Solaris Cluster.

De manera predeterminada, estos procedimientos inician y detienen un solo servidor de quórum predefinido, salvo que haya personalizado el contenido del archivo de configuración del servidor de quórum, `/etc/scqsd/scqsd.conf`. El servidor de quórum predeterminado está vinculado al puerto 9000 y usa el directorio `/var/scqsd` para la información de quórum.

Si desea obtener información sobre cómo personalizar el archivo de configuración del servidor de quórum, consulte [“Información general sobre el archivo de configuración del servidor de quórum” en la página 215](#). Si desea información sobre cómo instalar el software del servidor de quórum, consulte [“Instalación y configuración del software Servidor de quórum” de Guía de instalación del software Oracle Solaris Cluster](#).

### ▼ Inicio de un servidor de quórum

- 1 **Conviértase en superusuario del host en que desee iniciar el software Oracle Solaris Cluster.**
- 2 **Use el comando `clquorumserver start` para iniciar el software.**  

```
/usr/cluster/bin/clquorumserver start quorumserver
```



*servidor\_quórum* Identifica el servidor de quórum. Puede utilizar el número de puerto en el que el servidor de quórum escucha. Si ha especificado un nombre de instancia en el archivo de configuración, puede utilizar este nombre.

Para iniciar un solo servidor de quórum, especifique el nombre de la instancia o el número de puerto. Para iniciar todos los servidores de quórum si se han configurado varios, utilice el operando +.

### Ejemplo 6–11 Inicio de todos los servidores de quórum configurados

En el ejemplo siguiente se inician todos los servidores de quórum configurados.

```
/usr/cluster/bin/clquorumserver start +
```

### Ejemplo 6–12 Inicio de un servidor de quórum en concreto

En el ejemplo siguiente se inicia el servidor de quórum que escucha en el puerto número 2000.

```
/usr/cluster/bin/clquorumserver start 2000
```

## ▼ Detención de un servidor de quórum

- 1 Conviértase en superusuario del host en que desee iniciar el software Oracle Solaris Cluster.
- 2 Use el comando `clquorumserver stop` para detener el software.

```
/usr/cluster/bin/clquorumserver stop [-d] quorumserver
```

-d Controla si el servidor de quórum se inicia la próxima vez que arranque el equipo. Si especifica la opción -d, el servidor de quórum no se inicia la próxima vez que arranque el equipo.

*servidor\_quórum* Identifica el servidor de quórum. Puede utilizar el número de puerto en el que el servidor de quórum escucha. Si ha especificado un nombre de instancia en el archivo de configuración, puede utilizar ese nombre.

Para detener un solo servidor de quórum, especifique el nombre de la instancia o el número de puerto. Para detener todos los servidores de quórum si se han configurado varios, utilice el operando +.

**Ejemplo 6–13** Detención de todos los servidores de quórum configurados

En el ejemplo siguiente se detienen todos los servidores de quórum configurados.

```
/usr/cluster/bin/clquorumserver stop +
```

**Ejemplo 6–14** Detención de un servidor de quórum específico

En el ejemplo siguiente se detiene el servidor de quórum que escucha en el puerto número 2000.

```
/usr/cluster/bin/clquorumserver stop 2000
```

## Visualización de información sobre el servidor de quórum

Puede visualizar información de configuración sobre el servidor de quórum. En todos los clústers donde el servidor de quórum se configuró como dispositivo de quórum, este comando muestra el nombre del clúster correspondiente, el ID de clúster, la lista de claves de reserva y una lista con las claves de registro.

### ▼ Visualización de información sobre el servidor de quórum

- 1 **Conviértase en superusuario del host donde desee visualizar la información del servidor de quórum.**

Los usuarios de las demás categorías necesitan autorización de RBAC `solaris.cluster.read`. Si desea obtener más información sobre los perfiles de derechos de RBAC, consulte la página de comando `man rbac(5)`.

- 2 **Visualice la información de configuración del dispositivo de quórum mediante el comando `clquorumserver`.**

```
/usr/cluster/bin/clquorumserver show quorumserver
```

*servidor\_quórum*      Identifica uno o más servidores de quórum. Puede especificar el servidor de quórum por el nombre de instancia o por el número del puerto. Para mostrar la información de configuración para todos los servidores de quórum, use el operando `+`.

**Ejemplo 6–15** Visualización de la configuración de un servidor de quórum

En el ejemplo siguiente se muestra la información de configuración del servidor de quórum que usa el puerto 9000. El comando muestra información correspondiente a cada uno de los clústers

que tienen configurado el servidor de quórum como dispositivo de quórum. Esta información incluye el nombre del clúster y su ID, así como la lista de claves de registro y de reserva del dispositivo.

En el ejemplo siguiente, los nodos con ID 1, 2, 3 y 4 del clúster `bastille` han registrado sus claves en el servidor de quórum. Asimismo, debido a que el Nodo 4 es propietario de la reserva del dispositivo de quórum, su clave figura en la lista de reservas.

```
/usr/cluster/bin/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
```

#### **Ejemplo 6–16** Visualización de la configuración de varios servidores de quórum

En el ejemplo siguiente se muestra la información de configuración de tres servidores de quórum, `qs1`, `qs2` y `qs3`.

```
/usr/cluster/bin/clquorumserver show qs1 qs2 qs3
```

#### **Ejemplo 6–17** Visualización de la configuración de todos los servidores de quórum en ejecución

En el ejemplo siguiente se muestra la información de configuración de todos los servidores de quórum en ejecución:

```
/usr/cluster/bin/clquorumserver show +
```

# Limpieza de la información caducada sobre clústers del servidor de quórum

Para eliminar un dispositivo de quórum del tipo `quorumserver`, use el comando `clquorum remove` como se describe en [“Eliminación de un dispositivo de quórum” en la página 204](#). En condiciones normales de funcionamiento, este comando también elimina la información del servidor de quórum relativa al host del servidor de quórum. Sin embargo, si el clúster pierde la comunicación con el host del servidor de quórum, al eliminar el dispositivo de quórum dicha información no se limpia.

La información sobre los clústers del servidor de quórum perderá su validez en los casos siguientes:

- Cuando un clúster se anula sin que antes se haya eliminado dispositivo de quórum del clúster mediante el comando `clquorum remove`.
- Cuando un dispositivo de quórum del tipo `quorum_server` se elimina de un clúster mientras el host del servidor de quórum está detenido.



**Precaución** – Si un dispositivo de quórum del tipo `quorumserver` (servidor de quórum) aún no se ha eliminado del clúster, utilizar este procedimiento para limpiar un servidor de quórum válido puede afectar negativamente al quórum del clúster.



## Limpieza de la información de configuración del servidor de quórum

**Antes de empezar**

Quite el dispositivo de quórum del clúster de servidor de quórum como se describe en [“Eliminación de un dispositivo de quórum” en la página 204](#).



**Precaución** – Si el clúster sigue utilizando este servidor de quórum, efectuar este procedimiento afectará negativamente al quórum del clúster.

- 1 **Conviértase en superusuario del host del servidor de quórum.**
- 2 **Use el comando `clquorumserver clear` para limpiar el archivo de configuración.**  

```
clquorumserver clear -c clustername -I clusterID quorumserver [-y]
```

`-c nombre_clúster` El nombre del clúster que el servidor de quórum utilizaba anteriormente como dispositivo de quórum.

Puede obtener el nombre del clúster si ejecuta `cluster show` en un nodo del clúster.

`-I ID_clúster` ID del clúster.

	El ID del clúster es un número hexadecimal de ocho dígitos. Puede obtener el ID del clúster si ejecuta <code>cluster show</code> en un nodo del clúster.
<i>servidor_quórum</i>	Un identificador de uno o más servidores de quórum.
	El servidor de quórum se puede identificar mediante un número de puerto o un nombre de instancia. Los nodos del clúster usan el número del puerto para comunicarse con el servidor de quórum. El nombre de instancia aparece especificado en el archivo de configuración del servidor de quórum, <code>/etc/scqsd/scqsd.conf</code> .
-y	Fuerce el comando <code>clquorumserver clear</code> para limpiar la información del clúster del archivo de configuración sin solicitar antes la confirmación.
	Recurra a esta opción sólo si tiene la seguridad de que desea eliminar la información de clústers obsoleta del servidor de quórum.

- 3 (Opcional) Si no hay ningún otro dispositivo de quórum configurado en esta instancia del servidor, detenga el servidor de quórum.

#### **Ejemplo 6–18** Limpieza de la información obsoleta de clústers de la configuración del servidor de quórum

En este ejemplo se elimina información correspondiente al clúster denominado `sc-cluster` del servidor de quórum que emplea el puerto 9000.

```
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
```



# Administración de interconexiones de clústers y redes públicas

---

En este capítulo se presentan los procedimientos de software para administrar interconexiones entre Oracle Solaris Cluster y redes públicas.

Administrar las interconexiones entre los clústers y las redes públicas implica procedimientos de software y de hardware. Durante la instalación y configuración inicial del clúster se suelen configurar las interconexiones de los clústers y las redes públicas, incluidos los grupos de varias rutas IP. El método de ruta múltiple se instala automáticamente con el sistema operativo Oracle Solaris 10 y se debe habilitar para utilizarlo. Si después debe modificarse la configuración de las interconexiones entre redes y clúster, puede usar los procedimientos de software descritos en este capítulo. Si desea obtener información sobre cómo configurar grupos de varias rutas IP en un clúster, consulte la sección [“Administración de redes públicas” en la página 238](#).

En este capítulo hay información y procedimientos para los temas siguientes.

- [“Administración de interconexiones de clústers” en la página 223](#)
- [“Administración de redes públicas” en la página 238](#)

Si desea ver una descripción detallada de los procedimientos relacionados con este capítulo, consulte la [Tabla 7-1](#) y la [Tabla 7-3](#).

Consulte *Oracle Solaris Cluster Concepts Guide* para obtener información general y de referencia sobre las interconexiones entre clústeres y redes públicas.

## Administración de interconexiones de clústers

En esta sección se describen los procedimientos para reconfigurar las interconexiones de clústers, por ejemplo adaptador de transporte del clúster y cable de transporte de clústers. Para poder realizarlos, Oracle Solaris Cluster debe estar instalado.

Casi siempre se puede emplear la utilidad `clsetup` para administrar el transporte del clúster en la interconexión de clúster. Para obtener más información, consulte la página de comando `man`

`clsetup(1CL)`. Todos los comandos relacionados con las interconexiones de los clústeres deben ejecutarse en el nodo de votación del clúster global.

Si desea información sobre los procedimientos de instalación del software de clúster, consulte *Guía de instalación del software Oracle Solaris Cluster*. Para ver los procedimientos sobre tareas de mantenimiento y reparación de componentes de hardware de los clústeres, consulte *Oracle Solaris Cluster 3.3 Hardware Administration Manual*.

**Nota** – Durante los procedimientos de interconexión de clúster, en general es factible optar por usar el nombre de puerto predeterminado, si se da el caso. El nombre de puerto predeterminado es el mismo que el número de ID de nodo interno del nodo que aloja el extremo del cable donde se sitúa el adaptador.

TABLA 7-1 Lista de tareas: administrar la interconexión de clúster

Tarea	Instrucciones
Administrar el transporte del clúster con <code>clsetup(1CL)</code>	“Obtención de acceso a las utilidades de configuración del clúster” en la página 27
Comprobar el estado de la interconexión de los clústers con <code>clinterconnect status</code>	“Comprobación del estado de la interconexión de clúster” en la página 225
Agregar un cable de transporte de clúster, un adaptador de transporte o un conmutador mediante <code>clsetup</code>	“Adición de dispositivos de cable de transporte de clúster, adaptadores o conmutadores de transporte” en la página 226
Quitar un cable de transporte de clúster, un adaptador de transporte o un conmutador de transporte mediante <code>clsetup</code>	“Eliminación de cable de transporte de clúster, adaptadores de transporte y conmutadores de transporte” en la página 229
Habilitar un cable de transporte de clúster mediante <code>clsetup</code>	“Habilitación de un cable de transporte de clúster” en la página 232
Inhabilitar un cable de transporte de clúster mediante <code>clsetup</code>	“Inhabilitación de un cable de transporte de clúster” en la página 233
Determinar el número de instancia de un adaptador de transporte	“Determinación del número de instancia de un adaptador de transporte” en la página 235
Cambiar la dirección IP o el intervalo de direcciones de un clúster	“Modificación de la dirección de red privada o del intervalo de direcciones de un clúster” en la página 236



# Reconfiguración dinámica con interconexiones de clústers

Al completar operaciones de reconfiguración dinámica en las interconexiones de clústers es preciso tener en cuenta una serie de factores.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la reconfiguración dinámica de Oracle Solaris también son válidos para la reconfiguración dinámica en Oracle Solaris Cluster (excepto la inactividad del sistema operativo). Por lo tanto, consulte la documentación sobre la reconfiguración dinámica de Oracle Solaris *antes* de utilizarla con el software Oracle Solaris Cluster. Preste especial atención a los problemas que afecten a los dispositivos de E/S no conectados a red durante las operaciones de desconexión de reconfiguración dinámica.
- El software Oracle Solaris Cluster rechaza las operaciones de eliminación de tarjetas de reconfiguración dinámica realizadas en interfaces de interconexión privada activas.
- Debe eliminar del clúster por completo un adaptador activo, con el fin de desarrollar DR en una interconexión de clúster activa. Use el menú `cl setup` o los comandos correspondientes.



**Precaución** – El software Oracle Solaris Cluster necesita que cada nodo del clúster disponga al menos de una ruta que funcione y se comuniquen con cualquier otro nodo del clúster. No inhabilite una interfaz de interconexión privada que proporcione la última ruta a cualquier nodo del clúster.

Aplique los procedimientos siguientes en el orden indicado al realizar operaciones de reconfiguración dinámica en interfaces de redes públicas.

TABLA 7-2 Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas

Tarea	Instrucciones
1. Inhabilitar y eliminar la interfaz de la interconexión activa	<a href="#">“Reconfiguración dinámica con interfaces de red pública” en la página 240</a>
2. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	<i>Sun Enterprise 10000 Dynamic Reconfiguration Reference Manual</i> (de la recopilación <i>Solaris 10 on Sun Hardware</i> )

## ▼ Comprobación del estado de la interconexión de clúster

Este procedimiento también puede efectuarse con la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

Para llevar a cabo este procedimiento hace falta iniciar sesión como superusuario.

- 1 Compruebe el estado de la interconexión de clúster.  
% clinterconnect status
- 2 Consulte la tabla siguiente para ver los mensajes de estado comunes.

Mensaje de estado	Descripción y posibles acciones
Path online	La ruta funciona correctamente. No es necesario hacer nada.
Path waiting	La ruta se está inicializando. No es necesario hacer nada.
Faulted	La ruta no funciona. Puede tratarse de un estado transitorio cuando las rutas pasan del estado de espera al estado en línea. Si al volver a ejecutar clinterconnect status sigue apareciendo este mensaje, tome las medidas pertinentes para solucionarlo.

Ejemplo 7-1 Comprobación del estado de la interconexión de clúster

En el ejemplo siguiente se muestra el estado de una interconexión de clúster en funcionamiento.

```
% clinterconnect status
-- Cluster Transport Paths --
 Endpoint Endpoint Status

Transport path: phys-schost-1:qfe1 phys-schost-2:qfe1 Path online
Transport path: phys-schost-1:qfe0 phys-schost-2:qfe0 Path online
Transport path: phys-schost-1:qfe1 phys-schost-3:qfe1 Path online
Transport path: phys-schost-1:qfe0 phys-schost-3:qfe0 Path online
Transport path: phys-schost-2:qfe1 phys-schost-3:qfe1 Path online
Transport path: phys-schost-2:qfe0 phys-schost-3:qfe0 Path online
```

▼ Adición de dispositivos de cable de transporte de clúster, adaptadores o conmutadores de transporte

Si desea obtener información acerca de los requisitos necesarios para el transporte privado del clúster, consulte “[Interconnect Requirements and Restrictions](#)” de *Oracle Solaris Cluster 3.3 Hardware Administration Manual*.

También puede llevar a cabo este procedimiento mediante la interfaz gráfica de usuario (GUI) de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que se instale la parte física de cable de transporte de clústers.**  
Si desea información sobre cómo instalar un cable de transporte de clúster, consulte el [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#).
- 2 Conviértase en superusuario en un nodo de clúster.**
- 3 Inicie la utilidad `clsetup`.**  
`# clsetup`  
Aparece el menú principal.
- 4 Escriba el número correspondiente a la opción de mostrar el menú Interconexión del clúster.**
- 5 Escriba el número correspondiente a la opción de agregar un cable de transporte.**  
Siga las instrucciones y escriba la información que se solicite.
- 6 Escriba el número correspondiente a la opción de agregar el adaptador de transporte a un nodo.**  
Siga las instrucciones y escriba la información que se solicite.  
  
Si tiene pensado usar alguno de los adaptadores siguientes para la interconexión de clúster, agregue la entrada pertinente al archivo `/etc/system` en cada nodo del clúster. La entrada surtirá efecto la próxima vez que se arranque el sistema.

Adaptador	Entrada
ce	establecer <code>ce:ce_taskq_disable=1</code>
ipge	establecer <code>ipge:ipge_taskq_disable=1</code>
ixge	establecer <code>ixge:ixge_taskq_disable=1</code>

- 7 Escriba el número correspondiente a la opción de agregar el conmutador de transporte.**  
Siga las instrucciones y escriba la información que se solicite.

## 8 Compruebe que se haya agregado el cable de transporte de clúster, el adaptador o el conmutador de transporte.

```
clinterconnect show node:adapter,adapternode
clinterconnect show node:adapter
clinterconnect show node:switch
```

### Ejemplo 7-2 Adición de un cable de transporte de clúster, un adaptador de transporte o un conmutador de transporte

En el ejemplo siguiente se muestra cómo agregar un cable de transporte, un adaptador o un conmutador de transporte a un nodo mediante la utilidad `clsetup`.

```
[Ensure that the physical cable is installed.]
[Start the clsetup utility:]
clsetup
[Select Cluster interconnect]

[Select either Add a transport cable,
Add a transport adapter to a node,
or Add a transport switch.]
[Answer the questions when prompted.]
[You Will Need:]
[Information: Example:]
node names phys-schost-1
adapter names qfe2
switch names hub2
transport type dlpi
[Verify that the clinterconnect
command completed successfully:]Command completed successfully.
Quit the clsetup Cluster Interconnect Menu and Main Menu.
[Verify that the cable, adapter, and switch are added:]
clinterconnect show phys-schost-1:qfe2,hub2
===Transport Cables===
Transport Cable: phys-schost-1:qfe2@0,hub2
Endpoint1: phys-schost-2:qfe0@0
Endpoint2: ethernet-1@2 ????. Should this be hub2?
State: Enabled

clinterconnect show phys-schost-1:qfe2
=== Transport Adepters for qfe2
Transport Adapter: qfe2
Adapter State: Enabled
Adapter Transport Type: dlpi
Adapter Property (device_name): ce
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
```

**Pasos siguientes** Para comprobar el estado de la interconexión del cable de transporte de clúster, consulte [“Comprobación del estado de la interconexión de clúster” en la página 225.](#)

## ▼ Eliminación de cable de transporte de clúster, adaptadores de transporte y conmutadores de transporte

Este procedimiento también puede efectuarse mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Complete el procedimiento descrito a continuación para quitar cables, adaptadores y conmutadores de transporte de la configuración de un nodo. Cuando se inhabilita un cable, sus dos extremos continúan configurados. Un adaptador no se puede quitar si sigue en uso como extremo de un cable de transporte.



**Precaución** – Cada nodo del clúster debe contar al menos una ruta de transporte operativa que lo comuniquen con todos los demás nodos del clúster. No debe haber ningún par de nodos que estén aislados entre sí. Antes de inhabilitar un cable, compruebe siempre el estado de interconexión de clúster de un nodo. Sólo se debe inhabilitar una conexión por cable tras haber comprobado que sea redundante. Es decir, asegúrese de que haya disponible otra conexión. Si se inhabilita el último cable operativo de un nodo, dicho nodo deja de ser miembro del clúster.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario en un nodo de clúster.**
- 2 **Compruebe el estado de la ruta de transporte restante del clúster.**

```
clinterconnect status
```



**Precaución** – Si recibe un mensaje de error del tipo "ruta errónea" al intentar eliminar un nodo de un clúster de dos nodos, investigue el problema antes de seguir con este procedimiento. Un problema de ese tipo podría denotar que una ruta del nodo no está disponible. Si se elimina la última ruta operativa que quedaba, el nodo deja de pertenecer al clúster; como consecuencia, podría darse una reconfiguración del clúster.

---

**3 Inicie la utilidad `clsetup`.**

`# clsetup`

Aparece el menú principal.

**4 Escriba el número correspondiente a la opción de acceder al menú interconexión de clúster.**

**5 Escriba el número correspondiente a la opción de inhabilitar el cable de transporte.**

Siga las instrucciones y escriba la información que se solicite. Debe saber cuáles son los nombres válidos de nodos, adaptadores y conmutadores.

**6 Escriba el número correspondiente a la opción de quitar el cable de transporte.**

Siga las instrucciones y escriba la información que se solicite. Debe saber cuáles son los nombres válidos de nodos, adaptadores y conmutadores.

---

**Nota** – Si va a quitar un cable, desconecte el cable entre el puerto y el dispositivo de destino.

---

**7 Escriba el número correspondiente a la opción de quitar de un nodo el adaptador de transporte.**

Siga las instrucciones y escriba la información que se solicite. Debe saber cuáles son los nombres válidos de nodos, adaptadores y conmutadores.

---

**Nota** – Si va a quitar de un nodo un adaptador, consulte [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#) para ver los procedimientos de tareas de mantenimiento de hardware.

---

**8 Escriba el número correspondiente a la opción de quitar un conmutador de transporte.**

Siga las instrucciones y escriba la información que se solicite. Debe saber cuáles son los nombres válidos de nodos, adaptadores y conmutadores.

---

**Nota** – No es posible eliminar un conmutador si alguno de los puertos se sigue usando como extremo de algún cable de transporte.

---

**9 Compruebe que se haya quitado el cable, adaptador o conmutador.**

```
clinterconnect show node:adapter,adapternode
clinterconnect show node:adapter
clinterconnect show node:switch
```

El cable o adaptador de transporte eliminado del nodo correspondiente no debe aparecer en la salida de este comando.

### Ejemplo 7-3 Eliminación de un cable de transporte, un adaptador de transporte o un conmutador de transporte

En el ejemplo siguiente se muestra cómo quitar un cable, un adaptador o un conmutador de transporte mediante el comando `clsetup`.

```
[Become superuser on any node in the cluster.]
[Start the utility:]
clsetup
[Select Cluster interconnect.]
[Select either Remove a transport cable,
Remove a transport adapter to a node,
or Remove a transport switch.]
[Answer the questions when prompted.]
 You Will Need:
 Information Example:
 node names phys-schost-1
 adapter names qfe1
 switch names hub1
[Verify that the clinterconnect
command was completed successfully:]
Command completed successfully.
[Quit the clsetup utility Cluster Interconnect Menu and Main Menu.]
[Verify that the cable, adapter, or switch is removed:]
clinterconnect show phys-schost-1:qfe2,hub2
===Transport Cables===
Transport Cable: phys-schost-2:qfe2@0,hub2
 Cable Endpoint1: phys-schost-2:qfe0@0
 Cable Endpoint2: ethernet-1@2 ??? Should this be hub2???
 Cable State: Enabled

clinterconnect show phys-schost-1:qfe2
=== Transport Adepters for qfe2
Transport Adapter: qfe2
 Adapter State: Enabled
 Adapter Transport Type: dlpi
 Adapter Property (device_name): ce
 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 SState (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

## ▼ **Habilitación de un cable de transporte de clúster**

Este procedimiento también puede efectuarse mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Esta opción se utiliza para habilitar un cable de transporte de clúster que ya existe.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario en un nodo de clúster.**
- 2 Inicie la utilidad `clsetup`.**  
`# clsetup`  
Aparece el menú principal.
- 3 Escriba el número correspondiente a la opción de acceder al menú interconexión de clúster y pulse la tecla Intro.**
- 4 Escriba el número correspondiente a la opción de habilitar el cable de transporte y pulse la tecla Intro.**  
Siga las instrucciones cuando se solicite. Debe especificar los nombres del nodo y del adaptador de uno de los extremos del cable que está tratando de identificar.
- 5 Compruebe que el cable esté habilitado.**  
`# clinterconnect show node:adapter,adapternode`

### **Ejemplo 7-4** Habilitación de un cable de transporte de clúster

En este ejemplo se muestra cómo habilitar un cable de transporte de clúster en el adaptador `qfe-1`, ubicado en el nodo `phys - schost-2`.

```
[Become superuser on any node.]
[Start the clsetup utility:]
clsetup
```



```
[Select Cluster interconnect>Enable a transport cable.[

[Answer the questions when prompted.[
[You will need the following information.[
 You Will Need:
Information: Example:
 node names phys-schost-2
 adapter names qfe1
 switch names hub1
[Verify that the scinterconnect
command was completed successfully:]

clinterconnect enable phys-schost-2:qfe1

Command completed successfully.
[Quit the clsetup Cluster Interconnect Menu and Main Menu.]
[Verify that the cable is enabled:]
clinterconnect show phys-schost-1:qfe2,hub2
 Transport cable: phys-schost-2:qfe1@0 ethernet-1@2 Enabled
 Transport cable: phys-schost-3:qfe0@1 ethernet-1@3 Enabled
 Transport cable: phys-schost-1:qfe0@0 ethernet-1@1 Enabled
```

## ▼ Inhabilitación de un cable de transporte de clúster

Este procedimiento también puede efectuarse mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Quizá deba inhabilitar un cable de transporte de clúster para cerrar temporalmente una ruta de interconexión de clúster. Un cierre temporal es útil al buscar la solución a un problema de la interconexión de clúster o al sustituir hardware de interconexión de clúster.

Cuando se inhabilita un cable, sus dos extremos continúan configurados. Un adaptador no se puede quitar si sigue en uso como extremo de un cable de transporte.



**Precaución** – Cada nodo del clúster debe contar al menos una ruta de transporte operativa que lo comunique con todos los demás nodos del clúster. No debe haber ningún par de nodos que estén aislados entre sí. Antes de inhabilitar un cable, compruebe siempre el estado de interconexión de clúster de un nodo. Sólo se debe inhabilitar una conexión por cable tras haber comprobado que sea redundante. Es decir, asegúrese de que haya disponible otra conexión. Si se inhabilita el último cable operativo de un nodo, dicho nodo deja de ser miembro del clúster.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario en un nodo de clúster.**
- 2 **Compruebe el estado de la interconexión de clúster antes de inhabilitar un cable.**

```
clinterconnect status
```



**Precaución** – Si recibe un mensaje de error del tipo "ruta errónea" al intentar eliminar un nodo de un clúster de dos nodos, investigue el problema antes de seguir con este procedimiento. Un problema de ese tipo podría denotar que una ruta del nodo no está disponible. Si se elimina la última ruta operativa que quedaba, el nodo deja de pertenecer al clúster; como consecuencia, podría darse una reconfiguración del clúster.

- 3 **Inicie la utilidad clsetup.**  
# clsetup  
Aparece el menú principal.
- 4 **Escriba el número correspondiente a la opción de acceder al menú interconexión de clúster y pulse la tecla Intro.**
- 5 **Escriba el número correspondiente a la opción de inhabilitar el cable de transporte y pulse la tecla Intro.**

Siga las instrucciones y escriba la información que se solicite. Se inhabilitarán todos los componentes de la interconexión de este clúster. Debe especificar los nombres del nodo y del adaptador de uno de los extremos del cable que está tratando de identificar.

- 6 **Compruebe que se haya inhabilitado el cable.**  
# clinterconnect show node:adapter,adapternode

### Ejemplo 7-5 Inhabilitación de un cable de transporte de clúster

En este ejemplo se muestra cómo inhabilitar un cable de transporte de clúster en el adaptador qfe-1, ubicado en el nodo phys-schost-2.

```
[Become superuser on any node.]
[Start the clsetup utility:]
clsetup
[Select Cluster interconnect>Disable a transport cable.]
```

```
[Answer the questions when prompted.]
[You will need the following information.]
[You Will Need:]
```

Information:	Example:
node names	phys-schost-2
adapter names	qfe1
switch names	hub1

```
[Verify that the clinterconnect
command was completed successfully:]
```

```
Command completed successfully.
[Quit the clsetup Cluster Interconnect Menu and Main Menu.]
[Verify that the cable is disabled:]
clinterconnect show -p phys-schost-1:qfe2,hub2
Transport cable: phys-schost-2:qfe1@0 ethernet-1@2 Disabled
Transport cable: phys-schost-3:qfe0@1 ethernet-1@3 Enabled
Transport cable: phys-schost-1:qfe0@0 ethernet-1@1 Enabled
```

## ▼ Determinación del número de instancia de un adaptador de transporte

Debe determinarse el número de instancia de un adaptador de transporte para asegurarse de que se agregue y quite el adaptador correcto mediante el comando `clsetup`. El nombre del adaptador es una combinación del tipo de adaptador y del número de instancia de dicho adaptador.

**1 Según el número de ranura, busque el nombre del adaptador.**

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```
prtdiag
...
===== IO Cards =====
 Bus Max
 Bus Dev,
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
...
```

**2 Con la ruta del adaptador, busque el número de instancia del adaptador.**

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```
grep sci /etc/path_to_inst
"/xyz@1f,400/pci11c8,o@2" 0 "ttt"
"/xyz@1f,4000.pci11c8,0@4 "ttt"
```

**3 Mediante el nombre y el número de ranura del adaptador, busque su número de instancia.**

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```
prtconf
...
xyz, instance #0
 xyz11c8,0, instance #0
 xyz11c8,0, instance #1
...
```

## ▼ Modificación de la dirección de red privada o del intervalo de direcciones de un clúster

Siga este procedimiento para modificar una dirección de red privada, un intervalo de direcciones de red o ambas cosas.

### Antes de empezar

Compruebe que esté habilitado el acceso de shell remoto (rsh(1M)) o de shell seguro (ssh(1)) en todos los nodos de clúster para superusuario.

- 1 **Rearranque todos los nodos de clúster en un modo que no sea de clúster. Para ello, aplique los subpasos siguientes en cada nodo del clúster:**

- a. **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` en el clúster que se va a iniciar en un modo que no es de clúster.**

- b. **Cierre el nodo con los comandos `clnode evacuate` y `cluster shutdown`.**

El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos desde los nodos de votación o sin votación del nodo especificado hasta el siguiente nodo por orden de preferencia, sea de votación o sin votación.

```
clnode evacuate node
cluster shutdown -g0 -y
```

- 2 **Inicie la utilidad `clsetup` desde un nodo.**

Cuando se ejecuta en un modo que no sea de clúster, la utilidad `clsetup` muestra el menú principal para operaciones de un modo que no sea de clúster.

- 3 **Escriba el número correspondiente a la opción de cambiar el rango de direcciones IP y pulse la tecla Intro.**

La utilidad `clsetup` muestra la configuración de red privada actual; a continuación, pregunta si desea modificar esta configuración.

- 4 **Para cambiar la dirección IP de red privada o el intervalo de direcciones de red IP, escriba `yes` y pulse la tecla Intro.**

La utilidad `clsetup` muestra la dirección IP de red privada, `172.16.0.0` y pregunta si desea aceptarla de forma predeterminada.

## 5 Cambie o acepte la dirección IP de red privada.

- **Para aceptar la dirección IP de red privada predeterminada y cambiar el rango de direcciones IP, escriba *yes* y pulse la tecla *Intro*.**

La utilidad `clsetup` pregunta si desea aceptar la máscara de red predeterminada. Vaya al siguiente paso para escribir su respuesta.

- **Para cambiar la dirección IP de red privada predeterminada, siga los subpasos descritos a continuación.**

- a. **Escriba *no* como respuesta a la pregunta de la utilidad `clsetup` sobre si desea aceptar la dirección predeterminada; a continuación, pulse la tecla *Intro*.**

La utilidad `clsetup` solicita la nueva dirección IP de red privada.

- b. **Escriba la dirección IP nueva y pulse la tecla *Intro*.**

La utilidad `clsetup` muestra la máscara de red predeterminada; a continuación, pregunta si desea aceptar la máscara de red predeterminada.

## 6 Modifique o acepte el rango de direcciones IP de red privada predeterminado.

La máscara de red predeterminada es 255 . 255 . 240 . 0. Este rango de direcciones IP predeterminado admite hasta 64 nodos, 12 clústeres de zona y 10 redes privadas en el clúster.

- **Para aceptar el rango de direcciones IP predeterminado, escriba *yes* y pulse la tecla *Intro*.**

A continuación, vaya al siguiente paso.

- **Para cambiar el rango de direcciones IP, realice los subpasos siguientes.**

- a. **Escriba *no* como respuesta a la pregunta de la utilidad `clsetup` sobre si desea aceptar el rango de direcciones predeterminado; a continuación, pulse la tecla *Intro*.**

Si rechaza la máscara de red predeterminada, la utilidad `clsetup` solicita el número de nodos, redes privadas y clústeres de zona que tiene previsto configurar en el clúster.

- b. **Especifique el número de nodos, redes privadas y clústeres de zona que tiene previsto configurar en el clúster.**

A partir de estas cantidades, la utilidad `clsetup` calcula dos máscaras de red como propuesta:

- La primera máscara de red es la mínima para admitir el número de nodos, redes privadas y clústeres de zona que haya especificado.
- La segunda máscara de red admite el doble de nodos, redes privadas y clústeres de zona que haya especificado para asumir un posible crecimiento en el futuro.

- c. Especifique una de las máscaras de red, u otra diferente, que admita el número previsto de nodos, redes privadas y clústeres de zona.
- 7 Escriba **yes** como respuesta a la pregunta de la utilidad `c1setup` sobre si desea continuar con la actualización.
  - 8 Cuando haya finalizado, salga de la utilidad `c1setup`.
  - 9 Rearranque todos los nodos del clúster de nuevo en modo de clúster; para ello, siga estos subpasos en cada uno de los nodos:
    - a. Arranque el nodo.
      - En los sistemas basados en SPARC, ejecute el comando siguiente.  
`ok boot`
      - En los sistemas basados en x86, ejecute los comandos siguientes.  
 Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:
 

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```
  - 10 Compruebe que el nodo haya arrancado sin errores y esté en línea.  
`# cluster status -t node`

## Administración de redes públicas

El software Oracle Solaris Cluster admite la implementación del software de Oracle Solaris de IPMP (Internet Protocol network Multipathing) para las redes públicas. La administración básica de IPMP es igual en los entornos con clústeres y sin clústeres. El método de ruta múltiple se instala automáticamente al instalar el sistema operativo Oracle Solaris 10 y se debe habilitar para utilizarlo. La administración de varias rutas se trata en la documentación correspondiente del sistema operativo Oracle Solaris. No obstante, consulte las indicaciones siguientes antes de administrar IPMP en un entorno de Oracle Solaris Cluster.

## Administración de grupos de varias rutas de red IP en un clúster

Antes de realizar procedimientos relacionados con IPMP en un clúster, tenga en cuenta las indicaciones siguientes.

- Cada adaptador de red pública debe pertenecer a un grupo IPMP.
- La variable `local-mac-address?` debe presentar un valor `true` para los adaptadores de Ethernet.
- Puede utilizar grupos IPMP basados en sondeos o en vínculos en un clúster. Un grupo IPMP basado en sondeos prueba la dirección IP de destino y proporciona la mayor protección posible mediante el reconocimiento de más condiciones que pueden poner en peligro la disponibilidad.
- Debe configurar una dirección IP de prueba para cada adaptador en los tipos siguientes de grupos de varias rutas:
  - Todos los grupos de ruta múltiple con varios adaptadores requieren direcciones IP de prueba. Los grupos de ruta múltiple con un solo adaptador no requieren direcciones IP de prueba.
- Las direcciones IP de prueba de todos los adaptadores del mismo grupo de varias rutas deben pertenecer a una sola subred de IP.
- Las aplicaciones normales no deben utilizar direcciones IP, ya que no tienen alta disponibilidad.
- No hay restricciones respecto a la asignación de nombres en los grupos de varias rutas. No obstante, al configurar un grupo de recursos, la convención de asignación de nombres `netiflist` estipula que están formados por el nombre de cualquier ruta múltiple seguido del número de ID de nodo o del nombre de nodo. Por ejemplo, supongamos que hay un grupo de varias rutas denominado `sc_ipmp0`; el nombre `netiflist` asignado podría ser `sc_ipmp0@1` o `sc_ipmp0@phys-schost-1`, donde el adaptador está en el nodo `phys-schost-1`, que tiene un ID de nodo que es 1.
- Evite anular la configuración (desinstalar) de un adaptador o desconectarlo de un grupo de varias rutas de red IP sin haber conmutado primero las direcciones IP del adaptador para que se trasladen a otro adaptador alternativo del grupo mediante el comando `if_mpadm(1M)`.
- Evite volver a instalar los cables de los adaptadores en subredes distintas sin haberlos eliminado previamente de sus grupos de varias rutas respectivos.
- Las operaciones relacionadas con los adaptadores lógicos se pueden realizar en un adaptador, incluso si está activada la supervisión del grupo de varias rutas.
- Debe mantener al menos una conexión de red pública para cada nodo del clúster. Sin una conexión de red pública no es posible tener acceso al clúster.

- Para ver el estado de los grupos de varias rutas de red IP de un clúster, use el comando `clinterconnect status`.

Si desea obtener más información sobre varias rutas de red IP, consulte la documentación pertinente en el conjunto de documentación de administración de sistemas del sistema operativo Oracle Solaris.

TABLA 7-3 Mapa de tareas: administrar la red pública

Versión del sistema operativo Oracle Solaris	Instrucciones
Sistema operativo Oracle Solaris 10	Consulte el apartado relacionado a temas de varias rutas de redes IP en <i>Guía de administración del sistema: servicios IP</i> .

Si desea información sobre los procedimientos de instalación del software de clúster, consulte *Guía de instalación del software Oracle Solaris Cluster*. Para ver los procedimientos sobre tareas de mantenimiento y reparación de componentes de hardware de redes públicas, consulte *Oracle Solaris Cluster 3.3 Hardware Administration Manual*.

## Reconfiguración dinámica con interfaces de red pública

Debe tener en cuenta algunos aspectos al desarrollar operaciones de reconfiguración dinámica con interfaces de red pública en un clúster.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la reconfiguración dinámica de Oracle Solaris también son válidos para la reconfiguración dinámica en Oracle Solaris Cluster (excepto la inactividad del sistema operativo). Por lo tanto, consulte la documentación sobre la reconfiguración dinámica de Oracle Solaris *antes* de utilizarla con el software Oracle Solaris Cluster. Preste especial atención a los problemas que afecten a los dispositivos de E/S no conectados a red durante las operaciones de desconexión de reconfiguración dinámica.
- Las operaciones de eliminación de tarjetas de DR sólo pueden finalizar correctamente si las interfaces de red pública no están activas. Antes de eliminar una interfaz de red pública activa, conmute las direcciones IP del adaptador que se va a quitar a otro adaptador presente en el grupo de varias rutas con el comando `if_mpadm(1M)`.
- Si intenta eliminar una tarjeta de interfaz de red pública sin haberla inhabilitado como interfaz de red pública activa, Oracle Solaris Cluster rechaza la operación e identifica la interfaz que habría sido afectada por la operación.





**Precaución** – En el caso de grupos de varias rutas con dos adaptadores, si el adaptador de red restante sufre un error mientras se elimina la reconfiguración dinámica en el adaptador de red inhabilitado, la disponibilidad se verá afectada. El adaptador que se conserva no tiene ningún elemento al que migrar tras error mientras dure la operación de reconfiguración dinámica.

Aplice los procedimientos siguientes en el orden indicado al realizar operaciones de reconfiguración dinámica en interfaces de redes públicas.

**TABLA 7-4** Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas

Tarea	Instrucciones
1. Conmutar las direcciones IP del adaptador que se va a eliminar a otro adaptador del grupo de varias rutas mediante <code>if_mpadm</code>	Página de comando <code>man if_mpadm(1M)</code> . <a href="#">Parte VI, “IPMP” de Guía de administración del sistema: servicios IP</a>
2. Eliminar el adaptador del grupo de varias rutas mediante el comando <code>ifconfig</code>	Página de comando <code>man ifconfig(1M)</code> <a href="#">Parte VI, “IPMP” de Guía de administración del sistema: servicios IP</a>
3. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	<a href="#">Sun Enterprise 10000 DR Configuration Guide</a> y <a href="#">Sun Enterprise 10000 Dynamic Reconfiguration Reference Manual</a> (de la recopilación <i>Solaris 10 on Sun hardware</i> ).



## Adición y eliminación de un nodo

---

Este capítulo proporciona instrucciones sobre cómo agregar o eliminar un nodo de un clúster:

- “Adición de nodos a un clúster” en la página 243
- “Eliminación de nodos de un clúster” en la página 249

Para obtener información acerca de las tareas de mantenimiento del clúster, consulte el [Capítulo 9, “Administración del clúster”](#).

### Adición de nodos a un clúster

Esta sección describe cómo agregar un nodo a un clúster global o a un clúster de zona. Puede crear un nuevo nodo del clúster de zona sobre un nodo del clúster global que aloje como host el clúster de zona, siempre y cuando el nodo del clúster global no aloje ya un nodo de ese clúster de zona en concreto. No es posible convertir en un nodo de clúster de zona un nodo sin votación ya existente y que reside en un clúster global.

En este capítulo, `phys - schost#` refleja una solicitud de clúster global. El indicador de solicitud de shell interactivo de `clzonecluster` es `clzc: schost>`.

La tabla siguiente muestra una lista con las tareas que se deben realizar para agregar un nodo a un clúster ya existente. Efectúe las tareas en el orden en que se muestran.

**TABLA 8-1** Mapa de tareas: agregar un nodo a un clúster global o a un clúster de zona existente

Tarea	Instrucciones
Instalar el adaptador de host en el nodo y comprobar que las interconexiones ya existentes del clúster sean compatibles con el nuevo nodo	<i>Oracle Solaris Cluster 3.3 Hardware Administration Manual</i>
Agregar almacenamiento compartido	<i>Oracle Solaris Cluster 3.3 Hardware Administration Manual</i>

TABLA 8-1 Mapa de tareas: agregar un nodo a un clúster global o a un clúster de zona existente (Continuación)

Tarea	Instrucciones
Agregar el nodo a la lista de nodos autorizados con <code>clsetup</code>	<a href="#">“Adición de un nodo a la lista de nodos autorizados” en la página 244</a>
Instalar y configurar el software del nuevo nodo del clúster	Capítulo 2, “Instalación del software en los nodos del clúster global” de <i>Guía de instalación del software Oracle Solaris Cluster</i>
Agregar el nuevo nodo a un clúster existente	<a href="#">“Adición de nodos a un clúster” en la página 243</a>
Si el clúster se configura en asociación con Oracle Solaris Cluster Geographic Edition, configure el nuevo nodo como participante activo en la configuración.	<a href="#">“How to Add a New Node to a Cluster in a Partnership” de Oracle Solaris Cluster Geographic Edition System Administration Guide</a>

▼ Adición de un nodo a la lista de nodos autorizados

Antes de agregar un host de Oracle Solaris o una máquina virtual a un clúster global o a uno de zona que ya existe, asegúrese de que el nodo disponga de todos los elementos de hardware necesarios correctamente instalados y configurados, lo que incluye una conexión física operativa con la interconexión privada del clúster.

Para obtener información sobre la instalación de hardware, consulte [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#) o la documentación de hardware proporcionada con el servidor.

Este procedimiento habilita un equipo para que se instale a sí mismo en un clúster al agregar su nombre de nodo a la lista de nodos autorizados en ese clúster.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 En el caso de un miembro de clúster global actual, conviértase en superusuario del miembro de clúster actual. Lleve a cabo estos pasos desde un nodo de un clúster global.
- 2 Compruebe que haya finalizado correctamente todas las tareas de configuración e instalación de hardware que figuran como requisitos previos en una lista en el mapa de tareas de la [Tabla 8-1](#).

**3 Inicie la utilidad clsetup.**

```
phys-schost# clsetup
```

Aparece el menú principal.

---

**Nota** – Para agregar un nodo a un clúster de zona, use la utilidad `clzonecluster`. Consulte el Paso 9 para ver instrucciones sobre cómo agregar manualmente una zona a un clúster de zona.

---

**4 Escriba el número correspondiente a la opción de que se muestre el menú Nuevos nodos y pulse la tecla Intro.****5 Escriba el número correspondiente a la opción de modificar la lista autorizada y pulse la tecla Intro. Especifique el nombre de un equipo que pueda agregarse a sí mismo.**

Siga los indicadores para agregar el nombre del nodo al clúster. Se le solicita el nombre del nodo que se va a agregar.

**6 Verifique que la tarea se haya realizado correctamente.**

La utilidad `clsetup` imprime un mensaje del tipo "Comando finalizado correctamente" si consigue concluir la tarea sin errores.

**7 Para impedir que se agreguen equipos nuevos al clúster, escriba el número correspondiente a la opción que ordena al clúster que omita las solicitudes para agregar más equipos. Pulse la tecla Intro.**

Siga los indicadores de `clsetup`. Esta opción indica al clúster que omita todas las solicitudes llegadas a través de la red pública procedentes de cualquier equipo nuevo que intente agregarse a sí mismo al clúster.

**8 Cierre la utilidad clsetup.****9 Para agregar manualmente un nodo a un clúster de zona, debe especificar el host de Oracle Solaris y el nombre del nodo virtual. También debe especificar un recurso de red que se utilizará para la comunicación con red pública en cada nodo. En el ejemplo siguiente, el nombre de zona es `sczone` y `bge0` es el adaptador de red pública en ambos equipos.**

```
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-1
clzc:sczone:node>set hostname=hostname1
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname1
clzc:sczone:node:net>set physical=bge0
clzc:sczone:node:net>end
clzc:sczone:node>end
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-2
clzc:sczone:node>set hostname=hostname2
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname2
```

```
clzc:sczone:node:net>set physical=bge0
clzc:sczone:node:net>end
clzc:sczone:node>end
```

Para obtener instrucciones detalladas sobre cómo configurar el nodo, consulte [“Configuración de un clúster de zona” de Guía de instalación del software Oracle Solaris Cluster](#).

## 10 Instale y configure el software en el nuevo nodo del clúster.

Puede utilizar `cluster create` o software JumpStart para completar la instalación y la configuración del nuevo nodo, tal y como se describe en [Guía de instalación del software Oracle Solaris Cluster](#).

### Ejemplo 8-1 Adición de nodos de un clúster global a la lista de nodos autorizados

El ejemplo siguiente muestra cómo se agrega un nodo denominado `phys-schost-3` a la lista de nodos autorizados de un clúster ya existente.

```
[Become superuser and execute the clsetup utility.]
phys-schost# clsetup
[Select New nodes>Specify the name of a machine which may add itself.]
[Answer the questions when prompted.]
[Verify that the command completed successfully.]

claccess allow -h phys-schost-3

Command completed successfully.
[Select Prevent any new machines from being added to the cluster.]
[Quit the clsetup New Nodes Menu and Main Menu.]
[Install the cluster software.]
```

### Véase también [clsetup\(1CL\)](#)

Para obtener una lista completa de las tareas necesarias para agregar un nodo de clúster, consulte la [Tabla 8-1](#), "Mapa de tareas: agregar un nodo del clúster".

Para agregar un nodo a un grupo de recursos ya existente, consulte [Oracle Solaris Cluster Data Services Planning and Administration Guide](#).

## Creación de un nodo sin votación (zona) en un clúster global

Esta sección ofrece los procedimientos siguientes e información para crear un nodo sin votación, denominado *zona*, en un nodo de clúster global.

## ▼ Creación de un nodo sin votación en un clúster global

- 1 **Conviértase en superusuario en el nodo del clúster global en el que vaya a crear un nodo que no sea de votación.**

Debe encontrarse en la zona global.

- 2 **Compruebe en todos los nodos que los servicios multiusuario para la Utilidad de gestión de servicios (SMF) estén en línea.**

Si los servicios todavía no están en línea para un nodo, espere hasta que cambie el estado y aparezca como en línea antes de continuar con el siguiente paso.

```
phys-schost# svcs multi-user-server node
STATE STIME FMRI
online 17:52:55 svc:/milestone/multi-user-server:default
```

- 3 **Configure, instale e inicie la nueva zona.**

---

**Nota** – Debe establecer la propiedad autoboot en true (verdadero) para admitir el uso de las funciones de grupos de recursos en el nodo del clúster global que no sea de votación.

---

Siga los procedimientos descritos en la documentación de Solaris:

- a. **Realice los procedimientos descritos en Capítulo 18, “Planificación y configuración de zonas no globales (tareas)” de Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris.**
  - b. **Realice los procedimientos descritos en la sección “Instalación e inicio de zonas” de Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris.**
  - c. **Realice los procedimientos descritos en la sección “Cómo iniciar una zona” de Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris.**
- 4 **Asegúrese de que la zona presente el estado ready (listo).**

```
phys-schost# zoneadm list -v
ID NAME STATUS PATH
0 global running /
1 my-zone ready /zone-path
```

- 5 **(Opcional) En una zona de direcciones IP compartidas, asigne una dirección IP y un nombre de host privado a la zona.**

El siguiente comando permite seleccionar y asignar una dirección IP disponible del intervalo de direcciones IP privadas del clúster. También permite asignar el nombre de host privado (o alias de host) especificado a la zona y a la dirección IP privada indicada.

```
phys-schost# clnode set -p zprivatehostname=hostalias node:zone
```

<code>-p</code>	Especifica una propiedad.
<code>zprivatehostname=<i>alias_host</i></code>	Especifica el nombre de host privado o el alias de host de la zona.
<code>nodo</code>	El nombre del nodo.
<code>zona</code>	El nombre del nodo del clúster global que no es de votación.

## 6 Realice una configuración inicial de zonas internas.

Siga los procedimientos descritos en la sección “[Configuración inicial de la zona interna](#)” de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*. Seleccione uno de los siguientes métodos:

- Inicie una sesión en la zona.
- Utilice el archivo `/etc/sysidcfg`.

## 7 En el nodo que no es de votación, modifique el archivo `nsswitch.conf`.

Estos cambios permiten a la zona realizar búsquedas de direcciones IP y nombres de host específicos del clúster.

### a. Inicie una sesión en la zona.

```
phys-schost# zlogin -c zonename
```

### b. Abra el archivo `/etc/nsswitch.conf` para modificarlo.

```
sczone# vi /etc/nsswitch.conf
```

### c. Agregue el conmutador `cluster` al comienzo de las búsquedas de las entradas `hosts` y `netmasks` y, a continuación, el conmutador `files`.

Las entradas modificadas deberían presentar un aspecto similar al siguiente:

```
...
hosts: cluster files nis [NOTFOUND=return]
...
netmasks: cluster files nis [NOTFOUND=return]
...
```

### d. En las entradas restantes, asegúrese de que el conmutador `files` es el primero que se muestra en la entrada.

### e. Salga de la zona.

## 8 Si ha creado una zona de direcciones IP exclusivas, configure los grupos IPMP en cada archivo `/etc/hostname.interfaz` que se encuentre en esa zona.

Debe configurar un grupo de IPMP para cada adaptador de red pública que se utilice para el tráfico de servicios de datos en la zona. Esta información no se hereda de la zona global.



Consulte la sección “Redes públicas” de *Guía de instalación del software Oracle Solaris Cluster* para obtener más información sobre cómo configurar grupos IPMP en un clúster.

- 9 Configure las asignaciones de nombre y dirección para todos los nombres de host lógicos que utilice la zona.
  - a. Agregue las asignaciones de nombre y dirección al archivo `/etc/inet/hosts` en la zona.

Esta información no se hereda de la zona global.
  - b. Si utiliza un servidor de nombres, agregue la asignación de nombre y dirección.

## Eliminación de nodos de un clúster

Esta sección ofrece instrucciones sobre cómo eliminar un nodo de un clúster global o de un clúster de zona. También es posible eliminar un clúster de zona específico de un clúster global. La tabla siguiente muestra una lista con las tareas que se deben realizar para eliminar un nodo de un clúster ya existente. Efectúe las tareas en el orden en que se muestran.



**Precaución** – Si se elimina un nodo aplicando sólo este procedimiento para una configuración de RAC, el nodo podría experimentar un error grave al rearrancar. Para obtener instrucciones sobre cómo eliminar un nodo de una configuración de RAC, consulte “[Cómo eliminar Admisión de Oracle RAC de los nodos seleccionados](#)” de *Guía del servicio de datos de Oracle Solaris Cluster para Oracle Real Application Clusters (RAC)*. Después de finalizar el proceso, siga los pasos detallados a continuación que resulten pertinentes.

TABLA 8-2 Mapa de tareas: eliminar un nodo

Tarea	Instrucciones
Mover todos los grupos de recursos y grupos de dispositivos fuera del nodo que se va a eliminar	<code>clnode evacuate nodo</code>
Comprobar que sea posible eliminar el nodo mediante la verificación de los hosts permitidos	<code>claccess show nodo</code> <code>claccess allow -h nodo_que_eliminar</code>
Si no es posible eliminar el nodo, conceder acceso a la configuración del clúster	
Eliminar el nodo de todos los grupos de dispositivos	“ <a href="#">Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)</a> ” en la página 139

TABLA 8-2 Mapa de tareas: eliminar un nodo (Continuación)

Tarea	Instrucciones
Quitar todos los dispositivos de quórum conectados al nodo que se va a eliminar	<b>Este paso es opcional en caso de que se elimine un nodo de un clúster compuesto por dos nodos.</b>   <a href="#">“Eliminación de un dispositivo de quórum” en la página 204</a>   Aunque hay que quitar el dispositivo de quórum antes de eliminar el dispositivo de almacenamiento en el paso siguiente, inmediatamente después se puede volver a agregar el dispositivo de quórum.   <a href="#">“Eliminación del último dispositivo de quórum de un clúster” en la página 205</a>
Poner el nodo que se vaya a eliminar en un modo que no sea de clúster	<a href="#">“Colocación de un nodo en estado de mantenimiento” en la página 274</a>
Eliminar un nodo de un clúster de zona	<a href="#">“Eliminación de un nodo de un clúster de zona” en la página 250</a>
Eliminar un nodo de la configuración de software del clúster	<a href="#">“Eliminación de un nodo de la configuración de software del clúster” en la página 251</a>
(Opcional) Desinstalar el software Oracle Solaris Cluster de un nodo del clúster	<a href="#">“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287</a>
Eliminar todo un clúster de zona	<a href="#">“Eliminación de un clúster de zona” en la página 282</a>

## ▼ Eliminación de un nodo de un clúster de zona

Se puede eliminar un nodo de un clúster de zona si se detiene el nodo, se desinstala y se elimina de la configuración. Si posteriormente decidiese volver a agregar el nodo al clúster de zona, siga las instrucciones que se indican en la [Tabla 8-1](#). La mayoría de estos pasos se realizan desde el nodo del clúster global.

- 1 **Conviértase en superusuario en un nodo del clúster global.**
- 2 **Cierre el nodo del clúster de zona que desee eliminar; para ello, especifique el nodo y su clúster de zona.**  

```
phys-schost# clzonecluster halt -n node zoneclustername
```

También puede utilizar los comandos `clnode evacuate` y `shutdown` dentro del clúster de zona.
- 3 **Desinstale el nodo del clúster de zona.**  

```
phys-schost# clzonecluster uninstall -n node zoneclustername
```

#### 4 Elimine el nodo del clúster de zona de la configuración.

Use los comandos siguientes:

```
phys-schost# clzonecluster configure zoneclustername
```

```
clzc:sczone> remove node physical-host=zoneclusternodename
```

#### 5 Compruebe que el nodo se haya eliminado del clúster de zona.

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

## ▼ Eliminación de un nodo de la configuración de software del clúster

Siga este procedimiento para eliminar un nodo del clúster global.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que haya eliminado el nodo de todos los grupos de recursos, grupos de dispositivos y configuraciones de dispositivo de quórum, y establézcalo en estado de mantenimiento antes de seguir adelante con este procedimiento.
- 2 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo que desee eliminar. Siga todos los pasos de este procedimiento desde un nodo del clúster global.
- 3 Inicie el nodo del clúster global que desee eliminar en un modo que no sea el de clúster. Para un nodo de un clúster de zona, siga las instrucciones indicadas en [“Eliminación de un nodo de un clúster de zona” en la página 250](#) antes de realizar este paso.
  - En los sistemas basados en SPARC, ejecute el comando siguiente.  

```
ok boot -x
```
  - En los sistemas basados en x86, ejecute los comandos siguientes.

```
shutdown -g -y -i0
```

Press any key to continue

- a. En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba e para editar los comandos.

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

Para obtener más información sobre el arranque basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- b. En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba e para editarla.

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.
```

- c. Agregue -x al comando para especificar que el sistema arranque en un modo que no sea de clúster.

```
[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/multiboot -x
```

- d. Pulse la tecla Intro para aceptar el cambio y volver a la pantalla de los parámetros de arranque.

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -x |
+-----+
```

```
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

**e. Escriba `b` para arrancar el nodo en un modo que no sea de clúster.**

Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en un modo que no sea de clúster, realice estos pasos para agregar la opción `-x` al comando del parámetro de arranque del núcleo.

---

**Nota** – Si el nodo que se vaya a eliminar no está disponible o ya no se puede arrancar, ejecute el comando siguiente en cualquier nodo del clúster activo: **`clnode clear -F <nodo_que_eliminar>`**. Verifique la eliminación del nodo mediante la ejecución de **`clnode status <nombre_nodo>`**.

---

**4 Desde el nodo que desea eliminar, quite el nodo del clúster.**

```
phys-schost# clnode remove -F
```

Si el comando **`clnode remove`** no funciona correctamente y existe una referencia de nodo caducada, ejecute **`clnode clear -F nombre_nodo`** en un nodo activo.

---

**Nota** – Si va a eliminar el último nodo del clúster, éste debe establecerse en un modo que no sea de clúster de forma que no quede ningún nodo activo en el clúster.

---

**5 Compruebe la eliminación del nodo desde otro nodo del clúster.**

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

**6 Finalice la eliminación del nodo.**

- Si tiene la intención de desinstalar el software Oracle Solaris Cluster del nodo eliminado, continúe con [“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287](#).
- Si no tiene previsto desinstalar el software Oracle Solaris Cluster del nodo eliminado, puede quitar físicamente el nodo del clúster mediante la extracción de las conexiones de hardware, como se describe en [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#).

**Ejemplo 8-2** Eliminación de un nodo de la configuración de software del clúster

Este ejemplo muestra cómo eliminar un nodo (phys-schost-2) desde un clúster. El comando `clnode remove` se ejecuta en un modo que no sea de clúster desde el nodo que desea eliminar del clúster (phys-schost-2).

```
[Remove the node from the cluster:]
phys-schost-2# clnode remove
phys-schost-1# clnode clear -F phys-schost-2
[Verify node removal:]
phys-schost-1# clnode status
-- Cluster Nodes --
 Node name Status
 -
Cluster node: phys-schost-1 Online
```

**Véase también** Para desinstalar el software Oracle Solaris Cluster del nodo eliminado, consulte [“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287.](#)

Para obtener información sobre los procedimientos de hardware, consulte [Oracle Solaris Cluster 3.3 Hardware Administration Manual](#).

Para obtener una lista completa de las tareas necesarias para eliminar un nodo del clúster, consulte la [Tabla 8-2](#).

Para agregar un nodo a un clúster ya existente, consulte [“Adición de un nodo a la lista de nodos autorizados” en la página 244.](#)

▼ **Eliminación de un nodo sin votación (zona) de un clúster global**

- 1 Conviértase en superusuario en el nodo del clúster global donde haya creado el nodo sin votación.
- 2 Elimine el nodo sin votación del sistema.

Siga los procedimientos de [“Eliminación de una zona no global del sistema” de Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris.](#)

## ▼ Eliminación de la conectividad entre una matriz y un único nodo, en un clúster con conectividad superior a dos nodos

Use este procedimiento para desconectar una matriz de almacenamiento de un nodo único de un clúster, en un clúster que tenga conectividad de tres o cuatro nodos.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Realice copias de seguridad de todas las tablas de bases de datos, los servicios de datos y los volúmenes que estén asociados con la matriz de almacenamiento que vaya a eliminar.**
- 2 **Determine los grupos de recursos y grupos de dispositivos que se estén ejecutando en el nodo que se va a desconectar.**

```
phys-schost# clresourcegroup status
phys-schost# cldevicegroup status
```

- 3 **Si es necesario, traslade todos los grupos de recursos y grupos de dispositivos fuera del nodo que se vaya a desconectar.**



**Caution (SPARC only)** – Si el clúster ejecuta software de Oracle RAC, cierre la instancia de base de datos Oracle RAC que esté en ejecución en el nodo antes de mover los grupos y sacarlos fuera del nodo. Si desea obtener instrucciones, consulte *Oracle Database Administration Guide*.

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

El comando `clnode evacuate` conmuta todos los grupos de dispositivos desde el nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos desde los nodos de votación o sin votación del nodo especificado hasta el siguiente nodo por orden de preferencia, sea de votación o sin votación.

- 4 **Establezca los grupos de dispositivos en estado de mantenimiento.**

Para el procedimiento destinado a consentir la actividad E/S en los grupos de discos compartidos Veritas, consulte la documentación de VxVM.

Para el procedimiento de establecer un grupo de dispositivos en el estado de mantenimiento, consulte [“Colocación de un nodo en estado de mantenimiento” en la página 274](#).

**5 Elimine el nodo de los grupos de dispositivos.**

Si usa VxVM o un disco básico, utilice el comando `cldevicegroup(1CL)` para eliminar los grupos de dispositivos.

**6 Para cada grupo de recursos que contenga un recurso HASStoragePlus, elimine el nodo de la lista de nodos del grupo de recursos.**

```
phys-schost# clresourcegroup remove-node -z zone -n node + | resourcegroup
nodo El nombre del nodo.
```

zona El nombre del nodo sin votación que puede controlar el grupo de recursos. Indique la zona sólo si ha especificado un nodo sin votación al crear el grupo de recursos.

Consulte *Oracle Solaris Cluster Data Services Planning and Administration Guide* si desea más información sobre cómo cambiar la lista de nodos de un grupo de recursos.

---

**Nota** – Los nombres de los tipos, los grupos y las propiedades de recursos distinguen mayúsculas y minúsculas cuando se ejecuta `clresourcegroup`.

---

**7 Si la matriz de almacenamiento que va a eliminar es la última que está conectada al nodo, desconecte el cable de fibra óptica que comunica el nodo y el concentrador o conmutador que está conectado a esta matriz de almacenamiento. Si no es así, omita este paso.****8 Si va a quitar el adaptador de host del nodo que va a desconectar, cierre el nodo. Si va a eliminar el adaptador de host del nodo que vaya a desconectar, pase directamente al Paso 11.****9 Elimine el adaptador de host del nodo.**

Para el procedimiento de eliminar los adaptadores de host, consulte la documentación relativa al nodo.

**10 Encienda el nodo pero no lo haga arrancar.****11 Si se ha instalado software Oracle RAC, elimine el paquete de software Oracle RAC del nodo que va a desconectar.**

```
phys-schost# pkgrm SUNWscuclm
```



---

**Caution (SPARC only)** – Si no elimina el software Oracle RAC del nodo que desconectó, dicho nodo experimentará un error grave al volver a introducirlo en el clúster, lo que podría provocar una pérdida de la disponibilidad de los datos.

---

**12 Arranque el nodo en modo de clúster.**

- En los sistemas basados en SPARC, ejecute el comando siguiente.



ok **boot**

- En los sistemas basados en x86, ejecute los comandos siguientes.

Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

- 13 En el nodo, actualice el espacio de nombre del dispositivo mediante la actualización de las entradas `/devices` y `/dev`.**

```
phys-schost# devfsadm -C
cldevice refresh
```

- 14 Vuelva a poner en línea los grupos de dispositivos.**

Para obtener información sobre los procedimientos relativos a poner en línea un grupo de discos compartidos Veritas, consulte la documentación del administrador de volúmenes Veritas Volume Manager.

Si desea obtener información sobre cómo poner en línea un grupo de dispositivos, consulte [“Procedimiento para sacar un nodo del estado de mantenimiento” en la página 276](#).

## ▼ Corrección de mensajes de error

Para corregir cualquier mensaje de error que aparezca al intentar llevar a cabo alguno de los procedimientos de eliminación de nodos del clúster, siga el procedimiento detallado a continuación.

- 1 Intente volver a unir el nodo al clúster global. Realice este procedimiento sólo en un clúster global.**

```
phys-schost# boot
```

- 2 ¿Se ha vuelto a unir correctamente el nodo al clúster?**

- Si no es así, continúe en el [Paso b](#).

- Si ha sido posible, efectúe los pasos siguientes con el fin de eliminar el nodo de los grupos de dispositivos.
  - a. Si el nodo vuelve a unirse correctamente al clúster, elimine el nodo del grupo o los grupos de dispositivos restantes.

Siga los procedimientos descritos en [“Eliminación de un nodo de todos los grupos de dispositivos” en la página 138.](#)
  - b. Tras eliminar el nodo de todos los grupos de dispositivos, vuelva a [“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287](#) y repita el procedimiento.
- 3 Si no se ha podido volver a unir el nodo al clúster, cambie el nombre del archivo `/etc/cluster/ccr` del nodo por otro cualquiera, por ejemplo `ccr.old`.

```
mv /etc/cluster/ccr /etc/cluster/ccr.old
```
- 4 Vuelva a [“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287](#) y repita el procedimiento.

# Administración del clúster

---

Este capítulo presenta los procedimientos administrativos que afectan a todo un clúster global o a un clúster de zona:

- “Información general sobre la administración del clúster” en la página 259
- “Ejecución de tareas administrativas del clúster de zona” en la página 281
- “Solución de problemas” en la página 297

Si desea información sobre cómo agregar o quitar un nodo del clúster, consulte el [Capítulo 8](#), “Adición y eliminación de un nodo”.

## Información general sobre la administración del clúster

Esta sección describe cómo realizar tareas administrativas para todo el clúster global o de zona. La tabla siguiente enumera estas tareas administrativas y los procedimientos asociados. Las tareas administrativas del clúster suelen efectuarse en la zona global. Para administrar un clúster de zona, al menos un equipo que almacenará como host el clúster de zona debe estar activado en modo de clúster. No todos los nodos clúster de zona deben estar activados y en funcionamiento; Oracle Solaris Cluster reproduce cualquier cambio de configuración si el nodo que está fuera del clúster se vuelve a unir a éste.

---

**Nota** – De manera predeterminada, la administración de energía está deshabilitada para que no interfiera con el clúster. Si habilita la administración de energía en un clúster de un solo nodo, el clúster continuará ejecutándose pero podría no estar disponible durante unos segundos. La función de administración de energía intenta cerrar el nodo, pero no lo consigue.

---

En este capítulo, `phys - schost#` refleja una solicitud de clúster global. El indicador de solicitud de shell interactivo de `clzonecluster` es `clzc: schost>`.

TABLA 9-1 Lista de tareas: administrar el clúster

Tarea	Instrucciones
Agregar o eliminar un nodo de un clúster	<a href="#">Capítulo 8, “Adición y eliminación de un nodo”</a>
Cambiar el nombre del clúster	<a href="#">“Cambio del nombre del clúster” en la página 260</a>
Enumerar los ID de nodos y sus nombres de nodo correspondientes	<a href="#">“Asignación de un ID de nodo a un nombre de nodo” en la página 262</a>
Permitir o denegar que se agreguen nodos nuevos al clúster	<a href="#">“Uso de la autenticación del nodo del clúster nuevo” en la página 262</a>
Cambiar la hora de un clúster mediante el protocolo NTP	<a href="#">“Restablecimiento de la hora del día en un clúster” en la página 264</a>
Cerrar un nodo con el indicador ok de la PROM OpenBoot en un sistema basado en SPARC o con el mensaje Press any key to continue de un menú de GRUB en un sistema basado en x86	<a href="#">“SPARC: Visualización de la PROM OpenBoot en un nodo” en la página 266</a>
Agregar o cambiar el nombre de host privado	<a href="#">“Adición un nombre de host privado a un nodo sin votación en un clúster global” en la página 270</a> <a href="#">“Cambio del nombre de host privado de nodo” en la página 267</a>
Poner un nodo de clúster en estado de mantenimiento	<a href="#">“Colocación de un nodo en estado de mantenimiento” en la página 274</a>
Cambiar el nombre de un nodo	<a href="#">“Cómo cambiar el nombre de un nodo” en la página 272</a>
Sacar un nodo del clúster fuera del estado de mantenimiento	<a href="#">“Procedimiento para sacar un nodo del estado de mantenimiento” en la página 276</a>
Configurar los límites de carga de un nodo	<a href="#">“Cómo configurar límites de carga en un nodo” en la página 279</a>
Mover un clúster de zona, preparar un clúster de zona para aplicaciones, eliminar un clúster de zona	<a href="#">“Ejecución de tareas administrativas del clúster de zona” en la página 281</a>
Desinstalar el software Oracle Solaris Cluster desde un nodo	<a href="#">“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287</a>
Agregar y administrar una MIB de eventos de SNMP	<a href="#">“Habilitación de una MIB de eventos de SNMP” en la página 292</a> <a href="#">“Adición de un usuario de SNMP a un nodo” en la página 295</a>

## ▼ Cambio del nombre del clúster

Si es necesario, el nombre del clúster puede cambiarse tras la instalación inicial.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario en un nodo del clúster global.**
- 2 Inicie la utilidad `clsetup`.**  
`phys-schost# clsetup`  
Aparece el menú principal.
- 3 Para cambiar el nombre del clúster, escriba el número correspondiente a la opción Other Cluster Properties (Otras propiedades del clúster).**  
Aparece el menú Other Cluster Properties (Otras propiedades del clúster).
- 4 Seleccione desde el menú y siga las instrucciones que aparecen en la pantalla.**
- 5 Si desea que la etiqueta de servicio de Oracle Solaris Cluster refleje el nombre del clúster nuevo, elimine la etiqueta de Oracle Solaris Cluster y reinicie el clúster. Para eliminar la instancia de la etiqueta de servicio de Oracle Solaris Cluster, realice los subpasos siguientes en todos los nodos del clúster.**
  - a. Enumere todas las etiquetas de servicio.**  
`phys-schost# stclient -x`
  - b. Busque el número de instancia de la etiqueta de servicio de Oracle Solaris Cluster; a continuación, ejecute el comando siguiente.**  
`phys-schost# stclient -d -i service_tag_instance_number`
  - c. Rearranque todos los nodos del clúster.**  
`phys-schost# reboot`

### Ejemplo 9-1 Cambio del nombre del clúster

El ejemplo siguiente muestra el comando `cluster(1CL)` generado desde la utilidad `clsetup(1CL)` para cambiar al nombre nuevo del clúster, `dromedary`.

```
phys-schost# cluster -c dromedary
```

## ▼ Asignación de un ID de nodo a un nombre de nodo

Al instalar Oracle Solaris Cluster, de forma automática se asigna un número exclusivo de ID de nodo a todos los nodos. El número de ID de nodo se asigna a un nodo en el orden en que se une al clúster por primera vez. Una vez asignado, no se puede cambiar. El número de ID de nodo suele usarse en mensajes de error para identificar el nodo del clúster al que hace referencia el mensaje. Siga este procedimiento para determinar la asignación entre los ID y los nombres de los nodos.

Para visualizar la información sobre la configuración para un clúster global o de zona no hace falta ser superusuario. Un paso de este procedimiento se realiza desde un nodo de clúster global. El otro paso se efectúa desde un nodo de clúster de zona.

- 1 Use el comando **clnode(1CL)** para mostrar la información sobre configuración del clúster para el clúster global.

```
phys-schost# clnode show | grep Node
```

- 2 También puede mostrar los ID de nodo para un clúster de zona. El nodo de clúster de zona tiene el mismo ID que el nodo de clúster global donde se ejecuta.

```
phys-schost# zlogin szzone clnode -v | grep Node
```

### Ejemplo 9–2 Asignación de ID del nodo al nombre del nodo

El ejemplo siguiente muestra las asignaciones de ID de nodo para un clúster global.

```
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
```

## ▼ Uso de la autenticación del nodo del clúster nuevo

Oracle Solaris Cluster permite determinar si se pueden agregar nodos nuevos al clúster global y el tipo de autenticación que puede utilizar. Puede permitir que cualquier nodo nuevo se una al clúster por la red pública, denegar que los nodos nuevos se unan al clúster o indicar que un nodo en particular se una al clúster. Los nodos nuevos se pueden autenticar mediante la autenticación estándar de UNIX o Diffie-Hellman (DES). Si selecciona la autenticación DES, también debe configurar todas las pertinentes claves de cifrado antes de unir un nodo. Para obtener más información, consulte las páginas de comando man [keyserv\(1M\)](#) y [publickey\(4\)](#).

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario en un nodo del clúster global.**
- 2 Inicie la utilidad `clsetup(1CL)`.**  
`phys-schost# clsetup`  
 Aparece el menú principal.
- 3 Para trabajar con la autenticación del clúster, escriba el número correspondiente a la opción de nodos nuevos.**  
 Aparece el menú Nuevos nodos.
- 4 Seleccione desde el menú y siga las instrucciones que aparecen en la pantalla.**

#### **Ejemplo 9-3** Procedimiento para evitar que un equipo nuevo se agregue al clúster global

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que evita que equipos nuevos se agreguen al clúster.

```
phys-schost# claccess deny -h hostname
```

#### **Ejemplo 9-4** Procedimiento para permitir que todos los equipos nuevos se agreguen al clúster global

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que permite que todos los equipos nuevos se agreguen al clúster.

```
phys-schost# claccess allow-all
```

#### **Ejemplo 9-5** Procedimiento para especificar que se agregue un equipo nuevo al clúster global

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que permite que un solo equipo nuevo se agregue al clúster.

```
phys-schost# claccess allow -h hostname
```

#### **Ejemplo 9-6** Configuración de la autenticación en UNIX estándar

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que restablece la configuración de la autenticación UNIX estándar para los nodos nuevos que se unan al clúster.

```
phys-schost# claccess set -p protocol=sys
```

## Ejemplo 9-7 Configuración de la autenticación en DES

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que usa la autenticación DES para los nodos nuevos que se unan al clúster.

```
phys-schost# claccess set -p protocol=des
```

Cuando use la autenticación DES, también debe configurar todas las claves de cifrado necesarias antes de unir un nodo al clúster. Para obtener más información, consulte las páginas de comando `man keyserv(1M)` y `publickey(4)`.

## ▼ Restablecimiento de la hora del día en un clúster

El software Oracle Solaris Cluster usa el protocolo NTP para mantener la sincronización temporal entre los nodos del clúster. Los ajustes del clúster global se efectúa automáticamente según sea necesario cuando los nodos sincronizan su hora. Para obtener más información, consulte *Oracle Solaris Cluster Concepts Guide* y *Network Time Protocol User's Guide*.



**Precaución** – Si utiliza NTP, no intente ajustar la hora del clúster mientras se encuentre activo y en funcionamiento. No ajuste la hora mediante los comandos `date(1)`, `rdate(1M)`, `xntpd(1M)` ni `svcadm(1M)` de forma interactiva ni dentro de las secuencias de comandos `cron(1M)`.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario en un nodo del clúster global.**
- 2 **Cierre el clúster global.**  

```
phys-schost# cluster shutdown -g0 -y -i 0
```
- 3 **Compruebe que el nodo muestre el indicador ok en los sistemas basados en SPARC o el mensaje Press any key to continue en el menú de GRUB de los sistemas basados en x86.**
- 4 **Arranque el nodo en un modo que no sea de clúster.**
  - En los sistemas basados en SPARC, ejecute el comando siguiente.  

```
ok boot -x
```
  - En los sistemas basados en x86, ejecute los comandos siguientes.



```
shutdown -g -y -i0
```

Press any key to continue

- a. En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba **e** para editar los comandos.

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
```

```
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.  
Press enter to boot the selected OS, 'e' to edit the  
commands before booting, or 'c' for a command-line.

Para obtener más información sobre el arranque basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- b. En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba **e** para editarla.

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.  
Press 'b' to boot, 'e' to edit the selected command in the  
boot sequence, 'c' for a command-line, 'o' to open a new line  
after ('O' for before) the selected line, 'd' to remove the  
selected line, or escape to go back to the main menu.

- c. Agregue **-x** al comando para especificar que el sistema arranque en un modo que no sea de clúster.

```
[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/multiboot -x
```

- d. Pulse la tecla **Intro** para aceptar el cambio y volver a la pantalla de los parámetros de arranque.

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -x |
+-----+
```

```
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

**e. Escriba b para arrancar el nodo en un modo que no sea de clúster.**

---

**Nota** – Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en el modo que no es de clúster, realice estos pasos de nuevo para agregar la opción -x al comando del parámetro de arranque del núcleo.

---

**5 En un solo nodo, configure la hora del día; para ello, ejecute el comando date.**

```
phys-schost# date HHMM.SS
```

**6 En los otros equipos, sincronice la hora de ese nodo; para ello, ejecute el comando rdate(1M).**

```
phys-schost# rdate hostname
```

**7 Arranque cada nodo para reiniciar el clúster.**

```
phys-schost# reboot
```

**8 Compruebe que el cambio se haya hecho en todos los nodos del clúster.**

Ejecute el comando date en todos los nodos.

```
phys-schost# date
```

## ▼ SPARC: Visualización de la PROM OpenBoot en un nodo

Siga este procedimiento si debe configurar o cambiar la configuración de la PROM OpenBoot™.

phys-schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

**1 Conecte la consola al nodo para que se cierre.**

```
telnet tc_name tc_port_number
```

nombre\_ct Especifica el nombre del concentrador de terminales.

`tc_número_puerto` Especifica el número del puerto del concentrador de terminales. Los números de puerto dependen de la configuración. Los puertos 2 y 3 (5002 y 5003) suelen usarse para el primer clúster instalado en un sitio.

- 2 **Cierre el nodo del clúster mediante el comando `clnode evacuate` y, a continuación, mediante el comando `shutdown`. El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos del nodo especificado de votación o sin votación de clúster global por el siguiente nodo de votación o sin votación preferido.**

```
phys-schost# clnode evacuate node
shutdown -g0 -y
```




---

**Precaución** – No use `send brk` en una consola del clúster para cerrar el nodo del clúster.

---

- 3 **Ejecute los comandos de OBP.**

## ▼ Cambio del nombre de host privado de nodo

Use este procedimiento para cambiar el nombre de host privado de un clúster una vez finalizada la instalación.

Durante la instalación inicial del clúster se asignan nombre de host privados predeterminados. El nombre de host privado usa el formato `clusternode<id_nodo>-priv`; por ejemplo: `clusternode3-priv`. Cambie un nombre de host privado sólo si el nombre ya se utiliza en el dominio.




---

**Precaución** – No intente asignar direcciones IP a los nuevos nombres de host privados. Las asigna el software de agrupación en clúster.

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **En todos los nodos del clúster, inhabilite cualquier recurso de servicio de datos u otras aplicaciones que puedan almacenar en caché nombres de host privados.**

```
phys-schost# clresource disable resource[,...]
```

Incluya la información siguiente en las aplicaciones que inhabilite.

- Los servicios de HA-DNS y HA-NFS, si están configurados
- Cualquier aplicación configurada de forma personalizada para utilizar el nombre de host privado
- Cualquier aplicación que los clientes utilicen en la interconexión privada

Para obtener información sobre el uso del comando `clresource`, consulte la página de comando `man clresource(1CL)` y *Oracle Solaris Cluster Data Services Planning and Administration Guide*.

- 2 **Si el archivo de configuración NTP hace referencia al nombre de host privado que está cambiando, desactive el daemon del protocolo NTP en todos los nodos del clúster.**

Utilice el comando `svcadm` para detener el daemon del protocolo NTP. Consulte la página de comando `man svcadm(1M)` para obtener más información sobre el daemon del protocolo NTP.

```
phys-schost# svcadm disable ntp
```

- 3 **Ejecute la utilidad `clsetup(1CL)` para cambiar el nombre de host privado del nodo correspondiente.**

Ejecute la utilidad desde uno de los nodos del clúster.

---

**Nota** – Cuando seleccione un nombre nuevo para el sistema privado, compruebe que no se utilice en el nodo del clúster.

---

- 4 **Escriba el número correspondiente a la opción del nombre de host privado.**
- 5 **Escriba el número correspondiente a la opción de cambiar el nombre de host privado.**

Responda a las preguntas cuando se indique. Se le solicitará el nombre del nodo cuyo nombre de host privado desee cambiar (`clusternode< idnodo> -priv`), así como el nombre de host nuevo para el sistema privado.

- 6 **Purga la memoria caché del servicio de nombres.**

Efectúe este paso en todos los nodos del clúster. La purga evita que las aplicaciones del clúster y los servicios de datos intenten acceder al nombre de host privado anterior.

```
phys-schost# nscd -i hosts
```

- 7 Si ha cambiado un nombre de host privado en el archivo de configuración de NTP, actualice el archivo de configuración de NTP (`ntp.conf` o `ntp.conf.cluster`) en todos los nodos.
  - a. Use la herramienta de edición que prefiera.  
Si efectúa este paso durante la instalación, recuerde eliminar los nombres de los nodos que estén configurados. La plantilla predeterminada está preconfigurada con 16 nodos. En general, el archivo `ntp.conf.cluster` es el mismo en todos los nodos del clúster.
  - b. Compruebe que pueda realizar un ping correctamente en el nombre de host privado desde todos los nodos del clúster.
  - c. Reinicie el daemon de NTP.  
Efectúe este paso en todos los nodos del clúster.  
Utilice el comando `svcadm` para reiniciar el daemon de NTP.  

```
svcadm enable ntp
```

- 8 Habilite todos los recursos de servicio de datos y otras aplicaciones inhabilitados en el [Paso 1](#).

```
phys-schost# clresource enable resource[,...]
```

Para obtener información sobre el uso del comando `clresource`, consulte la página de comando `man clresource(1CL)` y *Oracle Solaris Cluster Data Services Planning and Administration Guide*.

### Ejemplo 9–8 Cambio del nombre de host privado

El ejemplo siguiente cambia el nombre de host privado de `clusternode2-priv` a `clusternode4-priv` en el nodo `phys-schost-2`.

```
[Disable all applications and data services as necessary.]
phys-schost-1# /etc/init.d/xntpd stop
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# vi /etc/inet/ntp.conf
...
peer clusternode1-priv
peer clusternode4-priv
peer clusternode3-priv
phys-schost-1# ping clusternode4-priv
phys-schost-1# /etc/init.d/xntpd start
[Enable all applications and data services disabled at the beginning of the procedure.]
```

## ▼ Adición un nombre de host privado a un nodo sin votación en un clúster global

Siga este procedimiento para agregar un nombre de host privado a un nodo sin votación en un clúster global una vez finalizada la instalación. En los procedimientos tratados en este capítulo, `phys - schost#` refleja una solicitud de clúster global. Realice este procedimiento sólo en un clúster global.

- 1 Ejecute la utilidad `clsetup(1CL)` para agregar un nombre de host privado a la zona adecuada.  
`phys - schost# clsetup`
- 2 Escriba el número correspondiente a la opción de nombres de host privados y pulse la tecla **Intro**.
- 3 Escriba el número correspondiente a la opción de agregar un nombre de host privado de zona y pulse la tecla **Intro**.  
Responda a las preguntas cuando se indique. No existe un nombre predeterminado para el sistema privado de un nodo sin votación del clúster global. Se debe proporcionar un nombre de host.

## ▼ Cambio del nombre de host privado en un nodo sin votación de un clúster global

Siga este procedimiento para cambiar el nombre de host privado de un nodo sin votación una vez finalizada la instalación.

Durante la instalación inicial del clúster se asignan nombre de host privados. El nombre de host privado usa el formato `clusternode< id_nodo>-priv`; por ejemplo: `clusternode3-priv`. Cambie un nombre de host privado sólo si el nombre ya se utiliza en el dominio.



---

**Precaución** – No intente asignar direcciones IP a los nuevos nombres de host privados. Las asigna el software de agrupación en clúster.

---

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **En todos los nodos del clúster global, inhabilite cualquier recurso de servicio de datos u otras aplicaciones que puedan almacenar en caché nombres de host privados.**

```
phys-schost# clresource disable resource1, resource2
```

Incluya la información siguiente en las aplicaciones que inhabilite.

- Los servicios de HA-DNS y HA-NFS, si están configurados
- Cualquier aplicación configurada de forma personalizada para utilizar el nombre de host privado
- Cualquier aplicación que los clientes utilicen en la interconexión privada

Para obtener información sobre el uso del comando `clresource`, consulte la página de comando `man clresource(1CL)` y *Oracle Solaris Cluster Data Services Planning and Administration Guide*.

- 2 **Ejecute la utilidad `clsetup(1CL)` para cambiar el nombre de host privado del nodo sin votación adecuado de clúster global.**

```
phys-schost# clsetup
```

Este paso debe hacerse desde un solo nodo del clúster.

---

**Nota** – Cuando seleccione un nombre nuevo para el sistema privado, compruebe que no se utilice en el clúster.

---

- 3 **Escriba el número correspondiente a la opción de nombres de host privados y pulse la tecla Intro.**
- 4 **Escriba el número correspondiente a la opción de agregar un nombre de host privado de zona y pulse la tecla Intro.**

No existe un nodo sin votación predeterminado del nombre de host privado de un clúster global. Se debe proporcionar un nombre de host.

- 5 **Escriba el número correspondiente a la opción de cambiar el nombre de host privado de zona.**

Responda a las preguntas cuando se indique. Se le solicitará el nombre del nodo sin votación cuyo nombre de host privado esté cambiando (`clusternode< id_nodo> -priv`), así como el nombre nuevo del sistema privado.

- 6 **Purgue la memoria caché del servicio de nombres.**

Efectúe este paso en todos los nodos del clúster. La purga evita que las aplicaciones del clúster y los servicios de datos intenten acceder al nombre de host privado anterior.

```
phys-schost# nscd -i hosts
```

- 7 **Habilite todos los recursos de servicio de datos y otras aplicaciones inhabilitados en el [Paso 1](#).**

## ▼ Eliminación del nombre de host privado para un nodo sin votación en un clúster global

Siga este procedimiento para eliminar un nombre de host privado para un nodo sin votación en un clúster global. Realice este procedimiento sólo en un clúster global.

- 1 Ejecute la utilidad `clsetup(1CL)` para eliminar un nombre de host privado en la zona adecuada.
- 2 Escriba el número correspondiente a la opción del nombre de host privado de zona.
- 3 Escriba el número correspondiente a la opción de eliminar un nombre de host privado de zona.
- 4 Escriba el nombre de host privado del nombre del nodo sin votación que va a eliminar.

## ▼ Cómo cambiar el nombre de un nodo

Puede cambiar el nombre de un nodo que forme parte de una configuración de Oracle Solaris Cluster. Debe cambiar el nombre del host de Oracle Solaris para poder cambiar el nombre del nodo. Utilice el `cnode rename` para cambiar el nombre del nodo.

Las instrucciones siguientes se aplican a cualquier aplicación que se esté ejecutando en un clúster global.

- 1 En el clúster global, conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.
- 2 Si va a cambiar el nombre de un nodo de un clúster de Oracle Solaris Cluster Geographic Edition que se encuentra en asociación con una configuración de Oracle Solaris, deberá llevar a cabo pasos adicionales. Si el clúster en el que está realizando el procedimiento de cambio de nombre es el clúster principal del grupo de protección y quiere que la aplicación del grupo de protección permanezca en línea, puede cambiar el grupo de protección al clúster secundario durante el procedimiento de cambio de nombre. Para obtener más información sobre los clústeres y nodos de Geographic Edition, consulte el [Capítulo 5, "Administering Cluster Partnerships" de Oracle Solaris Cluster Geographic Edition System Administration Guide](#).
- 3 Cambie los nombres de host de Oracle Solaris siguiendo los pasos indicados en ["How to Change a System's Host Name" de System Administration Guide: Advanced Administration](#), pero *no* reinicie el sistema al final del procedimiento. En su lugar, cierre el clúster una vez completados los pasos.
- 4 Inicie todos los nodos del clúster en un modo que no sea de clúster.

```
ok> boot -x
```



- 5 En un modo que no sea de clúster, en el nodo donde haya cambiado el nombre de host de Oracle Solaris, cambie el nombre del nodo y ejecute el comando `cmd` en cada uno de los host a los que haya cambiado el nombre. Cambie el nombre de los nodos uno por uno.  

```
clnode rename -n newnodename oldnodename
```
- 6 Actualice las referencias existentes al nombre de host anterior en las aplicaciones que se ejecuten en el clúster.
- 7 Compruebe que se haya cambiado el nombre del nodo consultando los mensajes de comando y los archivos de registro.
- 8 Reinicie todos los nodos en modo clúster.  

```
sync;sync;sync;/etc/reboot
```
- 9 Compruebe que el nodo muestre el nuevo nombre.  

```
clnode status -v
```
- 10 Si va a cambiar el nombre de un nodo de clúster de Geographic Edition y el clúster asociado del clúster que contiene el nodo al que se ha cambiado el nombre todavía hace referencia al nombre de nodo anterior, el estado de sincronización del grupo de protección se mostrará como un *error*. Debe actualizar el grupo de protección de un nodo del clúster asociado que contiene el nodo cuyo nombre se ha cambiado mediante el comando `geopg update <pg>`. Después de completar este paso, ejecute el comando `geopg start -e global <pg>`. Más adelante, puede volver a cambiar el grupo de protección al clúster que contiene el nodo cuyo nombre se ha cambiado.
- 11 Puede elegir si desea cambiar la propiedad `hostnameList` de recursos de nombre de host lógico. Consulte [“Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes” en la página 273](#) para obtener instrucciones sobre este paso opcional.

## ▼ Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes

Puede elegir si desea cambiar la propiedad `hostnameList` del recurso de nombre de host lógico antes o después de cambiar el nombre del nodo; para ello, siga los pasos descritos en [“Cómo cambiar el nombre de un nodo” en la página 272](#). Este paso es opcional.

- 1 En el clúster global, conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.

- 2 Si lo desea, puede cambiar los nombres de host lógicos utilizados por cualquiera de los recursos existentes de nombre de host lógico de Oracle Solaris Cluster.

En los pasos siguientes se describe cómo configurar el recurso `apache-lh-res` para que funcione con el nuevo nombre de host lógico. Se debe ejecutar en modo clúster.

- a. En el modo clúster, ponga fuera de línea los grupos de recursos de Apache que contengan los nombres de host lógicos.

```
clrg offline apache-rg
```

- b. Deshabilite los recursos de nombre de host lógico de Apache.

```
clrs disable apache-lh-res
```

- c. Proporcione la nueva lista de nombres de host.

```
clrs set -p HostnameList=test-2 apache-lh-res
```

- d. Cambie las referencias de la aplicación a entradas anteriores en la propiedad `hostnameList` para hacer referencia a las entradas nuevas.

- e. Habilitar los nuevos recursos de nombre de host lógico de Apache

```
clrs enable apache-lh-res
```

- f. Ponga en línea los grupos de recursos de Apache.

```
clrg online apache-rg
```

- g. Compruebe que la aplicación se haya iniciado correctamente; para ello, ejecute el comando siguiente para realizar la comprobación de un cliente.

```
clrs status apache-rs
```

## ▼ Colocación de un nodo en estado de mantenimiento

Ponga un nodo de clúster global en estado de mantenimiento si el nodo va a estar fuera de servicio durante mucho tiempo. De esta forma, el nodo no contribuye al número de quórum mientras sea objeto de tareas de mantenimiento o reparación. Para poner un nodo en estado de mantenimiento, éste tiene se debe cerrar con los comandos de evacuación `clnode(1CL)` y de cierre `cluster(1CL)`.

---

**Nota** – Use el comando `shutdown` de Oracle Solaris para cerrar un solo nodo. Use el comando `cluster shutdown` sólo cuando vaya a cerrar todo un clúster.

---

Cuando un clúster se cierra y se pone en estado de mantenimiento, el número de votos de quórum de todos los dispositivos de quórum configurados con puertos al nodo se reduce en

uno. Los números de votos del nodo y del dispositivo de quórum se incrementan en uno cuando el nodo sale del modo de mantenimiento y vuelve a estar en línea.

Use el comando de inhabilitación `clquorum(1CL)` para poner un nodo del clúster en estado de mantenimiento.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo del clúster global que va a poner en estado de mantenimiento.**
- 2 **Evacue del nodo todos los grupos de recursos y de dispositivos. El comando `clnode evacuate` conmuta todos los grupos de recursos y de dispositivos, incluidos todos los nodos de no votación, del nodo especificado al siguiente nodo por orden de preferencia.**

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

- 3 **Cierre el nodo que ha evacuado.**

```
phys-schost# shutdown -g0 -y-i 0
```

- 4 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en otro nodo del clúster y ponga el nodo cerrado en el Paso 3 en estado de mantenimiento.**

```
phys-schost# clquorum disable node
```

`node`                      Especifica el nombre de un nodo que desea poner en estado de mantenimiento.

- 5 **Compruebe que el nodo del clúster global esté en estado de mantenimiento.**

```
phys-schost# clquorum status node
```

El nodo puesto en estado de mantenimiento debe tener un Status de `offline` y `0` para los votos de quórum `Present` y `Possible`.

## Ejemplo 9-9 Colocación de un nodo del clúster global en estado de mantenimiento

En el ejemplo siguiente se pone un nodo del clúster en estado de mantenimiento y se comprueban los resultados. La salida de `clnode status` muestra los Node votes para que `phys-schost-1` sea `0` y el estado sea `Offline`. El Quorum Summary debe mostrar también el

número de votos reducido. Según la configuración, la salida de Quorum Votes by Device deber indicar que algunos dispositivos del disco también se encuentran fuera de línea.

```
[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

**Véase también** Para volver a poner en línea un nodo, consulte [“Procedimiento para sacar un nodo del estado de mantenimiento” en la página 276](#).

## ▼ Procedimiento para sacar un nodo del estado de mantenimiento

Complete el procedimiento descrito a continuación para volver a poner en línea un nodo del clúster global y restablecer el valor predeterminado en el número de votos de quórum. En los nodos del clúster, el número de quórum predeterminado es uno. En los dispositivos de quórum, el número de quórum predeterminado es  $N-1$ , donde  $N$  es el número de nodos con número de votos distinto de cero que tienen puertos conectados al dispositivo de quórum.

Si un nodo se pone en estado de mantenimiento, el número de votos de quórum se reduce en uno. Todos los dispositivos de quórum configurados con puertos conectados al nodo también ven reducido su número de votos. Al restablecer el número de votos de quórum y un nodo se quita del estado de mantenimiento, el número de votos de quórum y el del dispositivo de quórum se ven incrementados en uno.

Ejecute este procedimiento siempre que haya puesto el nodo del clúster global en estado de mantenimiento y vaya a sacarlo de él.



**Precaución** – Si no especifica ni la opción `globaldev` ni `node`, el recuento de quórum se restablece para todo el clúster.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC de `solaris.cluster.modify` en cualquier nodo del clúster global que no sea el que se está en estado de mantenimiento.**
- 2 **Según el número de nodos que tenga en la configuración del clúster global, realice uno de los pasos siguientes:**

- Si tiene dos nodos en la configuración del clúster, vaya al [Paso 4](#).
- Si tiene más de dos nodos en la configuración del clúster, vaya a [Paso 3](#).

- 3 **Si el nodo que va a quitar del estado de mantenimiento va a tener dispositivos de quórum, restablezca el número de votos de quórum del clúster desde un nodo que no se esté en estado de mantenimiento.**

El número de quórum debe restablecerse desde un nodo que no se esté en estado de mantenimiento antes de rearrancar el nodo; de lo contrario, el nodo puede bloquearse mientras espera el quórum.

```
phys-schost# clquorum reset
```

`reset` El indicador de cambio que restablece el quórum.

- 4 **Arranque el nodo que está quitando del estado de mantenimiento.**
- 5 **Compruebe el número de votos de quórum.**

```
phys-schost# clquorum status
```

El nodo que ha quitado del estado de mantenimiento deber tener el estado de `online` y mostrar el número de votos adecuado para los votos de quórum `Present` y `Possible`.

### **Ejemplo 9–10** Eliminación de un nodo del clúster del estado de mantenimiento y restablecimiento del número de votos de quórum

El ejemplo siguiente restablece el número de quórum para un nodo del clúster y sus dispositivos de quórum a la configuración predeterminada y comprueba el resultado. El resultado de `scstat -q` muestra los Node votes para que `phys-schost-1` sea 1 y que el estado sea `online`. La salida de Quorum Summary también debe mostrar un aumento en el número de votos.

```
phys-schost-2# clquorum reset
```

- En los sistemas basados en SPARC, ejecute el comando siguiente.  
ok **boot**
- En los sistemas basados en x86, ejecute los comandos siguientes.  
Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro. El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

```
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

## Configuración de los límites de carga

Puede habilitar la distribución automática de la carga del grupo de recursos entre los nodos o las zonas estableciendo límites de carga. Puede configurar un conjunto de límites de carga para cada nodo del clúster. Usted asigna factores de carga a los grupos de recursos y los factores de carga corresponden a los límites de carga definidos de los nodos. El funcionamiento predeterminado consiste en distribuir la carga del grupo de recursos de forma uniforme entre todos los nodos disponibles en la lista de nodos del grupo de recursos.

El RGM inicia los grupos de recursos en un nodo de la lista de nodos del grupo de recursos para que no se superen los límites de carga del nodo. Puesto que el RGM asigna grupos de recursos a los nodos, los factores de carga de los grupos de recursos de cada nodo se suman para proporcionar una carga total. La carga total se compara con los límites de carga del nodo.

Un límite de carga consta de los elementos siguientes:

- Un nombre asignado por el usuario.
- Un valor de límite flexible. Un límite de carga flexible se puede superar temporalmente.
- Un valor de límite fijo. Los límites de carga fijos no pueden superarse nunca y se aplican de forma estricta.

Puede definir tanto el límite fijo como el límite flexible con un solo comando. Si uno de los límites no se establece explícitamente, se utilizará el valor predeterminado. Los límites de carga fijos y flexibles de cada nodo se crean y modifican con los comandos `clnode create-loadlimit`, `clnode set-loadlimit` y `clnode delete-loadlimit`. Consulte la página de comando [man clnode\(1CL\)](#) para obtener más información.

Puede configurar un grupo de recursos para que tenga una mayor prioridad, de modo que sea menos probable que se desplace de un nodo específico. También puede establecer una propiedad `preemption_mode` para determinar si un grupo de recursos se apoderará de un nodo mediante un grupo de recursos de mayor prioridad debido a la sobrecarga de nodos. La propiedad `concentrate_load` también permite concentrar la carga del grupo de recursos en el menor número de nodos posible. El valor predeterminado de la propiedad `concentrate_load` es `FALSE`.

---

**Nota** – Puede configurar límites de carga en los nodos de un clúster global o de un clúster de zona. Puede utilizar la línea de comandos, la utilidad `clsetup` o la interfaz del Oracle Solaris Cluster Manager para configurar los límites de carga. En los procedimientos siguientes se explica cómo configurar límites de carga mediante la línea de comandos.

---

## ▼ Cómo configurar límites de carga en un nodo

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en cualquier nodo del clúster global.**
- 2 **Cree y configure un límite de carga para los nodos que quiera usar para el equilibrado de carga.**

```
clnode create-loadlimit -p limitname=mem_load -Z zc1 -p
softlimit=11 -p hardlimit=20 node1 node2 node3
```

En este ejemplo, el nombre del clúster de zona es `zc1`. La propiedad de ejemplo se llama `mem_load` y tiene un límite flexible de 11 y un límite de carga fijo de 20. Los límites fijos y

flexibles son argumentos opcionales y, si no se definen de forma específica, el valor predeterminado será ilimitado. Consulte la página de comando [man clnode\(1CL\)](#) para obtener más información.

### 3 Asigne valores de factor de carga a cada grupo de recursos.

```
clresourcegroup set -p load_factors=mem_load@50,factor2@1 rg1 rg2
```

En este ejemplo, los factores de carga se definen en los dos grupos de recursos, rg1 y rg2. La configuración de factor de carga se corresponde con los límites de carga definidos para los nodos. También puede realizar este paso durante la creación del grupo de recursos con el comando `clresourcegroup create`. Para obtener más información, consulte la página de comando [man clresourcegroup\(1CL\)](#).

### 4 Si lo desea, puede redistribuir la carga existente (clrg remaster).

```
clresourcegroup remaster rg1 rg2
```

Este comando puede mover los grupos de recursos de su nodo maestro actual a otros nodos para conseguir una distribución uniforme de la carga.

### 5 Si lo desea, puede conceder a algunos grupos de recursos una mayor prioridad que a otros.

```
clresourcegroup set -p priority=600 rg1
```

La prioridad predeterminada es 500. Los grupos de recursos con valores de prioridad mayores tienen preferencia en la asignación de nodos por encima de los grupos de recursos con valores de prioridad menores.

### 6 Si lo desea, puede definir la propiedad `Preemption_mode`.

```
clresourcegroup set -p Preemption_mode=No_cost rg1
```

Consulte la página de comando [man clresourcegroup\(1CL\)](#) para obtener más información sobre las opciones `HAS_COST`, `NO_COST` y `NEVER`.

### 7 Si lo desea, también puede definir el indicador `Concentrate_load`.

```
cluster set -p Concentrate_load=TRUE
```

### 8 Si lo desea, puede especificar una afinidad entre grupos de recursos.

Una afinidad negativa o positiva fuerte tiene preferencia por encima de la distribución de carga. Las afinidades fuertes no pueden infringirse, del mismo modo que los límites de carga fijos. Si define afinidades fuertes y límites de carga fijos, es posible que algunos grupos de recursos estén obligados a permanecer sin conexión, si no pueden cumplirse ambas restricciones.

En el ejemplo siguiente se especifica una afinidad positiva fuerte entre el grupo de recursos rg1 del clúster de zona zc1 y el grupo de recursos rg2 del clúster de zona zc2.

```
clresourcegroup set -p RG_affinities=++zc2:rg2 zc1:rg1
```



**9 Compruebe el estado de todos los nodos de clúster global y de clúster de zona del clúster.**

```
clnode status -Z all -v
```

El resultado incluye todos los valores de límite de carga que hay definidos en el nodo o en sus zonas no globales.

## Ejecución de tareas administrativas del clúster de zona

En un clúster de zona puede efectuar otras tareas administrativas, por ejemplo mover la ruta de zona, preparar un clúster de zona para que ejecuten aplicaciones y clonar un clúster de zona. Todos estos comandos deben ejecutarse desde el nodo de votación del clúster global.

**Nota** – Los comandos de Oracle Solaris Cluster que sólo ejecuta desde el nodo de votación en el clúster global no son válidos para usarlos con los clústeres de zona. Consulte la página de comando `man` de Oracle Solaris Cluster para obtener información sobre la validez del uso de un comando en zonas.

**TABLA 9-2** Otras tareas del clúster de zona

Tarea	Instrucciones
Mover la ruta de zona a una ruta nueva	<code>clzonecluster move -f ruta_zona nombre_clúster_zona</code>
Preparar el clúster de zona para que ejecute aplicaciones	<code>clzonecluster ready -n nombre_nodo nombre_clúster_zona</code>
Clonar un clúster de zona	<code>clzonecluster clone -Z origen_nombre_clúster_zona [-m método_copia] nombre_clúster_zona</code>  Detenga el clúster de zona de origen antes de usar el subcomando <code>clone</code> . El clúster de zona de destino ya debe estar configurado.
Eliminar un clúster de zona	<a href="#">“Eliminación de un clúster de zona” en la página 282</a>
Eliminar un sistema de archivos de un clúster de zona	<a href="#">“Eliminación de un sistema de archivos desde un clúster de zona” en la página 283</a>
Eliminar un dispositivo de almacenamiento desde un clúster de zona	<a href="#">“Eliminación de un dispositivo de almacenamiento desde un clúster de zona” en la página 285</a>
Solucionar el problema de la desinstalación de un nodo	<a href="#">“Solución de problemas de la desinstalación de nodos” en la página 289</a>

TABLA 9-2 Otras tareas del clúster de zona (Continuación)	
Tarea	Instrucciones
Crear, configurar y administrar la MIB de eventos SNMP de Oracle Solaris Cluster	“Creación, configuración y administración de MIB de eventos de SNMP de Oracle Solaris Cluster” en la página 291

## ▼ Eliminación de un clúster de zona

Puede borrar un clúster de zona específico o usar un comodín para eliminar todos los clústers de zona configurados en el clúster global. El clúster de zona debe estar configurado antes de eliminarlo.

- 1

**Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo del clúster global. Siga todos los pasos de este procedimiento desde un nodo del clúster global.**
- 2

**Elimine todos los grupos de recursos y los recursos del clúster de zona.**  
`phys-schost# clresourcegroup delete -F -Z zoneclustername +`

**Nota** – Este paso se efectúa en un nodo del clúster global. Para llevar a cabo este paso en un nodo del clúster de zona, inicie sesión en el nodo del clúster de zona y omita `-Z clúster_zona` desde el comando.

- 3

**Detenga el clúster de zona.**  
`phys-schost# clzonecluster halt zoneclustername`
- 4

**Desinstale el clúster de zona.**  
`phys-schost# clzonecluster uninstall zoneclustername`
- 5

**Anule la configuración del clúster de zona.**  
`phys-schost# clzonecluster delete zoneclustername`

### Ejemplo 9-11 Eliminación de un clúster de zona de un clúster global

```
phys-schost# clresourcegroup delete -F -Z sczone +

phys-schost# clzonecluster halt sczone

phys-schost# clzonecluster uninstall sczone

phys-schost# clzonecluster delete sczone
```

## ▼ Eliminación de un sistema de archivos desde un clúster de zona

Para exportar un sistema de archivos a un clúster zona, utilice un montaje directo o un montaje de bucle.

Los clústeres de zona admiten montajes directos para:

- UFS
- Vxfs
- Sistema de archivos compartidos QFS independiente
- ZFS (exportado como conjunto de datos)

Los clústeres de zona pueden administrar montajes de bucle para:

- UFS
- Vxfs
- Sistema de archivos compartidos QFS independiente
- Sistema de archivos compartidos QFS
- PxFs en UFS
- PxFs en Vxfs

Para obtener instrucciones sobre cómo agregar un sistema de archivos a un clúster de zona, consulte [“Adición de sistemas de archivos a un clúster de zona” de Guía de instalación del software Oracle Solaris Cluster](#).

`phys-schost#` refleja un indicador de clúster global. Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario en un nodo del clúster global que aloje el clúster de zona. Algunos pasos del procedimiento se realizan desde un nodo de clúster global. Otros pasos se efectúan desde un nodo del clúster de zona.**
- 2 **Elimine los recursos relacionados con el sistema de archivos que va a eliminar.**
  - a. **Identifique y elimine los tipos de recursos de Oracle Solaris Cluster, como `HASStoragePlus` y `SUNW.ScalMountPoint`, configurados para el sistema de archivos de clúster de zona que va a quitar.**

```
phys-schost# clresource delete -F -Z zoneclustername fs_zone_resources
```

- b. **Si procede, identifique y elimine los recursos de Oracle Solaris Cluster de tipo `SUNW.qfs` configurados en el clúster global para el sistema de archivos que va a quitar.**

```
phys-schost# clresource delete -F fs_global_resources
```

Use la opción `-F` con cuidado porque fuerza el borrado de todos los recursos que especifique, incluso si no los ha inhabilitado primero. Todos los recursos especificados se eliminan de la configuración de la dependencia de recursos de otros recursos, que pueden provocar la pérdida de servicio en el clúster. Los recursos dependientes que no se borren pueden quedar en estado no válido o de error. Para obtener más información, consulte la página de comando `man clresource(1CL)`.

---

**Consejo** – Si el grupo de recursos del recurso eliminado se vacía posteriormente, puede eliminarlo de forma segura.

---

**3 Determine la ruta del directorio de punto de montaje de sistemas de archivos. Por ejemplo:**

```
phys-schost# clzonecluster configure zoneclustername
```

**4 Elimine el sistema de archivos de la configuración del clúster de zona.**

```
phys-schost# clzonecluster configure zoneclustername
```

```
clzc:zoneclustername> remove fs dir=filesystemdirectory
```

```
clzc:zoneclustername> commit
```

El punto de montaje de sistemas de archivos está especificado por `dir=`.

**5 Compruebe que haya eliminado el sistema de archivos.**

```
phys-schost# clzonecluster show -v zoneclustername
```

## Ejemplo 9–12 Eliminación de un sistema de archivos de alta disponibilidad en un clúster de zona

Este ejemplo muestra cómo eliminar un sistema de archivos con un directorio de punto de montaje (`/local/ufs-1`) configurado en un clúster de zona denominado `sczone`. Este recurso es `hasp-rs` y es del tipo de `HASStoragePlus`.

```
phys-schost# clzonecluster show -v sczone
```

```
...
Resource Name: fs
dir: /local/ufs-1
special: /dev/md/dsk/d0
raw: /dev/md/dsk/r0
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
```

### Ejemplo 9–13 Eliminación de un sistema de archivos ZFS de alta disponibilidad en un clúster de zona

Este ejemplo muestra cómo eliminar los sistemas de archivos de ZFS en una agrupación de ZFS llamada HAZpool, configurada en el clúster de zona sczone del recurso hasp-rs de tipo SUNW.HASStoragePlus.

```
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
```

## ▼ Eliminación de un dispositivo de almacenamiento desde un clúster de zona

Puede eliminar dispositivos de almacenamiento, por ejemplo grupos de discos de SVM y los dispositivos de DID, de un clúster de zona. Siga este procedimiento para eliminar un dispositivo de almacenamiento desde un clúster de zona.

- 1 **Conviértase en superusuario en un nodo del clúster global que aloje el clúster de zona. Algunos pasos del procedimiento se realizan desde un nodo de clúster global. Otros pasos se efectúan desde un nodo del clúster de zona.**
- 2 **Elimine los recursos relacionados con los dispositivos que se van a eliminar. Identifique y elimine los tipos de recursos de Oracle Solaris Cluster, como SUNW.HASStoragePlus y SUNW.ScalDeviceGroup, configurados para los dispositivos del clúster de zona que va a eliminar.**

```
phys-schost# clresource delete -F -Z zoneclustername dev_zone_resources
```

- 3 **Determine la entrada de coincidencia de los dispositivos que se van a eliminar.**

```
phys-schost# clzonecluster show -v zoneclustername
...
Resource Name: device
match: <device_match>
...
```

- 4 **Elimine los dispositivos de la configuración del clúster de zona.**

```
phys-schost# clzonecluster configure zoneclustername
clzc:zoneclustername> remove device match=<devices_match>
clzc:zoneclustername> commit
clzc:zoneclustername> end
```

**5 Rearranque el clúster de zona.**

```
phys-schost# clzonecluster reboot zoneclustername
```

**6 Compruebe que se hayan eliminado los dispositivos.**

```
phys-schost# clzonecluster show -v zoneclustername
```

**Ejemplo 9–14 Eliminación de un conjunto de discos de SVM de un clúster de zona**

Este ejemplo muestra cómo eliminar un conjunto de discos de SVM denominado `apachedg` configurado en un clúster de zona llamado `sczone`. El número establecido del conjunto de discos `apachedg` es 3. El recurso `zc_rs` configurado en el clúster utiliza los dispositivos.

```
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
```

**Ejemplo 9–15 Eliminación de un dispositivo de DID de un clúster de zona**

Este ejemplo muestra cómo eliminar los dispositivos de DID `d10` y `d11`, configurados en un clúster de zona denominado `sczone`. El recurso `zc_rs` configurado en el clúster utiliza los dispositivos.

```
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# clzonecluster reboot sczone
phys-schost# clzonecluster show -v sczone
```

## ▼ Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster

Aplique este procedimiento para desinstalar el software Oracle Solaris Cluster desde un nodo del clúster global antes de desconectarlo de una configuración del clúster completamente establecida. Puede seguir este procedimiento para desinstalar el software desde el último nodo de un clúster.

---

**Nota** – No siga este procedimiento para desinstalar el software Oracle Solaris Cluster desde un nodo que todavía no se haya unido al clúster o que aún esté en modo de instalación. En lugar de eso, vaya a "Cómo anular la configuración del software Oracle Solaris Cluster para solucionar problemas de instalación" en [Guía de instalación del software Oracle Solaris Cluster](#).

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Compruebe que haya completado correctamente todas las tareas de los requisitos previos del mapa de tareas para eliminar un nodo del clúster.**

Consulte la [Tabla 8–2](#).

---

**Nota** – Compruebe que haya eliminado el nodo de la configuración del clúster mediante `clnode remove` antes de continuar con este procedimiento.

---

- 2 Conviértase en superusuario en un miembro activo del clúster global *que no sea el nodo del clúster global que va a desinstalar*. Realice este procedimiento desde un nodo del clúster global.**
- 3 En el miembro activo del clúster, agregue el nodo que desea desinstalar a la lista de autenticación de nodos del clúster.**

`phys-schost# claccess allow -h hostname`

`-h` Especifica el nombre del nodo que se va a agregar a la lista de autenticación de nodos.

También puede usar la utilidad `clsetup(1CL)`. Consulte los procedimientos en [“Adición de un nodo a la lista de nodos autorizados” en la página 244](#).

- 4 Conviértase en superusuario en el nodo que vaya a desinstalar.**

**5 Si tiene un clúster de zona, desinstálelo.**

```
phys-schost# clzonecluster uninstall -F zoneclustername
```

Para consultar los pasos específicos, “[Eliminación de un clúster de zona](#)” en la página 282.

**6 Si el nodo tiene una partición dedicada para el espacio de nombres de dispositivos globales, re arranque el nodo del clúster global en un modo que no sea de clúster.**

- Ejecute el comando siguiente en un sistema basado en SPARC.

```
shutdown -g0 -y -i0ok boot -x
```

- Ejecute los comandos siguientes en un sistema basado en x86.

```
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
```

**7 En el archivo /etc/vfstab, elimine todas las entradas de sistema de archivos montadas globalmente *excepto* los montajes globales de /global/.devices.****8 Si quiere volver a instalar el software Oracle Solaris Cluster en este nodo, elimine la entrada de Oracle Solaris Cluster del registro de productos de Sun Java Enterprise System (Java ES).**

Si en el registro de productos de Java ES aparece que se ha instalado el software Oracle Solaris Cluster, el programa de instalación muestra el componente de Oracle Solaris Cluster atenuado y no permite la reinstalación.

**a. Inicie el programa de desinstalación de Java ES.**

Ejecute el comando siguiente, donde *ver* es la versión de distribución de Java ES desde la que instaló el software Oracle Solaris Cluster.

```
/var/sadm/prod/SUNWentsys ver /uninstall
```

**b. Siga los indicadores para seleccionar Oracle Solaris Cluster para desinstalar.**

Para obtener más información sobre cómo utilizar el comando `uninstall`, consulte el [Capítulo 8, “Uninstalling” de Sun Java Enterprise System 5 Update 1 Installation Guide for UNIX](#).



- 9 Si no tiene previsto volver a instalar el software Oracle Solaris Cluster en este clúster, desconecte los cables y el conmutador de transporte, si los hubiera, de los otros dispositivos del clúster.
  - a. Si el nodo desinstalado está conectado a un dispositivo de almacenamiento que usa una interfaz SCSI paralelo, instale un terminador de SCSI al conector de SCSI del dispositivo de almacenamiento después de haber desconectado los cables de transporte.  
Si el nodo desinstalado está conectado a un dispositivo de almacenamiento que usa interfaces de canal de fibra, el cierre no es necesario.
  - b. Consulte la documentación suministrada con el adaptador del sistema y el servidor respecto a los procedimientos de desconexión.

---

**Consejo** – Si usa un dispositivo de interfaz de archivo de bucle invertido (lofi), el programa de desinstalación de Java ES elimina de forma automática el archivo lofi, que se denomina `/globaldevices`. Para obtener más información cómo migrar un espacio de nombres de dispositivos globales a un lofi, consulte [“Migración del espacio de nombre de dispositivos globales” en la página 128](#).

---

## Solución de problemas de la desinstalación de nodos

Esta sección describe los mensajes de error que pueden generarse al ejecutar el comando `clnode remove` y las medidas correctivas que debe utilizar.

### Entradas no eliminadas del sistema de archivos de clúster

Los mensajes de error siguientes indican que el nodo del clúster global que ha eliminado todavía tiene sistemas de archivo del clúster referenciados en el archivo `vfstab`.

```
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.
```

Para corregir este error, vuelva a [“Desinstalación del software Oracle Solaris Cluster desde un nodo del clúster” en la página 287](#) y repita el procedimiento. Compruebe que haya realizado correctamente el [Paso 7](#) del procedimiento antes de volver a ejecutar el comando `clnode remove`.

## No supresión de la lista de grupos de dispositivos

Los mensajes de error siguientes indican que el nodo que ha eliminado todavía está en la lista con un grupo de dispositivos.

```
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.
```

## Falta la secuencia de comandos de desinstalación

Si no ha utilizado el programa `installer` para instalar o actualizar el software Sun Cluster 3.1 o 3.2 que ahora quiere eliminar, no habrá ninguna secuencia de comandos de desinstalación disponible para esta versión del software. En su lugar, siga los pasos indicados a continuación para desinstalar el software.

---

**Nota** – El software Oracle Solaris Cluster 3.3 se instala con el programa de instalación, por lo que estos pasos no se aplican a esta versión del software.

---

### ▼ Cómo desinstalar el software Sun Cluster 3.1 y 3.2 sin una secuencia de comandos de desinstalación

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Acceda a un directorio que no esté asociado con ningún paquete de Sun Cluster.**  

```
cd /directory
```
- 3 **Desinstale el software de Sun Cluster del nodo.**  

```
scinstall -r
```
- 4 **Cambie el nombre del archivo `productregistry` para permitir la posible reinstalación del software en el futuro.**  

```
mv /var/sadm/install/productregistry /var/sadm/install/productregistry.sav
```

## Creación, configuración y administración de MIB de eventos de SNMP de Oracle Solaris Cluster

Esta sección describe cómo crear, configurar y administrar la base de información de administración (MIB) del evento del protocolo SNMP. Esta sección también describe cómo habilitar, inhabilitar y cambiar la MIB de eventos de SNMP de Oracle Solaris Cluster.

El software Oracle Solaris Cluster admite actualmente una MIB, la MIB de eventos. El software del administrador de SNMP intercepta los eventos del clúster en tiempo real. Si se habilita, el administrador de SNMP envía automáticamente notificaciones de captura a todos los sistemas definidos en el comando `clsnmp host`. La MIB mantiene una tabla de sólo lectura con los 50 eventos más actuales. Debido a que los clústers generan numerosas notificaciones, sólo se envían como notificaciones de capturas los eventos con cierto grado de warning (advertencia) o superior. Esta información no se mantiene después de rearrancar.

La MIB de eventos de SNMP está definida en el archivo `sun-cluster-event-mib.mib` y se ubica en el directorio `/usr/cluster/lib/mib`. Esta definición puede usarse para interpretar la información de captura de SNMP.

El número de puerto predeterminado del módulo SNMP del evento es 11161 y el puerto predeterminado de las capturas de SNMP es 11162. Estos números de puerto se pueden cambiar modificando el archivo de propiedad de Common Agent Container, `/etc/cacao/instances/default/private/cacao.properties`.

Crear, configurar y administrar una MIB de eventos de SNMP de Oracle Solaris Cluster puede implicar las tareas siguientes.

**TABLA 9-3** Mapa de tareas: crear, configurar y administrar la MIB de eventos de SNMP de Oracle Solaris Cluster

Tarea	Instrucciones
Habilitar una MIB de eventos de SNMP	“Habilitación de una MIB de eventos de SNMP” en la página 292
Inhabilitar una MIB de eventos de SNMP	“Inhabilitación de una MIB de eventos de SNMP” en la página 292
Cambiar una MIB de eventos de SNMP	“Cambio de una MIB de eventos de SNMP” en la página 293
Agregar un host de SNMP a la lista de hosts que recibirán las notificaciones de captura de las MIB	“Habilitación de un host de SNMP para que reciba capturas de SNMP en un nodo” en la página 294
Quitar un host de SNMP	“Inhabilitación de la recepción de capturas de SNMP en un nodo por parte del host de SNMP” en la página 294
Agregar un usuario de SNMP	“Adición de un usuario de SNMP a un nodo” en la página 295

TABLA 9-3    Mapa de tareas: crear, configurar y administrar la MIB de eventos de SNMP de Oracle Solaris Cluster  
(Continuación)

Tarea	Instrucciones
Suprimir un usuario de SNMP	<a href="#">“Eliminación de un usuario de SNMP de un nodo” en la página 296</a>

▼ **Habilitación de una MIB de eventos de SNMP**

Este procedimiento muestra cómo habilitar una MIB de eventos de SNMP.

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1    **Conviértase en superusuario o asuma una función que cuente con autorización RBAC solaris.cluster.modify.**
- 2    **Habilite la MIB de eventos de SNMP.**  
phys-schost-1# **clsnmpmib enable [-n node] MIB**  
[-n nodo]                    Especifica el *nodo* en el que se ubica la MIB del evento que desea habilitar. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.  
  
MIB                            Especifica el nombre de la MIB que desea habilitar. En este caso, el nombre de la MIB debe ser event.

▼ **Inhabilitación de una MIB de eventos de SNMP**

Este procedimiento muestra cómo inhabilitar una MIB de eventos de SNMP.

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1    **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC solaris.cluster.modify.**
- 2    **Inhabilite la MIB de eventos de SNMP.**  
phys-schost-1# **clsnmpmib disable -n node MIB**

<code>-n nodo</code>	Especifica el <i>nodo</i> en el que se ubica la MIB del evento que desea inhabilitar. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.
<i>MIB</i>	Especifica el tipo de MIB que desea inhabilitar. En este caso, debe especificar el tipo event.

## ▼ Cambio de una MIB de eventos de SNMP

Este procedimiento muestra cómo cambiar una MIB de eventos de SNMP.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Cambie el protocolo de la MIB de eventos de SNMP.**

```
phys-schost-1# clsnmpmib set -n node -p version=valor MIB
```

`-n nodo`

Especifica el *nodo* en el que se ubica la MIB del evento que desea cambiar. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.

`-p version=valor`

Especifica la versión del protocolo SNMP que desea usar con las MIB. Especifique el *valor* tal y como se indica a continuación:

- `version=SNMPv2`
- `version=snmpv2`
- `version=2`
- `version=SNMPv3`
- `version=snmpv3`
- `version=3`

*MIB*

Especifica el nombre de la MIB o las MIB a las que hay que aplicar el subcomando. En este caso, debe especificar el tipo event. Si no especifica este operando, el subcomando usa el signo más predeterminado (+), que equivale a todas las MIB. Si usa el operando de *MIB*, especifique la MIB en una lista delimitada por espacios después de todas las demás opciones de línea de comandos.

## ▼ **Habilitación de un host de SNMP para que reciba capturas de SNMP en un nodo**

Este procedimiento muestra cómo agregar un host de SNMP de un nodo a la lista de sistemas que recibirá las notificaciones de captura para las MIB.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**

- 2 **Agregue el host a la lista de hosts de SNMP de una comunidad en otro nodo.**

```
phys-schost-1# clsnmhost add -c SNMPcommunity [-n node] host
```

**-c *comunidad\_SNMP***

Especifica el nombre de la comunidad de SNMP que se usa en combinación con el nombre del host.

Se debe especificar el nombre de comunidad de SNMP *comunidad\_SNMP* cuando agregue un host que no sea `public` a la comunidad. Si usa el subcomando `add` sin la opción `-c`, el subcomando utiliza `public` como nombre de comunidad predeterminado.

Si el nombre de comunidad especificado no existe, este comando crea la comunidad.

**-n *nodo***

Especifica el nombre del *nodo* del host de SNMP al que se ha proporcionado acceso a las MIB de SNMP en el clúster. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.

***host***

Especifica el nombre, la dirección IP o la dirección IPv6 del host al que se ha proporcionado acceso a las MIB de SNMP en el clúster.

## ▼ **Inhabilitación de la recepción de capturas de SNMP en un nodo por parte del host de SNMP**

Este procedimiento muestra cómo eliminar un host de SNMP de un nodo de la lista de hosts que recibirán las notificaciones de captura para las MIB.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**

- 2 **Elimine el host de la lista de hosts de SNMP de una comunidad en el nodo especificado.**

```
phys-schost-1# clsnmphost remove -c SNMPcommunity -n node host
```

*remove*

Elimina el host de SNMP especificado del nodo indicado.

*-c comunidad\_SNMP*

Especifica el nombre de la comunidad de SNMP de la que se ha eliminado el host de SNMP.

*-n nodo*

Especifica el nombre del *nodo* en el que se ha eliminado de la configuración el host de SNMP. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.

*host*

Especifica el nombre, la dirección IP o la dirección IPv6 del host que se ha eliminado de la configuración.

Para eliminar todos los sistemas de la comunidad SNMP especificada, use un signo más (+) para el *host* con la opción *-c*. Para eliminar todos los sistemas, use el signo más (+) para el *host*.

## ▼ Adición de un usuario de SNMP a un nodo

Este procedimiento muestra cómo agregar un usuario de SNMP a la configuración de usuarios de SNMP en un nodo.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**

- 2 **Agregue el usuario de SNMP.**

```
phys-schost-1# clsnmpuser create -n node -a authentication \
-f password user
```

-n <i>nodo</i>	Especifica el nodo en el que se ha agregado el usuario de SNMP. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.
-a <i>autenticación</i>	Especifica el protocolo de autenticación que se utiliza para autorizar al usuario. El valor del protocolo de autenticación puede ser SHA o MD5.
-f <i>contraseña</i>	Especifica un archivo con las contraseñas de usuario de SNMP. Si no especifica esta opción al crear un usuario, el comando solicita una contraseña. Esta opción sólo es válida con el subcomando add.  Las contraseñas de los usuarios deben especificarse en líneas distintas y con el formato siguiente: <i>user:password</i>  Las contraseñas no pueden tener espacios ni los caracteres siguientes:     ▪ ; (punto y coma)    ▪ : (dos puntos)    ▪ \ (barra diagonal inversa)    ▪ \n (línea nueva)
<i>usuario</i>	Especifica el nombre del usuario de SNMP que desea agregar.

▼ **Eliminación de un usuario de SNMP de un nodo**

Este procedimiento muestra cómo eliminar un usuario de SNMP de la configuración de usuarios de SNMP en un nodo.

*phys-schost#* refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Quite el usuario de SNMP.**

<i>phys-schost-1# clnmpuser delete -n node user</i>	
-n <i>nodo</i>	Especifica el nodo desde el que se ha eliminado el usuario de SNMP. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual.
<i>usuario</i>	Especifica el nombre del usuario de SNMP que desea eliminar.



# Solución de problemas

Esta sección contiene procedimientos de solución de problemas válidos para tareas de prueba.

## Ejecución de una aplicación fuera del clúster global

### ▼ Ejecución de un metaconjunto de Solaris Volume Manager de los nodos que han arrancado en un modo que no es de clúster

Siga este procedimiento para ejecutar una aplicación fuera del clúster global para tareas de prueba.

- 1 **Determine si el dispositivo de quórum se utiliza en el metaconjunto de Solaris Volume Manager y si usa reservas de SCSI2 y SCSI3.**

```
phys-schost# clquorum show
```

- a. Si el dispositivo de quórum se encuentra en el metaconjunto de Solaris Volume Manager, agregue un dispositivo de quórum nuevo que no sea parte del metaconjunto para ponerlo más tarde en un modo que no sea de clúster.

```
phys-schost# clquorum add did
```

- b. Elimine el dispositivo de quórum antiguo.

```
phys-schost# clqorum remove did
```

- c. Si el dispositivo de quórum usa una reserva de SCSI2, quite la reserva de SCSI2 de quórum antiguo y compruebe que no queden reservas de SCSI2.

Para saber cómo ejecutar los comandos `pgre`, es necesario instalar y utilizar el paquete Diagnostic Toolkit (SUNWscdtk) proporcionado por el equipo de asistencia de Oracle.

- 2 **Evacue el nodo de clúster global que desea arrancar en un modo que no sea de clúster.**

```
phys-schost# clresourcegroup evacuate -n targetnode
```

- 3 **Ponga fuera de línea todos los grupos de recursos que contengan recursos HASStorage o HASStoragePlus y que contengan dispositivos o sistemas de archivos afectados por el metaconjunto que desee poner en un modo que no sea de clúster más tarde.**

```
phys-schost# clresourcegroup offline resourcegroupname
```

- 4 **Inhabilite todos los recursos de los grupos de recursos puestos fuera de línea.**

```
phys-schost# clresource disable resourcename
```

- 5 **Anule la administración de los grupos de recursos.**  
`phys-schost# clresourcegroup unmanage resourcegroupname`
- 6 **Ponga fuera de línea el grupo o los grupos de dispositivos correspondientes.**  
`phys-schost# cldevicegroup offline devicegroupname`
- 7 **Inhabilite el grupo o los grupos de dispositivos.**  
`phys-schost# cldevicegroup disable devicegroupname`
- 8 **Arranque el nodo pasivo en un modo que no sea de clúster.**  
`phys-schost# reboot -x`
- 9 **Antes de seguir, compruebe que haya finalizado el proceso de arranque en el nodo pasivo.**  
`phys-schost# svcs -x`
- 10 **Determine si hay reservas de SCSI3 en los discos de los metaconjuntos. Ejecute el comando siguiente en todos los discos de los metaconjuntos.**  
`phys-schost# /usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/dids2`
- 11 **Si hay reservas de SCSI3 en los discos, quítelas.**  
`phys-schost# /usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/dids2`
- 12 **Utilice el metaconjunto del nodo evacuado.**  
`phys-schost# metaset -s name -C take -f`
- 13 **Monte el sistema o los sistemas de archivos que contengan el dispositivo definido en el metaconjunto.**  
`phys-schost# mount device mountpoint`
- 14 **Inicie la aplicación y realice la prueba que desee. Cuando finalice la prueba, detenga la aplicación.**
- 15 **Rearranque el nodo y espere hasta que el proceso de arranque haya terminado.**  
`phys-schost# reboot`
- 16 **Ponga en línea el grupo o los grupos de dispositivos.**  
`phys-schost# cldevicegroup online -e devicegroupname`
- 17 **Inicie el grupo o los grupos de recursos.**  
`phys-schost# clresourcegroup online -eM resourcegroupname`

## Restauración de un conjunto de discos dañado

Siga este procedimiento si un conjunto de discos está dañado o se encuentra en un estado en que los nodos del clúster no pueden hacerse propietarios del conjunto de discos. Si los intentos de borrar el estado han sido en vano, en última instancia siga este procedimiento para corregir el conjunto de discos.

Estos procedimientos funcionan para los metaconjuntos de Solaris Volume Manager y los metaconjuntos de varios propietarios de Solaris Volume Manager.

### ▼ Cómo guardar la configuración de software de Solaris Volume Manager

Restaurar un conjunto de discos desde cero puede llevar mucho tiempo y causar errores. La mejor alternativa es utilizar el comando `metastat` para realizar copias de seguridad de las réplicas de forma regular o utilizar el explorador (SUNWexplo) para crear una copia de seguridad. A continuación, puede utilizar la configuración guardada para volver a crear el conjunto de discos. Debe guardar la configuración actual en archivos (mediante el comando `prvtoc` y los comandos `metastat`) y, a continuación, vuelva a crear el conjunto de discos y sus componentes. Consulte [“Cómo volver a crear la configuración de software de Solaris Volume Manager” en la página 300](#).

#### 1 Guarde la tabla de particiones para cada uno de los discos del conjunto de discos.

```
/usr/sbin/prvtoc /dev/global/rdisk/diskname > /etc/lvm/diskname.vtoc
```

#### 2 Guarde la configuración de software de Solaris Volume Manager.

```
/bin/cp /etc/lvm/md.tab /etc/lvm/md.tab_ORIGINAL
```

```
/usr/sbin/metastat -p -s setname >> /etc/lvm/md.tab
```

---

**Nota** – Otros archivos de configuración, como el archivo `/etc/vfstab`, pueden hacer referencia al software Solaris Volume Manager. Este procedimiento presupone que se reconstruye una configuración de software de Solaris Volume Manager idéntica y, por tanto, la información de montaje es la misma. Si el explorador (SUNWexplo) se ejecuta en un nodo propietario del conjunto, recuperará la información del comando `prvtoc` y `metaset -p`.

---

### ▼ Cómo purgar el conjunto de discos dañado

Al purgar un conjunto de un nodo o de todos los nodos se elimina la configuración. Para depurar un conjunto de discos de un nodo, el nodo no debe ser propietario del conjunto de discos.

#### 1 Ejecute el comando purge en todos los nodos.

```
/usr/sbin/metaset -s setname -P
```

Al ejecutar este comando, se elimina la información del conjunto de discos de las réplicas de base de datos, así como del depósito de Oracle Solaris Cluster. Las opciones -P y -C permiten purgar un conjunto de discos sin tener que reconstruir completamente el entorno de Solaris Volume Manager.

---

**Nota** – Si se purga un conjunto de discos de varios propietarios mientras los nodos se están iniciando fuera del modo clúster, es posible que tenga que instalar y utilizar el paquete Diagnostic Toolkit (SUNWscdtk) proporcionado por el equipo de asistencia de Oracle. El kit de herramientas elimina la información de los archivos de configuración dcs. Consulte el [Paso 2](#).

---

**2 Si sólo desea eliminar la información del conjunto de discos de las réplicas de base de datos, utilice el comando siguiente.**

```
/usr/sbin/metaset -s setname -C purge
```

Por lo general, debería utilizar la opción -P, en lugar de la opción -C. Utilizar la opción -C puede causar problemas al volver a crear el conjunto de discos, porque el software Oracle Solaris Cluster sigue reconociendo el conjunto de discos.

- a. Si ha utilizado la opción -C con el comando `metaset`, primero cree el conjunto de discos para ver si se produce algún problema.
- b. En caso afirmativo, utilice el paquete Diagnostic Toolkit (SUNWscdtk) para eliminar la información de los archivos de configuración dcs.

Si las opciones del comando `purge` no se completan correctamente, compruebe si tiene instalados los parches más recientes del núcleo y del metadispositivo, y contacte con Oracle Solaris Cluster.

## ▼ **Cómo volver a crear la configuración de software de Solaris Volume Manager**

Siga este procedimiento únicamente si pierde por completo la configuración de software de Solaris Volume Manager. En estos pasos se presupone que ha guardado la configuración actual de Solaris Volume Manager y todos sus componentes, además de haber purgado el conjunto de discos dañado.

---

**Nota** – Los mediadores sólo deben utilizarse en clústeres de dos nodos.

---

**1 Cree un conjunto de discos.**

```
/usr/sbin/metaset -s setname -a -h nodename1 nodename2
```

Si se trata de un conjunto de discos de varios propietarios, utilice el comando siguiente para crear un conjunto de discos.

```
/usr/sbin/metaset -s setname -aM -h nodename1 nodename2
```

- 2 En el mismo host donde se haya creado el conjunto, agregue los hosts mediadores si es preciso (sólo dos nodos).**

```
/usr/sbin/metaset -s setname -a -m nodename1 nodename2
```

- 3 Vuelva a agregar los mismos discos al conjunto de discos de este mismo host.**

```
/usr/sbin/metaset -s setname -a /dev/did/rdisk/diskname /dev/did/rdisk/diskname
```

- 4 Si ha purgado el conjunto de discos y está volviendo a crearlo, el índice de contenido del volumen (VTOC) debe permanecer en los discos para que pueda omitir este paso. Sin embargo, si está volviendo a crear un conjunto que se va a recuperar, debe dar formato a los discos según una configuración guardada en el archivo `/etc/lvm/nombre_disco.vtoc`. Por ejemplo:**

```
/usr/sbin/fmthard -s /etc/lvm/d4.vtoc /dev/global/rdsk/d4s2
```

```
/usr/sbin/fmthard -s /etc/lvm/d8.vtoc /dev/global/rdsk/d8s2
```

Este comando se puede ejecutar en cualquier nodo.

- 5 Compruebe la sintaxis en el archivo `/etc/lvm/md.tab` existente para cada metadispositivo.**

```
/usr/sbin/metainit -s setname -n -a metadvice
```

- 6 Cree cada metadispositivo a partir de una configuración guardada.**

```
/usr/sbin/metainit -s setname -a metadvice
```

- 7 Si ya existe un sistema de archivos en el metadispositivo, ejecute el comando `fsck`.**

```
/usr/sbin/fsck -n /dev/md/setname/rdsk/metadvice
```

Si el comando `fsck` sólo muestra algunos errores, como el recuento de superbloqueos, probablemente el dispositivo se haya reconstruido de forma correcta. A continuación, puede ejecutar el comando `fsck` sin la opción `-n`. Si surgen varios errores, compruebe si el metadispositivo se ha reconstruido correctamente. En caso afirmativo, revise los errores del comando `fsck` para determinar si se puede recuperar el sistema de archivos. Si no se puede recuperar, deberá restablecer los datos a partir de una copia de seguridad.

- 8 Concatene el resto de metaconjuntos en todos los nodos del clúster con el archivo `/etc/lvm/md.tab` y, a continuación, concatene el conjunto de discos local.**

```
/usr/sbin/metastat -p >> /etc/lvm/md.tab
```



## Configuración del control del uso de la CPU

---

Si desea controlar el uso de la CPU, configure la función de control de la CPU. Para obtener más información sobre cómo configurar la función de control de la CPU, consulte la página de comando `man rg_properties(5)`. Este capítulo proporciona información sobre los temas siguientes:

- “Introducción al control de la CPU” en la página 303
- “Configuración del control de la CPU” en la página 305

### Introducción al control de la CPU

El software Oracle Solaris Cluster permite controlar el uso de la CPU.

La función de control de la CPU se basa en las funciones disponibles en el sistema operativo Oracle Solaris. Para obtener más información sobre zonas, proyectos, agrupaciones de recursos, conjuntos de procesadores y clases de planificación, consulte *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

En el sistema operativo Oracle Solaris se pueden realizar las tareas siguientes:

- Asignar recursos compartidos de CPU a grupos de recursos.
- Asignar procesadores a grupos de recursos.

### Elección de una situación hipotética

Según las opciones de configuración y la versión del sistema operativo que elija, puede tener niveles distintos de control de la CPU. Todos los aspectos de control de la CPU que se describen en este capítulo dependen de la propiedad del grupo de recursos `RG_SLM_TYPE` establecida en `automated`.

La [Tabla 10–1](#) describe las diferentes situaciones hipotéticas de configuración posibles.

TABLA 10-1 Situaciones hipotéticas de control de la CPU

Descripción	Instrucciones
El grupo de recursos se ejecuta en el nodo de votación del clúster global.   Asigne recursos compartidos de CPU a los grupos y las zonas de recursos; para ello, proporcione valores para <code>project.cpu-shares</code> y <code>zone.cpu-shares</code>.   Este procedimiento es válido con nodos sin votación de clúster global configurados o no configurados.	“Control del uso de la CPU en el nodo de votación de un clúster global” en la página 305
El grupo de recursos se ejecuta en una zona sin votación de clúster global mediante el conjunto de procesadores predeterminado.   Asigne recursos compartidos de CPU a los grupos y las zonas de recursos; para ello, proporcione valores para <code>project.cpu-shares</code> y <code>zone.cpu-shares</code>.   Efectúe este procedimiento si no necesita controlar el tamaño del conjunto de procesadores.	“Control del uso de la CPU en un nodo sin votación de clúster global con el conjunto de procesadores predeterminado” en la página 307
El grupo de recursos se ejecuta en un nodo sin votación de clúster global con un conjunto de procesadores dedicado.   Asigne recursos compartidos de CPU a grupos de recursos; para ello, proporcione valores para <code>project.cpu-shares</code>, <code>zone.cpu-shares</code> y el número máximo de procesadores en un conjunto de procesadores dedicado.   Establezca el número mínimo de conjuntos de procesadores en un conjunto de procesadores dedicado.   Siga este procedimiento si desea controlar los recursos compartidos de CPU y el tamaño de un conjunto de procesadores. Sólo puede tener este control en un nodo sin votación de clúster global mediante un conjunto de procesadores dedicado.	“Control del uso de la CPU en un nodo sin votación de clúster global con un conjunto de procesadores dedicado” en la página 310

## Planificador por reparto equitativo

El primer paso del procedimiento para asignar recursos compartidos de CPU a grupos de recursos es configurar un planificador del sistema para que sea el planificador por reparto equitativo (FSS, Fair Share Scheduler). De forma predeterminada, la clase de programación



para el sistema operativo Oracle Solaris es la planificación de tiempo compartido (TS, Timesharing Schedule). Configure el planificador para que sea de clase FSS y que se aplique la configuración de los recursos compartidos.

Puede crear un conjunto de procesadores dedicado sea cual sea el planificador que se elija.

## Configuración del control de la CPU

Esta sección incluye los procedimientos siguientes:

- “Control del uso de la CPU en el nodo de votación de un clúster global” en la página 305
- “Control del uso de la CPU en un nodo sin votación de clúster global con el conjunto de procesadores predeterminado” en la página 307
- “Control del uso de la CPU en un nodo sin votación de clúster global con un conjunto de procesadores dedicado” en la página 310

### ▼ Control del uso de la CPU en el nodo de votación de un clúster global

Realice este procedimiento para asignar los recursos compartidos de CPU a un grupo de recursos que se ejecutará en un nodo de votación de clúster global.

Si un grupo de recursos se asigna a los recursos compartidos de CPU, el software Oracle Solaris Cluster realiza las tareas siguientes al iniciar un recurso del grupo en un nodo de votación de clúster global:

- Aumenta el número de recursos compartidos de CPU asignados al nodo de votación (`zone.recursos_compartidos_cpu`) con el número especificado de recursos compartidos de CPU, si todavía no se ha hecho.
- Crea un proyecto denominado `SCSLM_nombre_grupo_recursos` en el nodo de votación, si todavía no se ha hecho. Este proyecto es específico del grupo de recursos y se le asigna el número específico de recursos compartidos de CPU (`project.recursos_compartidos_cpu`).
- Inicia el recurso en el proyecto `SCSLM_nombre_grupo_recursos`.

Para obtener más información sobre cómo configurar la función de control de la CPU, consulte la página de comando `man rg_properties(5)`.

#### 1 Configure FSS como tipo de programador predeterminado para el sistema.

```
dispadmin -d FSS
```

FSS se convierte en el programador predeterminado a partir del siguiente rearranque. Para que esta configuración se aplique inmediatamente, use el comando `priocntl`.

```
priocntl -s -C FSS
```

Mediante la combinación de los comandos `priocntl` y `dispadm`, se asegura de que FSS se convierta en el programador predeterminado inmediatamente y de que permanezca después del rearranque. Para obtener más información sobre cómo configurar una clase de programación, consulte las páginas de comando `man dispadm(1M)` y `priocntl(1)`.

---

**Nota** – Si FSS no es el programador predeterminado, la asignación de recursos compartidos de CPU no se lleva a cabo.

---

**2 En todos los nodos que usen el control de la CPU, configure el número de recursos compartidos para los nodos de votación del clúster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.**

La configuración de estos parámetros protege procesos que se ejecuten en los nodos de votación de competir por las CPU con los procesos que se ejecuten en nodos sin votación. Si no asigna un valor a las propiedades `globalzoneshares` y `defaulttpsetmin`, asumen sus valores predeterminados.

```
clnode set [-p globalzoneshares=integer] \
[-p defaulttpsetmin=integer] \
node
```

`-p defaulttpsetmin= entero_mínimo_pestablecido_predeterminado`

Establece el número mínimo de recursos compartidos de CPU disponibles en el conjunto de procesadores predeterminado. El valor predeterminado es 1.

`-p globalzoneshares= entero`

Configura el número de recursos compartidos asignados al nodo de votación. El valor predeterminado es 1.

*nodo*

Especifica los nodos donde van a configurarse las propiedades.

Al configurar estas propiedades, también se establecen las del nodo de votación. Si no configura estas propiedades, no se puede beneficiar de la propiedad de `RG_SLM_PSET_TYPE` en los nodos sin votación.

**3 Compruebe si ha configurado correctamente estas propiedades.**

```
clnode show node
```

Para el nodo que especifique, el comando `clnode` imprime las propiedades establecidas y los valores definidos para dichas propiedades. Si no configura las propiedades de control de la CPU con `clnode`, asumen el valor predeterminado.

**4 Configure la función de control de la CPU.**

```
cresourcegroup create -p RG_SLM_TYPE=automated \
[-p RG_SLM_CPU_SHARES=value] resource_group_name
```

<code>-p RG_SLM_TYPE=automated</code>	Permite controlar el uso de la CPU y automatiza algunos pasos para configurar el sistema operativo Oracle Solaris con el fin de administrar recursos del sistema.
<code>-p RG_SLM_CPU_SHARES= valor</code>	Especifica el número de recursos compartidos de CPU asignados al proyecto específico del grupo de recursos, <code>project.cpu-shares</code> , y determina el número de recursos compartidos CPU asignados al nodo de votación <code>zone.cpu-shares</code> .
<code>nombre_grupo_recursos</code>	Especifica el nombre del grupo de recursos.

En este procedimiento, no se configura la propiedad `RG_SLM_PSET_TYPE`. En el nodo de votación, esta propiedad asume el valor `default`.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos.

## 5 Active el cambio de configuración.

```
clresourcegroup online -M resource_group_name
```

`nombre_grupo_recursos` Especifica el nombre del grupo de recursos.

---

**Nota** – No elimine ni modifique el proyecto de `SCSLM_nombre_grupo_recursos`. Puede agregar más control de recursos de forma manual al proyecto; por ejemplo, puede configurar la propiedad `project.max-lwps`. Para obtener más información, consulte la página de comando [man projmod\(1M\)](#).

---

## ▼ Control del uso de la CPU en un nodo sin votación de clúster global con el conjunto de procesadores predeterminado

Siga este procedimiento si desea asignar recursos compartidos de CPU a grupos de recursos en un nodo sin votación de clúster global pero sin tener que crear un conjunto de procesadores dedicado.

Si se asigna un grupo de recursos a los recursos compartidos de CPU, el software Oracle Solaris Cluster realiza las tareas siguientes al iniciar un recurso de ese grupo en un nodo sin votación:

- Crea una agrupación con el nombre `SCSLM_nombre_grupo_recursos`, si todavía no se ha hecho.
- Asocia la agrupación `SCSLM_pool_nombre_zona` con el conjunto de recursos predeterminado.

- Vincula de forma dinámica el nodo sin votación con la agrupación `SCSLM_pool_nombre_zona`.
- Aumenta el número de recursos compartidos de CPU asignados al nodo sin votación (`zone.cpu-shares`) con el número especificado de recursos compartidos de CPU, si todavía no se ha hecho.
- Crea un proyecto con el nombre `SCSLM_nombre_grupo_recursos` en el nodo sin votación, si todavía no se ha hecho. Este proyecto es específico del grupo de recursos y se le asigna el número especificado de recursos compartidos de CPU (`proyecto.recursos_compartidos_cpu`).
- Inicia el recurso en el proyecto de `SCSLM_nombre_grupo_recursos`.

Para obtener más información sobre cómo configurar la función de control de la CPU, consulte la página de comando `man rg_properties(5)`.

## 1 Configure FSS como tipo de programador predeterminado para el sistema.

```
dispadmin -d FSS
```

FSS se convierte en el programador predeterminado a partir del siguiente rearranque. Para que esta configuración se aplique de forma inmediata, use el comando `priocntl`:

```
priocntl -s -C FSS
```

Mediante la combinación de los comandos `priocntl` y `dispadmin`, se asegura de que FSS se convierta en el programador predeterminado inmediatamente y de que permanezca después del rearranque. Para obtener más información sobre cómo configurar un tipo de planificación, consulte las páginas de comando `man dispadmin(1M)` y `priocntl(1)`.

---

**Nota** – Si FSS no es el programador predeterminado, la asignación de recursos compartidos de CPU no se lleva a cabo.

---

## 2 En todos los nodos que usen un control de la CPU, configure el número de recursos compartidos para el nodo de votación de clúster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.

La configuración de estos parámetros protege los procesos que se estén ejecutando en el nodo de votación de competir por las CPU con los procesos que se estén ejecutando en nodos sin votación de clúster global. Si no asigna un valor a las propiedades `globalzoneshares` y `defaultpsetmin`, asumen sus valores predeterminados.

```
clnode set [-p globalzoneshares=integer] \
[-p defaultpsetmin=integer] \
node
```

```
-p globalzoneshares=entero
```

Configura el número de recursos compartidos asignados al nodo de votación. El valor predeterminado es 1.

`-p defaulttpsetmin= minintero_pestablecido_predeterminado`  
 Establece el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado. El valor predeterminado es 1.

*nodo*

Identifica los nodos donde se deben configurar propiedades.

Al configurar estas propiedades, también se establecen las del nodo de votación.

### 3 Compruebe si ha configurado estas propiedades correctamente:

```
clnode show node
```

Para el nodo que especifique, el comando `clnode` imprime las propiedades establecidas y los valores definidos para dichas propiedades. Si no configura las propiedades de control de la CPU con `clnode`, asumen el valor predeterminado.

### 4 Configure la función de control de la CPU.

```
clresourcegroup create -p RG_SLM_TYPE=automated \
 [-p RG_SLM_CPU_SHARES=valor] resource_group_name
```

`-p RG_SLM_TYPE=automated` Permite controlar el uso de la CPU y automatiza algunos pasos para configurar el sistema operativo Oracle Solaris con el fin de administrar recursos del sistema.

`-p RG_SLM_CPU_SHARES= valor` Especifica el número de recursos compartidos de CPU asignados al proyecto del grupo específico de recursos (`project.cpu-shares`) y determina el número de recursos compartidos de CPU asignados al nodo sin votación de clúster global (`zone.cpu-shares`).

*nombre\_grupo\_recursos* Especifica el nombre del grupo de recursos.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos.

No puede definir `RG_SLM_TYPE` como `automated` en un nodo sin votación si otra agrupación que no sea la predeterminada está en la configuración de la zona o si la zona está vinculada a otra agrupación que no sea la predeterminada. Consulte las páginas de comando [man zonecfg\(1M\)](#) y [poolbind\(1M\)](#) para obtener más información sobre la configuración de zonas y la vinculación de agrupaciones, respectivamente. Visualice la configuración de la zona como se indica a continuación:

```
zonecfg -z zone_name info pool
```

---

**Nota** – Se ha configurado un recurso, como `HASStoragePlus` o `LogicalHostname`, para que se inicie en un nodo sin votación de clúster global, pero con la propiedad `GLOBAL_ZONE` establecida en `TRUE` se inicia en el nodo de votación. Aunque establezca la propiedad `RG_SLM_TYPE` en `automated` este recurso no se beneficia de la configuración de los recursos compartidos de CPU y se lo considera como si estuviera en un grupo de recursos con `RG_SLM_TYPE` establecida en `manual`.

---

En este procedimiento, no se configura la propiedad `RG_SLM_PSET_TYPE`. Oracle Solaris Cluster usa el conjunto de procesadores predeterminado.

## 5 Active el cambio de configuración.

```
clresourcegroup online -M resource_group_name
```

*nombre\_grupo\_recursos*      Especifica el nombre del grupo de recursos.

Si establece `RG_SLM_PSET_TYPE` en `default`, Oracle Solaris Cluster crea una agrupación, `SCSLM_pool_nombre_zona`, pero no un conjunto de procesadores. En este caso, `SCSLM_pool_nombre_zona` se asocia con el conjunto de procesadores predeterminado.

Si los grupos de recursos en línea ya no están configurados para el control de la CPU en un nodo sin votación, el valor de los recursos compartidos de CPU para el nodo sin votación asume el valor de `zone.cpu-shares` de la configuración de la zona. Este parámetro tiene un valor predeterminado de 1. Para obtener más información sobre la configuración de zonas, consulte la página de comando [man zonecfg\(1M\)](#).

---

**Nota** – No elimine ni modifique el proyecto `SCSLM_nombre_grupo_recursos`. Puede agregar más control de recursos de forma manual al proyecto, por ejemplo configurando la propiedad `project.max-lwps`. Para obtener más información, consulte la página de comando [man projmod\(1M\)](#).

---

## ▼ Control del uso de la CPU en un nodo sin votación de clúster global con un conjunto de procesadores dedicado

Siga este procedimiento si desea que el grupo de recursos se ejecute en un conjunto de procesadores dedicado.

Si un grupo de recursos se configura para ejecutarse en un conjunto de procesadores dedicado, el software Oracle Solaris Cluster realiza las tareas siguientes al iniciar un recurso del grupo en un nodo sin votación de clúster global:

- Crea una agrupación denominada `SCSLM_pool_nombre_zona`, si todavía no se ha hecho.

- Crea un conjunto de procesadores dedicado. El tamaño del conjunto de procesadores se determina mediante las propiedades `RG_SLM_CPU_SHARES` y `RG_SLM_PSET_MIN`.
- Asocia la agrupación `SCSLM_pool_nombre_zona` con el conjunto de procesadores creado.
- Vincula de forma dinámica el nodo sin votación con la agrupación `SCSLM_pool_nombre_zona`.
- Aumenta el número de recursos compartidos de CPU asignados al nodo sin votación con el número específico de recursos compartidos de CPU, si todavía no se ha hecho.
- Crea un proyecto con el nombre `SCSLM_nombre_grupo_recursos` en el nodo sin votación, si todavía no se ha hecho. Este proyecto es específico del grupo de recursos y se le asigna el número especificado de recursos compartidos de CPU (`proyecto.recursos_compartidos_cpu`).
- Inicia el recurso en el proyecto de `SCSLM_nombre_grupo_recursos`.

## 1 Configure FSS como tipo de programador para el sistema.

```
dispadmin -d FSS
```

FSS se convierte en el programador a partir del siguiente rearranque. Para que esta configuración se aplique inmediatamente, use el comando `priocntl`.

```
priocntl -s -C FSS
```

Mediante la combinación de los comandos `priocntl` y `dispadmin`, se asegura de que FSS se convierta en el programador predeterminado inmediatamente y de que permanezca después del rearranque. Para obtener más información sobre cómo configurar un tipo de planificación, consulte las páginas de comando `man dispadmin(1M)` y `priocntl(1)`.

---

**Nota** – Si FSS no es el programador predeterminado, la asignación de recursos compartidos de CPU no se lleva a cabo.

---

## 2 En todos los nodos que usen un control de la CPU, configure el número de recursos compartidos para el nodo de votación de clúster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.

La configuración de estos parámetros protege los procesos que se estén ejecutando en el nodo de votación de competir por las CPU con los procesos que se estén ejecutando en nodos sin votación. Si no asigna un valor a las propiedades `globalzoneshares` y `defaultpsetmin`, asumen sus valores predeterminados.

```
clnode set [-p globalzoneshares=integer] \
[-p defaultpsetmin=integer] \
node
```

```
-p defaultpsetmin= minintero_p_establecido_predeterminado
```

Establece el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado. El número predeterminado es 1.

-p globalzoneshares= *entero*

Configura el número de recursos compartidos asignados al nodo de votación. El número predeterminado es 1.

*nodo*

Identifica los nodos donde se deben configurar propiedades.

Al configurar estas propiedades, también se establecen las del nodo de votación.

### 3 Compruebe si ha configurado estas propiedades correctamente:

```
clnode show node
```

Para el nodo que especifique, el comando `clnode` imprime las propiedades establecidas y los valores definidos para dichas propiedades. Si no configura las propiedades de control de la CPU con `clnode`, asumen el valor predeterminado.

### 4 Configure la función de control de la CPU.

```
clresourcegroup create -p RG_SLM_TYPE=automated \
 [-p RG_SLM_CPU_SHARES=value] \
-p -y RG_SLM_PSET_TYPE=value \
[-p RG_SLM_PSET_MIN=value] resource_group_name
```

-p RG\_SLM\_TYPE=automated      Permite controlar el uso de la CPU y automatiza algunos pasos para configurar el sistema operativo Oracle Solaris con el fin de administrar recursos del sistema.

-p RG\_SLM\_CPU\_SHARES= *valor*      Especifica el número de recursos compartidos de CPU asignados al proyecto específico del grupo de recursos (`project.cpu-shares`) y determina el número de recursos compartidos de CPU asignados al nodo sin votación (`zone.cpu-shares`) y el número máximo de procesadores en un conjunto.

-p RG\_SLM\_PSET\_TYPE= *valor*      Habilita la creación de un conjunto de procesadores dedicado. Para tener un conjunto de procesadores dedicado, puede establecer esta propiedad como `strong` o `weak`. Los valores `strong` y `weak` se excluyen mutuamente. Es decir, no puede configurar grupos de recursos en la misma zona para que algunos sean `strong` y otros `weak`.

-p RG\_SLM\_PSET\_MIN= *valor*      Determina el número mínimo de procesadores en el conjunto.

*nombre\_grupo\_recursos*      Especifica el nombre del grupo de recursos.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos.

No puede definir `RG_SLM_TYPE` como `automated` en un nodo sin votación si otra agrupación que no sea la predeterminada está en la configuración de la zona o si la zona está vinculada a otra



agrupación que no sea la predeterminada. Consulte las páginas de comando man [zonecfg\(1M\)](#) y [poolbind\(1M\)](#) para obtener más información sobre la configuración de zonas y la vinculación de agrupaciones, respectivamente. Visualice la configuración de la zona como se indica a continuación:

```
zonecfg -z zone_name info pool
```

---

**Nota** – Se ha configurado un recurso, como `HASStoragePlus` o `LogicalHostname`, para que se inicie en un nodo sin votación de clúster global, pero con la propiedad `GLOBAL_ZONE` establecida en `TRUE` se inicia en el nodo de votación. Incluso si establece la propiedad `RG_SLM_TYPE` en `automated`, este recurso no se beneficia de los recursos compartidos de CPU ni de la configuración del conjunto de procesadores dedicado y se lo considera como si fuera parte de un grupo de recursos con la propiedad `RG_SLM_TYPE` establecida en `manual`.

---

## 5 Active el cambio de configuración.

*nombre\_grupo\_recursos*      Especifica el nombre del grupo de recursos.

---

**Nota** – No elimine ni modifique el proyecto `SCSLM_nombre_grupo_recursos`. Puede agregar más control de recursos de forma manual al proyecto, por ejemplo configurando la propiedad `project.max-lwps`. Para obtener más información, consulte la página de comando man [projmod\(1M\)](#).

---

Los cambios realizados en `RG_SLM_CPU_SHARES` y `RG_SLM_PSET_MIN` aunque el grupo de recursos esté en línea se tienen en cuenta de forma dinámica. Sin embargo, si `RG_SLM_PSET_TYPE` está establecida en `strong` y no hay CPU suficientes disponibles para aceptar el cambio, no se aplica el cambio solicitado para `RG_SLM_PSET_MIN`. En este caso, aparece un mensaje de advertencia. En la siguiente conmutación, puede haber errores por CPU insuficientes si no hay bastantes CPU disponibles para confirmar los valores configurados para `RG_SLM_PSET_MIN`.

Si un grupo de recursos en línea ya no está configurado para el control de la CPU en el nodo sin votación, el valor de los recursos compartidos de CPU para el nodo sin votación asume el valor de `zone.cpu-shares`. Este parámetro tiene un valor predeterminado de 1.



# Aplicación de parches en el software y el firmware de Oracle Solaris Cluster

---

Este capítulo proporciona los procedimientos para agregar y quitar parches en una configuración de Oracle Solaris Cluster de las secciones siguientes.

- [“Información general sobre aplicación de parches de Oracle Solaris Cluster” en la página 315](#)
- [“Aplicación de parches del software Oracle Solaris Cluster” en la página 317](#)

## Información general sobre aplicación de parches de Oracle Solaris Cluster

Debido a la naturaleza de un clúster, para que el clúster funcione adecuadamente todos sus nodos deben estar en el mismo nivel de parche. En ocasiones, al aplicar parches en un nodo con un parche de Oracle Solaris Cluster, antes de instalarlo temporalmente es posible que se deba eliminar la pertenencia de un nodo al clúster o detener el clúster completo. Esta sección describe estos pasos.

Antes de aplicar un parche de Oracle Solaris Cluster, consulte el archivo README del parche. Asimismo, compruebe los requisitos de actualización para determinar el método de parche que necesitan los dispositivos de almacenamiento.

---

**Nota** – En lo concerniente a los parches de Oracle Solaris Cluster, consulte siempre el archivo README del parche y SunSolve para obtener instrucciones que reemplacen procedimientos de este capítulo.

---

La instalación del parche en todos los nodos del clúster se puede describir con una de las situaciones siguientes:

Parche de rearranque (nodo)

Antes de poder aplicar el parche o el firmware, se debe arrancar un nodo en modo monousuario con el comando `boot -sx` o `shutdown -g -y -i0`; a continuación, se debe

rearrancar para que se una al clúster. En primer lugar, el nodo debe ponerse en estado "silencioso" conmutando todos los grupos de recursos o de dispositivos del nodo al que se vaya a aplicar parches como miembro de otro clúster. Asimismo, aplique el parche o el firmware nodo a nodo para evitar cerrar todo el clúster.

El clúster permanece disponible mientras se aplica este tipo de parche aunque los nodos individuales no estén disponibles temporalmente. Un nodo al que se han aplicado parches se puede volver a unir a un clúster como miembro aunque otros nodos no estén todavía en el mismo nivel de parche.

#### Parche de rearranque (clúster)

Para aplicar el parche del software o firmware, el clúster debe detenerse y todos los nodos deben arrancarse en modo monousuario con el comando `boot -sx` o `shutdown -g -y -i0`. A continuación, rearranque los nodos para que vuelvan a unirse al clúster. Para este tipo de parches, el clúster no está disponible cuando se aplican los parches.

#### Parche que no es de rearranque

Al aplicar el parche, un nodo no debe estar en estado "silencioso" (todavía puede estar controlando grupos de recursos o de dispositivos) ni debe rearrancarse. Ahora bien, el parche debe aplicarse a los nodos de uno en uno y comprobar que el parche funcione antes de aplicar el parche a otro nodo.

---

**Nota** – Los protocolos del clúster subyacente no cambian debido a un parche.

---

Use el comando `patchadd` para aplicar un parche al clúster y `patchrm` para eliminar un parche (si es posible).

## Sugerencias sobre los parches de Oracle Solaris Cluster

Las sugerencias siguientes ayudan a administrar los parches de Oracle Solaris Cluster de forma más eficaz:

- Antes de aplicar el parche, lea siempre el archivo README del parche.
- Compruebe los requisitos de actualización de los dispositivos de almacenamiento para determinar el método de parche que se necesita.

- Aplique todos los parches (necesarios y recomendados) antes de ejecutar el clúster en un entorno de producción.
- Compruebe los niveles de hardware del firmware e instale todas las actualizaciones de firmware que puedan ser necesarias.
- Todos los nodos que actúen como miembros del clúster deben tener los mismos parches.
- Mantenga actualizados los parches de subsistema del clúster. Entre otros, estos parches incluyen la administración de volúmenes, el firmware de dispositivos de almacenamiento y el transporte del clúster.
- Revise con regularidad los informes sobre parches, por ejemplo una vez por trimestre, y aplique el conjunto de parches recomendados en la configuración de Oracle Solaris Cluster.
- Aplique parches selectivos conforme a las recomendaciones de Enterprise Services.
- Pruebe la migración tras error tras haber aplicado actualizaciones de parches importantes. Prepárese para retirar el parche si disminuye el rendimiento del clúster o si no funciona correctamente.

## Aplicación de parches del software Oracle Solaris Cluster

TABLA 11-1 Mapa de tareas: aplicar parches en el clúster

Tarea	Instrucciones
Aplique parches de Oracle Solaris Cluster que no sean de re arranque nodo por nodo y sin detener el nodo	<a href="#">“Aplicación de un parche de Oracle Solaris Cluster que no sea de re arranque” en la página 325</a>
Aplique un parche de re arranque de Oracle Solaris Cluster tras poner el miembro del clúster en un modo que no sea de clúster	<a href="#">“Aplicación de un parche de re arranque (nodo)” en la página 317</a> <a href="#">“Aplicación de un parche de re arranque (clúster)” en la página 322</a>
Aplique un parche en modo monousuario si el clúster tiene nodos de migración tras error	<a href="#">“Aplicación de parches en modo monousuario con nodos de migración tras error” en la página 326</a>
Quitar un parche de Oracle Solaris Cluster	<a href="#">“Cambio de un parche de Oracle Solaris Cluster” en la página 330</a>

### ▼ Aplicación de un parche de re arranque (nodo)

Aplique el parche nodo a nodo del clúster para mantenerlo operativo durante el proceso de aplicación de parches. Con este procedimiento, antes de aplicar el parche primero debe detener el nodo del clúster y arrancarlo en modo monousuario con el comando `boot -sx` o `shutdown -g -y -i0`.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Antes de aplicar el parche, visite el sitio web del producto de Oracle Solaris Cluster y compruebe si hay instrucciones especiales sobre preinstalación o postinstalación.**
- 2 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.admin` en el nodo en el que va a aplicar el parche.**
- 3 **Enumere los grupos de recursos y dispositivos del nodo al que va a aplicar el parche.**

```
clresourcegroup status -n node
cldevicegroup status -n node
```

- 4 **Conmute todos los grupos de recursos, recursos y grupos de dispositivos del nodo al que se va a aplicar el parche a otros miembros del clúster global.**

```
clnode evacuate -n node
```

`evacuate` Evacua todos los grupos de dispositivos y de recursos, incluidos todos los nodos sin votación de clúster global.

`-n node` Especifica el nodo desde el que está conmutando los grupos de recursos y de dispositivos.

- 5 **Cierre el nodo.**

```
shutdown -g0 [-y]
[-i0]
```

- 6 **Arranque el nodo en un modo monousuario que no sea de clúster.**

- En los sistemas basados en SPARC, ejecute el comando siguiente.

```
ok boot -sx
```

- En los sistemas basados en x86, ejecute los comandos siguientes.

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

Press any key to continue

- a. **En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba `e` para editar los comandos.**

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
+-----+
```

```
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

Para obtener más información sobre el arranque basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- b. En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba e para editarla.**

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.
```

- c. Agregue -x al comando para especificar que el sistema arranque en un modo que no sea de clúster.**

```
[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/multiboot -sx
```

- d. Pulse la tecla Intro para aceptar el cambio y volver a la pantalla de los parámetros de arranque.**

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -sx |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

- e. Escriba b para arrancar el nodo en un modo que no sea de clúster.**

---

**Nota** – Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en el modo que no es de clúster, realice estos pasos de nuevo para agregar la opción `-sx` al comando del parámetro de arranque del núcleo.

---

**7 Aplique el parche de software o firmware.**

```
patchadd -M patch-dir patch-id
```

*dir\_parche*               Especifica la ubicación del directorio del parche.

*id\_parche*               Especifica el número de parche de un parche concreto.

---

**Nota** – Siga siempre las instrucciones del directorio de parches, que sustituyen los procedimientos de este capítulo.

---

**8 Compruebe que el parche se haya instalado correctamente.**

```
showrev -p | grep patch-id
```

**9 Rearranque el nodo en el clúster.**

```
reboot
```

**10 Compruebe que funcione el parche, y que el nodo y el clúster funcionen correctamente.**

**11 Repita del [Paso 2](#) al [Paso 10](#) con todos los demás nodos del clúster.**

**12 Conmute los grupos de recursos y de dispositivos como sea necesario.**

Después de haber re arrancado todos los nodos, el último de ellos no tiene los grupos de recursos ni de dispositivos en línea.

```
cldevicegroup switch -n node + | devicegroup ...
```

```
clresourcegroup switch -n node[:zone][,...] + | resource-group ...
```

*nodo*           El nombre del nodo al que está conmutando los grupos de recursos y de dispositivos.

*zona*           El nombre del nodo sin votación de clúster global (node) que puede controlar el grupo de recursos. Especifique la zona sólo si ha proporcionado un nodo sin votación al crear el grupo de recursos.

**13 Compruebe si debe confirmar el software del parche mediante el comando `scversions`.**

```
/usr/cluster/bin/scversions
```

Aparece uno de los resultados siguientes:

```
Upgrade commit is needed.
```



Upgrade commit is NOT needed. All versions match.

#### 14 Si se necesita confirmación, confirme el software del parche.

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, al menos hay una reconfiguración de CMM; depende de la situación.

---

### Ejemplo 11–1 Aplicación de un parche de re arranque (nodo)

El ejemplo siguiente muestra la aplicación de un parche de re arranque de Oracle Solaris Cluster en un nodo.

```
clresourcegroup status -n rg1
...Resource Group Resource

rg1 rs-2
rg1 rs-3
...
cldevicegroup status -n nodedg-schost-1
...
Device Group Name: dg-schost-1
...
clnode evacuate phys-schost-2
shutdown -g0 -y -i0
...
```

Arranque el nodo en un modo monousuario que no sea de clúster.

- SPARC: tipo:
 

```
ok boot -sx
```
- x86: arranque el nodo en un modo monousuario que no sea de clúster. Consulte los pasos de arranque en el procedimiento siguiente.

```
patchadd -M /var/tmp/patches 234567-05
...
showrev -p | grep 234567-05
...
reboot
...
cldevicegroup switch -n phys-schost-1 dg-schost-1
clresourcegroup switch -n phys-schost-1 schost-sa-1
scversions
Upgrade commit is needed.
scversions -c
```

**Véase también** Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 330.](#)

## ▼ Aplicación de un parche de rearranque (clúster)

Con este procedimiento, antes de aplicar el parche debe detener el clúster y arrancar cada nodo en modo monousuario con el comando `boot -sx` o `shtudown -g -y -i0`.

- 1 **Antes de aplicar el parche, visite el sitio web del producto de Oracle Solaris Cluster y compruebe si hay instrucciones especiales sobre preinstalación o postinstalación.**
- 2 **Conviértase en superusuario en un nodo de clúster.**

### 3 Cierre el clúster.

```
cluster shutdown -y -g grace-period "message"
```

`-y` Especifica responder *sí* al indicador de confirmación.

`-g período_gracia` Especifica, en segundos, el tiempo de espera antes de cerrar. El período de gracia predeterminado es de 60 segundos.

`mensaje` Especifica el mensaje de advertencia que transmitir. Use comillas si el `mensaje` contiene varias palabras.

### 4 Arranque todos los nodos en un modo monousuario que no sea de clúster.

En la consola de cada nodo, ejecute los comandos siguientes.

- En los sistemas basados en SPARC, ejecute el comando siguiente.

```
ok boot -sx
```

- En los sistemas basados en x86, ejecute los comandos siguientes.

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

Press any key to continue

#### a. En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba e para editar los comandos.

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
```

```
+-----+
| Solaris 10 /sol_10_x86
| Solaris failsafe
|
+-----+
```

Use the ^ and v keys to select which entry is highlighted.  
Press enter to boot the selected OS, 'e' to edit the  
commands before booting, or 'c' for a command-line.

Para obtener más información sobre el arranque basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- b. En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba **e** para editarla.

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.
```

- c. Agregue **-x** al comando para especificar que el sistema arranque en un modo que no sea de clúster.

[ 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/multiboot -sx
```

- d. Pulse la tecla **Intro** para aceptar el cambio y volver a la pantalla de los parámetros de arranque.

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -sx |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

- e. Escriba **b** para arrancar el nodo en un modo que no sea de clúster.

---

**Nota** – Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en el modo que no es de clúster, realice estos pasos de nuevo para agregar la opción **-sx** al comando del parámetro de arranque del núcleo.

---

## 5 Aplique el parche de software o firmware.

Ejecute el comando siguiente, nodo por nodo.

```
patchadd -M patch-dir patch-id
```

<i>dir_parche</i>	Especifica la ubicación del directorio del parche.
<i>id_parche</i>	Especifica el número de parche de un parche concreto.

---

**Nota** – Siga siempre las instrucciones del directorio de parches, que sustituyen los procedimientos de este capítulo.

---

**6 Compruebe que el parche se haya instalado correctamente en todos los nodos.**

```
showrev -p | grep patch-id
```

**7 Después de aplicar el parche en todos los nodos, re arranque los nodos en el clúster.**

Ejecute el comando siguiente en todos los nodos.

```
reboot
```

**8 Compruebe si debe confirmar el software del parche con el comando `scversions`.**

```
/usr/cluster/bin/scversions
```

Aparece uno de los resultados siguientes:

Upgrade commit is needed.

Upgrade commit is NOT needed. All versions match.

**9 Si se necesita confirmación, confirme el software del parche.**

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, al menos hay una reconfiguración de CMM; depende de la situación.

---

**10 Compruebe que funcione el parche, y que los nodos y el clúster funcionen correctamente.**

**Ejemplo 11–2 Aplicación de un parche de re arranque (clúster)**

El ejemplo siguiente muestra la aplicación de un parche de re arranque de Oracle Solaris Cluster en un clúster.

```
cluster shutdown -g0 -y
...
```

Arranque el clúster en un modo monousuario que no sea de clúster.

- SPARC: tipo:

```
ok boot -sx
```

- x86: arranque todos los nodos en un modo monousuario que no sea de clúster. Consulte los pasos en el siguiente procedimiento.

```
...
patchadd -M /var/tmp/patches 234567-05
(Apply patch to other cluster nodes)
...
showrev -p | grep 234567-05
reboot
scversions
Upgrade commit is needed.
scversions -c
```

**Véase también** Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 330.](#)

## ▼ Aplicación de un parche de Oracle Solaris Cluster que no sea de re arranque

Aplique el parche nodo por nodo en el clúster. Si aplica un parche que no es de reinicio, el nodo que recibe el parche no se debe re arrancar primero.

- 1 Antes de aplicar el parche, visite la página web del producto de Oracle Solaris Cluster y compruebe si hay instrucciones especiales sobre preinstalación o postinstalación.

- 2 Aplique el parche en un solo nodo.

```
patchadd -M patch-dir patch-id
```

*dir\_parche* Especifica la ubicación del directorio del parche.

*id\_parche* Especifica el número de parche de un parche concreto.

- 3 Compruebe que el parche se haya instalado correctamente.

```
showrev -p | grep patch-id
```

- 4 Compruebe que funcione el parche, y que el nodo y el clúster funcionen correctamente.

- 5 Repita del [Paso 2](#) al [Paso 4](#) con todos los demás nodos del clúster.

- 6 Compruebe si debe confirmar el software del parche con el comando `scversions`.

```
/usr/cluster/bin/scversions
```

Aparece uno de los resultados siguientes:

Upgrade commit is needed.

Upgrade commit is NOT needed. All versions match.

**7 Si se necesita confirmación, confirme el software del parche.**

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, al menos hay una reconfiguración de CMM; depende de la situación.

---

**Ejemplo 11–3** Aplicación de un parche de Oracle Solaris Cluster que no sea de rearranque

```
patchadd -M /tmp/patches 234567-05
...
showrev -p | grep 234567-05
scversions
Upgrade commit is needed.
scversions -c
```

**Véase también** Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 330.](#)

## ▼ Aplicación de parches en modo monousuario con nodos de migración tras error

Realice esta tarea para aplicar parches en modo monousuario con nodos de migración tras error. Este método de parche es necesario si utiliza los servicios de datos de Oracle Solaris Cluster para contenedores Solaris en una configuración de migración tras error con el software Oracle Solaris Cluster.

- 1 Compruebe que el dispositivo de quórum no se configure para uno de los LUN usado como almacenamiento compartido perteneciente a los grupos de discos que contienen la ruta de zona que se usa de forma manual en este procedimiento.**
  - a. Determine si el dispositivo de quórum se utiliza en los grupos de discos que contienen las rutas de zona y si el dispositivo de quórum usa las reservas de SCSI2 o de SCSI3.**

```
clquorum show
```

- b. Si el dispositivo de quórum está entre un LUN de los grupos de discos, agregue un LUN nuevo como dispositivo de quórum que no forme parte de ningún conjunto de discos que contengan la ruta de zona.

```
clquorum add new-didname
```

- c. Elimine el dispositivo de quórum antiguo.

```
clquorum remove old-didname
```

- d. Si las reservas de SCSI2 se usan para el dispositivo de quórum antiguo, quite la reserva de SCSI2 de quórum antiguo y compruebe que no queden reservas de SCSI2.

Para saber cómo ejecutar los comandos `pgre`, es necesario instalar y utilizar el paquete Diagnostic Toolkit (SUNWscdtk) proporcionado por el equipo de asistencia de Oracle.

---

**Nota** – Si involuntariamente ha quitado las claves de reserva del dispositivo de quórum activo, elimine y vuelva a agregar el dispositivo de quórum para poner claves de reserva nuevas en dicho dispositivo.

---

- 2 Evacue el nodo al que desea aplicar parches.

```
clresourcegroup evacuate -n node1
```

- 3 Ponga fuera de línea los grupos de recursos que contengan recursos de contenedor Solaris de alta disponibilidad.

```
clresourcegroup offline resourcegroupname
```

- 4 Inhabilite todos los recursos de los grupos de recursos puestos fuera de línea.

```
clresource disable resourcename
```

- 5 Anule la administración de los grupos de recursos puestos fuera de línea.

```
clresourcegroup unmanage resourcegroupname
```

- 6 Ponga fuera de línea el grupo o los grupos de dispositivos correspondientes.

```
cldevicegroup offline cldevicegroupname
```

- 7 Inhabilite los grupos de dispositivos que haya puesto fuera de línea.

```
cldevicegroup disable devicegroupname
```

- 8 Arranque el nodo pasivo fuera del clúster.

```
reboot -- -x
```

- 9 Antes de continuar, compruebe que los métodos de inicio de SMF hayan finalizado en el nodo pasivo.

```
svcs -x
```

**10 Compruebe que haya finalizado el proceso de reconfiguración del nodo activo.**

```
cluster status
```

**11 Determine si hay reservas de SCSI-2 en el disco del conjunto de discos y libere las claves. Siga estas instrucciones para determinar si hay reservas de SCSI-2; a continuación, libérelas.**

- En todos los discos del conjunto, ejecute el comando: `/usr/cluster/lib/sc/scsi -c disfailfast -d /dev/did/rdisk/d#s2`.
- Si las claves están enumeradas, libérelas ejecutando el comando:  
`/usr/cluster/lib/sc/scsi -c release -d /dev/did/rdisk/d#s2`.

Cuando termine de liberar las claves de reserva, omita el Paso 12 y siga en el Paso 13.

**12 Determine si hay reservas de SCSI-3 en los discos de los grupos de discos.**

**a. Ejecute el comando siguiente en todos los discos de los conjuntos de discos.**

```
/usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/didnames2
```

**b. Si las claves están enumeradas, quítelas.**

```
/usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/didnames2
```

**13 Pase a ser propietario del metaconjunto en el nodo pasivo.**

```
metaset -s disksetname -C take -f
```

**14 Monte el sistema o los sistemas de archivos que contenga la ruta de zona en el nodo pasivo.**

```
mount device mountpoint
```

**15 Cambie al modo monousuario en el nodo pasivo.**

```
init s
```

**16 Detenga todas las zonas arrancadas posibles que no estén bajo el servicio de datos de Oracle Solaris Cluster para el control del contenedor Solaris.**

```
zoneadm -z zonename halt
```

**17 (Opcional) Si desinstala varios parches, por motivos de rendimiento puede elegir arrancar todas las zonas configuradas en modo monousuario.**

```
zoneadm -z zonename boot -s
```

**18 Aplique los parches.**



- 19 Rearranque el nodo y espere hasta que hayan finalizado todos los métodos de inicio SMF. Ejecute el comando `svcs -a` sólo después de haber rearrancado el nodo.  

```
reboot
```

```
svcs -a
```

El primer nodo ya está preparado.
- 20 Evacue el segundo nodo al que desea aplicar parches.  

```
clresourcegroup evacuate -n node2
```
- 21 Repita los pasos del 8 al 13 para el segundo nodo.
- 22 Desconecte las zonas a las que ya haya aplicado parches para acelerar el proceso.  

```
zoneadm -z zonename detach
```
- 23 Cambie al modo monousuario en el nodo pasivo.  

```
init s
```
- 24 Detenga todas las zonas arrancadas posibles que no estén bajo el servicio de datos de Oracle Solaris Cluster para el control del contenedor Solaris.  

```
zoneadm -z zonename halt
```
- 25 (Opcional) Si desinstala varios parches, por motivos de rendimiento puede elegir arrancar todas las zonas configuradas en modo monousuario.  

```
zoneadm -z zonename boot -s
```
- 26 Aplique los parches.
- 27 Conecte las zonas que haya desconectado.  

```
zoneadm -z zonename attach -F
```
- 28 Rearranque el nodo en modo de clúster.  

```
reboot
```
- 29 Ponga en línea el grupo o los grupos de dispositivos.
- 30 Inicie los grupos de recursos.
- 31 Compruebe si debe confirmar el software del parche con el comando `scversions`.  

```
/usr/cluster/bin/scversions
```

Aparece uno de los resultados siguientes:

Upgrade commit is needed.

Upgrade commit is NOT needed. All versions match.

**32 Si se necesita confirmación, confirme el software del parche.**

**# scversions -c**

---

**Nota** – Al ejecutar `scversions`, al menos hay una reconfiguración de CMM; depende de la situación.

---

## Cambio de un parche de Oracle Solaris Cluster

Para quitar un parche de Oracle Solaris Cluster al clúster, primero debe eliminar el parche nuevo de Oracle Solaris Cluster y, a continuación, volver a aplicar el parche anterior o la versión actualizada. Para quitar el parche de Oracle Solaris Cluster, consulte los procedimientos siguientes. Para volver a aplicar un parche de Oracle Solaris Cluster, consulte uno de los procedimientos siguientes:

- “Aplicación de un parche de re arranque (nodo)” en la página 317
- “Aplicación de un parche de re arranque (clúster)” en la página 322
- “Aplicación de un parche de Oracle Solaris Cluster que no sea de re arranque” en la página 325

---

**Nota** – Antes de aplicar un parche de Oracle Solaris Cluster, consulte el archivo README del parche.

---

### ▼ Eliminación de un parche de Oracle Solaris Cluster que no sea de re arranque

- 1 Conviértase en superusuario en un nodo de clúster.
- 2 Elimine el parche que no sea de re arranque.

**# patchrm *patchid***

### ▼ Eliminación de un parche de Oracle Solaris Cluster que sea de re arranque

- 1 Conviértase en superusuario en un nodo de clúster.

- 2 Arranque el clúster en un modo que no sea de clúster. Para obtener más información sobre cómo arrancar un nodo en un modo que no sea de clúster, consulte [“Rearranque de un nodo en un modo que no sea de clúster” en la página 79](#).
- 3 Elimine el parche de re arranque.  
`# patchrm patchid`
- 4 Rearranque el clúster en modo de clúster.  
`# reboot`
- 5 Repita los pasos del 2 al 4 para todos los nodos del clúster.



# Copias de seguridad y restauraciones de clústers

Este capítulo contiene las secciones siguientes.

- “Copia de seguridad de un clúster” en la página 333
- “Restauración de los archivos del clúster” en la página 345

## Copia de seguridad de un clúster

TABLA 12-1 Mapa de tareas: hacer una copia de seguridad de los archivos del clúster

Tarea	Instrucciones
Buscar los nombres de los sistemas de archivos de los que desee realizar una copia de seguridad	“Búsqueda de los nombres de sistema de archivos para hacer copias de seguridad” en la página 334
Calcular las cintas necesarias para guardar una copia de seguridad completa	“Determinación del número de cintas necesarias para una copia de seguridad completa” en la página 334
Hacer una copia de seguridad del sistema de archivos root	“Copias de seguridad del sistema de archivos root (/)” en la página 335
Realizar una copia de seguridad en línea para sistemas de archivos duplicados o plexados	“Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)” en la página 338   “Copias de seguridad en línea para volúmenes (Veritas Volume Manager)” en la página 341
Hacer una copia de seguridad de la configuración del clúster	“Copias de seguridad de la configuración del clúster” en la página 345
Realizar una copia de seguridad de la configuración de partición del disco de almacenamiento	Consulte la documentación del disco de almacenamiento

## ▼ Búsqueda de los nombres de sistema de archivos para hacer copias de seguridad

Siga este procedimiento para determinar los nombres de los sistemas de archivos de los que desea realizar una copia de seguridad.

- 1 **Visualice el contenido del archivo `/etc/vfstab`.**  
Para ejecutar este comando no se debe ser superusuario ni asumir una función equivalente.  
`# more /etc/vfstab`
- 2 **En la columna de puntos de montaje, busque el nombre del sistema de archivos del que efectúa una copia de seguridad.**  
Use este nombre cuando realice la copia de seguridad del sistema de archivos.  
`# more /etc/vfstab`

### Ejemplo 12-1 Búsqueda de nombres de sistemas de archivos para hacer copias de seguridad

El ejemplo siguiente muestra los nombres de los sistemas de archivos disponibles enumerados en el archivo `/etc/vfstab`.

```
more /etc/vfstab
#device device mount FS fsck mount mount
#to mount to fsck point type pass at boot options
#
#/dev/dsk/cl1d0s2 /dev/rdisk/cl1d0s2 /usr ufs 1 yes -
f - /dev/fd fd - no -
/proc - /proc proc - no -
/dev/dsk/cl1t6d0s1 - - swap - no -
/dev/dsk/cl1t6d0s0 /dev/rdisk/cl1t6d0s0 / ufs 1 no -
/dev/dsk/cl1t6d0s3 /dev/rdisk/cl1t6d0s3 /cache ufs 2 yes -
swap - /tmp tmpfs - yes -
```

## ▼ Determinación del número de cintas necesarias para una copia de seguridad completa

Siga este procedimiento para calcular el número de cintas que necesita para realizar la copia de seguridad de un sistema de archivos.

- 1 **Conviértase en superusuario o asuma una función equivalente en el nodo del clúster del que esté haciendo la copia de seguridad.**
- 2 **Calcule el tamaño de la copia de seguridad en bytes.**

```
ufsdump S filesystem
```

<b>S</b>	Muestra el cálculo del número de bytes que se necesitan para realizar la copia de seguridad.
<b>sistema_archivos</b>	Especifica el nombre del sistema de archivos del que desea hacer una copia de seguridad.

- 3 Divida el tamaño estimado por la capacidad de la cinta para saber la cantidad de cintas que necesita.**

### **Ejemplo 12-2** Determinación el número de cintas que se necesitan

En el ejemplo siguiente, el tamaño del sistema de archivos de 905.881.620 bytes cabe fácilmente en una cinta de 4 Gigabytes (905.881.620 ÷ 4.000.000.000).

```
ufsdump S /global/phys-schost-1
905881620
```

## ▼ **Copias de seguridad del sistema de archivos root (/)**

Siga estos pasos para realizar una copia de seguridad del sistema de archivos root (/) de un nodo del clúster. Antes de efectuar la copia de seguridad, compruebe que el clúster se ejecute sin errores.

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo del clúster del que va a realizar la copia de seguridad.**
- 2 Conmute a otro nodo del clúster todos los servicios de datos en ejecución del nodo del que se va a realizar la copia de seguridad.**

```
clnode evacuate node
```

**nodo** Especifica el nodo desde el que esté conmutando los grupos de recursos y de dispositivos.

- 3 Cierre el nodo.**  

```
shutdown -g0 -y -i0
```
- 4 Rearranque el nodo en un modo que no sea de clúster.**

- En los sistemas basados en SPARC, ejecute el comando siguiente.

```
ok boot -xs
```

- En los sistemas basados en x86, ejecute los comandos siguientes.

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

```
Press any key to continue
```

- En el menú de GRUB, utilice las teclas de flecha para seleccionar la correspondiente entrada de Oracle Solaris y escriba e para editar los comandos.**

El menú de GRUB que aparece es similar al siguiente:

```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86 |
| Solaris failsafe |
| |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```

Para obtener más información sobre el arranque basado en GRUB, consulte [“Booting an x86 Based System by Using GRUB \(Task Map\)”](#) de *System Administration Guide: Basic Administration*.

- En la pantalla de los parámetros de arranque, utilice las teclas de flecha para seleccionar la entrada de núcleo y escriba e para editarla.**

La pantalla de los parámetros de arranque de GRUB que aparece es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.
```

- Agregue -x al comando para especificar que el sistema se arranque en un modo que no sea de clúster.**

```
[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/multiboot -x
```



- d. Pulse la tecla Intro para aceptar el cambio y volver a la pantalla de los parámetros de arranque.

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -x |
| module /platform/i86pc/boot_archive |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.  
Press 'b' to boot, 'e' to edit the selected command in the  
boot sequence, 'c' for a command-line, 'o' to open a new line  
after ('O' for before) the selected line, 'd' to remove the  
selected line, or escape to go back to the main menu.-

- e. Escriba b para arrancar el nodo en un modo que no sea de clúster.

---

**Nota** – Este cambio en el comando del parámetro de arranque del núcleo no se conserva tras arrancar el sistema. La próxima vez que re arranque el nodo, lo hará en modo de clúster. Para arrancarlo en el modo que no es de clúster, realice estos pasos de nuevo para agregar la opción -x al comando del parámetro de arranque del núcleo.

---

- 5 Realice la copia de seguridad del sistema de archivos root (/) creando una instantánea de UFS.

- a. Compruebe que el sistema de archivos tenga espacio suficiente en el disco para el archivo de almacenamiento de seguridad.

```
df -k
```

- b. Compruebe que ya no exista el archivo de almacenamiento de seguridad con el mismo nombre y la misma ubicación.

```
ls /backing-store-file
```

- c. Cree la instantánea de UFS.

```
fssnap -F ufs -o bs=/backing-store-file /file-system
```

- d. Compruebe que la instantánea se haya creado.

```
/usr/lib/fs/ufs/fssnap -i /file-system
```

- 6 Realice una copia de seguridad de la instantánea del sistema de archivos.

```
ufsdump 0ucf /dev/rmt/0 snapshot-name
```

Por ejemplo:

```
ufsdump 0ucf /dev/rmt/0 /dev/rfssnap/1
```

- 7 Compruebe que la copia de seguridad de la instantánea se haya completado correctamente.

```
ufsrestore ta /dev/rmt/0
```

- 8 Rearranque el nodo en modo de clúster.

```
init 6
```

### Ejemplo 12-3 Copias de seguridad del sistema de archivos root (/)

En el ejemplo siguiente, se guarda una instantánea del sistema de archivos root (/) en /scratch/usr.back.file en el directorio /usr. ‘

```
fssnap -F ufs -o bs=/scratch/usr.back.file /usr
/dev/fssnap/1
```

## ▼ Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)

Se puede realizar una copia de seguridad de un volumen de Solaris Volume Manager sin desmontarlo o poner la duplicación completa fuera de línea. Una de las subduplicaciones debe ponerse temporalmente fuera de línea; si bien se pierde el duplicado, puede ponerse en línea y volver a sincronizarse en cuanto la copia de seguridad esté terminada sin detener el sistema ni denegar el acceso del usuario a los datos. El uso de duplicaciones para realizar copias de seguridad en línea crea una copia de seguridad que es una "instantánea" de un sistema de archivos activo.

Si un programa escribe datos en el volumen justo antes de ejecutarse el comando `lockfs` puede causar problemas. Para evitarlo, detenga temporalmente todos los servicios que se estén ejecutando en este nodo. Asimismo, antes de efectuar la copia de seguridad, compruebe que el clúster se ejecute sin errores.

`phys -s host#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función equivalente en el nodo del clúster del que esté haciendo la copia de seguridad.
- 2 Use el comando `metaset(1M)` para determinar el nodo que es propietario del volumen de la copia de seguridad.

```
metaset -s setname
```

-s                      Especifica el nombre del conjunto de discos.  
*nombre\_conjunto*

- 3 Use el comando **lockfs(1M)** con la opción -w para bloquear la escritura en el sistema de archivos.

```
lockfs -w mountpoint
```

---

**Nota** – El sistema de archivos debe bloquearse sólo si el sistema de archivos de UFS reside en el duplicado. Por ejemplo, si el volumen de Solaris Volume Manager se configura como un dispositivo básico para el software de administración de la base de datos u otra aplicación, no se debe usar el comando `lockfs`. Sin embargo, podría ejecutar la utilidad adecuada según el proveedor para purgar los búferes y bloquear el acceso.

---

- 4 Use el comando **metastat(1M)** para determinar los nombres de las subduplicaciones.

```
metastat -s setname -p
```

-p                      Muestra el estado en un formato similar al del archivo `md.tab`.

- 5 Use el comando **metadetach(1M)** para poner una subduplicación fuera de línea desde la duplicación.

```
metadetach -s setname mirror submirror
```

---

**Nota** – Las lecturas se siguen efectuando desde otras subduplicaciones. Sin embargo, la subduplicación fuera de línea pierde la sincronización al realizarse la primera escritura en la duplicación. Esta incoherencia se corrige al poner la subduplicación en línea de nuevo. El comando `fsck` no debe ejecutarse.

---

- 6 Desbloquee los sistemas y permita que se pueda seguir escribiendo mediante el comando **lockfs** con la opción -u.

```
lockfs -u mountpoint
```

- 7 Realice una comprobación del sistema de archivos.

```
fsck /dev/md/diskset/rdisk/submirror
```

- 8 Haga una copia de seguridad de la subduplicación fuera de línea en una cinta o en otro medio.

Use el comando **ufsdump(1M)** o la utilidad de copia de seguridad que use normalmente.

```
ufsdump 0ucf dump-device submirror
```

---

**Nota** – Use el nombre del dispositivo básico (`/rdsk`) para la subduplicación, en lugar del nombre del dispositivo de bloqueo (`/dsk`).

---

- 9 Use el comando **metattach(1M)** para volver a poner en línea el metadispositivo o el volumen.

```
metattach -s setname mirror submirror
```

Cuando se pone en línea un metadispositivo o un volumen, se vuelve a sincronizar automáticamente con la duplicación.

- 10 Use el comando **metastat** para comprobar que la subduplicación se vuelve a sincronizar.

```
metastat -s setname mirror
```

#### Ejemplo 12–4 Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)

En el ejemplo siguiente, el nodo del clúster `phys-schost-1` es el propietario del metaconjunto `schost-1`, de modo que el procedimiento de copia de seguridad se realiza desde `phys-schost-1`. La duplicación `/dev/md/schost-1/dsk/d0` consta de las subduplicaciones `d10`, `d20` y `d30`.

```
[Determine the owner of the metaset:]
metaset -s schost-1
Set name = schost-1, Set number = 1
Host Owner
 phys-schost-1 Yes
...
[Lock the file system from writes:]
lockfs -w /global/schost-1
[List the submirrors:]
metastat -s schost-1 -p
schost-1/d0 -m schost-1/d10 schost-1/d20 schost-1/d30 1
schost-1/d10 1 1 d4s0
schost-1/d20 1 1 d6s0
schost-1/d30 1 1 d8s0
[Take a submirror offline:]
metadetach -s schost-1 d0 d30
[Unlock the file system:]
lockfs -u /
[Check the file system:]
fsck /dev/md/schost-1/rdisk/d30
[Copy the submirror to the backup device:]
ufsdump 0ucf /dev/rmt/0 /dev/md/schost-1/rdisk/d30
DUMP: Writing 63 Kilobyte records
DUMP: Date of this level 0 dump: Tue Apr 25 16:15:51 2000
DUMP: Date of last level 0 dump: the epoch
DUMP: Dumping /dev/md/schost-1/rdisk/d30 to /dev/rdisk/c1t9d0s0.
...
DUMP: DUMP IS DONE
[Bring the submirror back online:]
metattach -s schost-1 d0 d30
schost-1/d0: submirror schost-1/d30 is attached
[Resynchronize the submirror:]
metastat -s schost-1 d0
schost-1/d0: Mirror
 Submirror 0: schost-0/d10
 State: Okay
 Submirror 1: schost-0/d20
```

```

State: Okay
Submirror 2: schost-0/d30
State: Resyncing
Resync in progress: 42% done
Pass: 1
Read option: roundrobin (default)
...

```

## ▼ Copias de seguridad en línea para volúmenes (Veritas Volume Manager)

Veritas Volume Manager identifica un volumen duplicado como bidireccional. Se puede realizar una copia de seguridad de una unidad sin desmontarla o poniendo el volumen entero fuera de línea. Este resultado se obtiene creando una copia de instantánea del volumen y realizando una copia de seguridad de este volumen temporal sin detener el sistema o denegar el acceso del usuario a los datos.

Antes de efectuar la copia de seguridad, compruebe que el clúster se ejecute sin errores.

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Inicie sesión en cualquier nodo del clúster y conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.admin` en el nodo primario actual para el grupo de discos del clúster.**
- 2 **Muestre en una lista la información del grupo de discos.**  
`# vxprint -g diskgroup`
- 3 **Determine qué nodo tiene el grupo de discos importado actualmente e indique que es el nodo primario para el grupo de discos.**  
`# cldevicegroup status`
- 4 **Cree una instantánea del volumen.**  
`# vxassist -g diskgroup snapstart volume`

---

**Nota** – Crear una instantánea puede tardar mucho tiempo, en función del tamaño del volumen.

---

**5 Asegúrese de que se ha creado el volumen nuevo.**

```
vxprint -g diskgroup
```

Cuando la instantánea haya finalizado, aparece un estado de Snapdone (instantánea hecha) en el campo State (estado) para el grupo de discos seleccionado.

**6 Detenga todos los servicios de datos que estén accediendo al sistema de archivos.**

```
clresourcegroup offline resource-group
```

---

**Nota** – Detenga todos los servicios de datos para asegurarse de que se realice la copia de seguridad del sistema de archivos de datos adecuadamente. Si no se está ejecutando ningún servicio de datos, no debe realizar los pasos [Paso 6](#) y [Paso 8](#).

---

**7 Cree un volumen de copia de seguridad con el nombre bkup-vol y conéctelo con el volumen de la instantánea.**

```
vxassist -g diskgroup snapshot volume bkup-vol
```

**8 Reinicie los servicios de datos que detuvo en el [Paso 6](#) con el comando clresourcegroup.**

```
clresourcegroup online -zone -n node resourcegroup
```

*nodo* El nombre del nodo.

*zona* El nombre del nodo sin votación de clúster global (*nodo*) que puede controlar el grupo de recursos. Especifique la *zona* únicamente si ha señalado un nodo sin votación al crear el grupo de recursos.

**9 Compruebe si el volumen está ya conectado al volumen nuevo bkup-vol.**

```
vxprint -g diskgroup
```

**10 Registre el cambio en la configuración del grupo de dispositivos.**

```
cldevicegroup sync diskgroup
```

**11 Compruebe el volumen de la copia de seguridad.**

```
fsck -y /dev/vx/rdisk/diskgroup/bkup-vol
```

**12 Realice una copia de seguridad para copiar el volumen bkup-vol en una cinta u otro medio.**

Use el comando `ufsdump(1M)` u otra utilidad de copia de seguridad que utilice normalmente.

```
ufsdump 0ucf dump-device /dev/vx/dsk/diskgroup/bkup-vol
```

**13 Elimine el volumen temporal.**

```
vxedit -rf rm bkup-vol
```

**14 Registre los cambios en la configuración del grupo de discos.**

```
cldevicegroup sync diskgroup
```

## Ejemplo 12-5 Copias de seguridad en línea para los volúmenes (Veritas Volume Manager)

En el ejemplo siguiente, el nodo del clúster `phys-schost-2` es el propietario principal del grupo de dispositivos `schost-1`. De modo que el proceso de copia de seguridad se realiza desde `phys-schost-2`. El volumen `/vol01` se copia y, a continuación, se asocia con un volumen nuevo: `bkup-vol`.

```
[Become superuser or assume a role that provides solaris.cluster.admin RBAC authorization on the primary node.]
```

```
[Identify the current primary node for the device group:]
```

```
cldevicegroup status
```

```
-- Device Group Servers --
```

	Device Group	Primary	Secondary
	-----	-----	-----
Device group servers:	rmt/1	-	-
Device group servers:	schost-1	phys-schost-2	phys-schost-1

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

	Device Group	Status
	-----	-----
Device group status:	rmt/1	Offline
Device group status:	schost-1	Online

```
[List the device group information:]
```

```
vxprint -g schost-1
```

TY	NAME	ASSOC	KSTATE	LENGTH	PLOFFS	STATE	TUTIL0	PUTIL0
dg	schost-1	schost-1	-	-	-	-	-	-
dm	schost-101	clt1d0s2	-	17678493	-	-	-	-
dm	schost-102	clt2d0s2	-	17678493	-	-	-	-
dm	schost-103	c2t1d0s2	-	8378640	-	-	-	-
dm	schost-104	c2t2d0s2	-	17678493	-	-	-	-
dm	schost-105	clt3d0s2	-	17678493	-	-	-	-
dm	schost-106	c2t3d0s2	-	17678493	-	-	-	-
v	vol01	gen	ENABLED	204800	-	ACTIVE	-	-
pl	vol01-01	vol01	ENABLED	208331	-	ACTIVE	-	-
sd	schost-101-01	vol01-01	ENABLED	104139	0	-	-	-
sd	schost-102-01	vol01-01	ENABLED	104139	0	-	-	-
pl	vol01-02	vol01	ENABLED	208331	-	ACTIVE	-	-
sd	schost-103-01	vol01-02	ENABLED	103680	0	-	-	-
sd	schost-104-01	vol01-02	ENABLED	104139	0	-	-	-
pl	vol01-03	vol01	ENABLED	LOGONLY	-	ACTIVE	-	-
sd	schost-103-02	vol01-03	ENABLED	5	LOG	-	-	-

```
[Start the snapshot operation:]
```

```
vxassist -g schost-1 snapstart vol01
```

```
[Verify the new volume was created:]
```

```
vxprint -g schost-1
```

TY	NAME	ASSOC	KSTATE	LENGTH	PLOFFS	STATE	TUTIL0	PUTIL0
dg	schost-1	schost-1	-	-	-	-	-	-
dm	schost-101	clt1d0s2	-	17678493	-	-	-	-
dm	schost-102	clt2d0s2	-	17678493	-	-	-	-
dm	schost-103	c2t1d0s2	-	8378640	-	-	-	-
dm	schost-104	c2t2d0s2	-	17678493	-	-	-	-
dm	schost-105	clt3d0s2	-	17678493	-	-	-	-
dm	schost-106	c2t3d0s2	-	17678493	-	-	-	-

```

v vol01 gen ENABLED 204800 - ACTIVE - -
pl vol01-01 vol01 ENABLED 208331 - ACTIVE - -
sd schost-101-01 vol01-01 ENABLED 104139 0 - - -
sd schost-102-01 vol01-01 ENABLED 104139 0 - - -
pl vol01-02 vol01 ENABLED 208331 - ACTIVE - -
sd schost-103-01 vol01-02 ENABLED 103680 0 - - -
sd schost-104-01 vol01-02 ENABLED 104139 0 - - -
pl vol01-03 vol01 ENABLED LOGONLY - ACTIVE - -
sd schost-103-02 vol01-03 ENABLED 5 LOG - - -
pl vol01-04 vol01 ENABLED 208331 - SNAPDONE - -
sd schost-105-01 vol01-04 ENABLED 104139 0 - - -
sd schost-106-01 vol01-04 ENABLED 104139 0 - - -
[Stop data services, if necessary:]
clresourcegroup offline nfs-rg
[Create a copy of the volume:]
vxassist -g schost-1 snapshot vol01 bkup-vol
[Restart data services, if necessary:]
clresourcegroup online -n phys-schost-1 nfs-rg
[Verify bkup-vol was created:]
vxprint -g schost-1
TY NAME ASSOC KSTATE LENGTH PLOFFS STATE TUTILO PUTILO
dg schost-1 schost-1 - - - - -
...

dm schost-101 clt1d0s2 - 17678493 - - -
...

v bkup-vol gen ENABLED 204800 - ACTIVE - -
pl bkup-vol-01 bkup-vol ENABLED 208331 - ACTIVE - -
sd schost-105-01 bkup-vol-01 ENABLED 104139 0 - - -
sd schost-106-01 bkup-vol-01 ENABLED 104139 0 - - -

v vol01 gen ENABLED 204800 - ACTIVE - -
pl vol01-01 vol01 ENABLED 208331 - ACTIVE - -
sd schost-101-01 vol01-01 ENABLED 104139 0 - - -
sd schost-102-01 vol01-01 ENABLED 104139 0 - - -
pl vol01-02 vol01 ENABLED 208331 - ACTIVE - -
sd schost-103-01 vol01-02 ENABLED 103680 0 - - -
sd schost-104-01 vol01-02 ENABLED 104139 0 - - -
pl vol01-03 vol01 ENABLED LOGONLY - ACTIVE - -
sd schost-103-02 vol01-03 ENABLED 5 LOG - - -
[Synchronize the disk group with cluster framework:]
cldevicegroup sync schost-1
[Check the file systems:]
fsck -y /dev/vx/rdisk/schost-1/bkup-vol
[Copy bkup-vol to the backup device:]
ufsdump 0ucf /dev/rmt/0 /dev/vx/rdisk/schost-1/bkup-vol
DUMP: Writing 63 Kilobyte records
DUMP: Date of this level 0 dump: Tue Apr 25 16:15:51 2000
DUMP: Date of last level 0 dump: the epoch
DUMP: Dumping /dev/vx/dsk/schost-2/bkup-vol to /dev/rmt/0.
...
DUMP: DUMP IS DONE
[Remove the bkup-volume:]
vxedit -rf rm bkup-vol
[Synchronize the disk group:]
cldevicegroup sync schost-1

```



## ▼ Copias de seguridad de la configuración del clúster

Para asegurarse de que se archive la configuración del clúster y para facilitar su recuperación, realice periódicamente copias de seguridad de la configuración del clúster. Oracle Solaris Cluster proporciona la capacidad de exportar la configuración del clúster a un archivo eXtensible Markup Language (XML).

- 1 **Inicie sesión en cualquier nodo del clúster y conviértase en superusuario o asuma una función que le proporcione la autorización de RBAC `solaris.cluster.read`.**

- 2 **Exporte la información de configuración del clúster a un archivo.**

```
/usr/cluster/bin/cluster export -o configfile
```

*archivo\_conf* El nombre del archivo de configuración XML al que el comando del clúster va a exportar la información de configuración del clúster. Para obtener información sobre el archivo de configuración XML, consulte [clconfiguration\(5CL\)](#)

- 3 **Compruebe que la información de configuración del clúster se haya exportado correctamente al archivo XML.**

```
vi configfile
```

## Restauración de los archivos del clúster

El comando `ufsrestore(1M)` copia en disco los archivos relativos al directorio actual en funcionamiento, desde copias de seguridad creadas con el comando `ufsdump(1M)`. Puede usar `ufsrestore` para recargar una jerarquía completa del archivo de sistemas desde el volcado de nivel 0 y los volcados que le siguen, o para restaurar al menos un archivo individual desde cualquier cinta de volcado. Si `ufsrestore` se ejecuta como superusuario o asume una función equivalente, los archivos se restauran con el propietario original, la última hora de modificación y el modo (permisos).

Antes de empezar a restaurar los archivos o los sistemas de archivos, debe conocer la información siguiente:

- Las cintas que necesita
- El nombre del dispositivo básico en el que va a restaurar el sistema de archivos
- El tipo de unidad de cinta que va a utilizar
- El nombre del dispositivo (local o remoto) para la unidad de cinta
- El esquema de partición de cualquier disco donde haya habido un error, porque las particiones y los sistemas de archivos deben estar duplicados exactamente en el disco de reemplazo

TABLA 12-2 Mapa de tareas: restaurar los archivos del clúster

Tarea	Instrucciones
Para Solaris Volume Manager, restaurar los archivos de forma interactiva	“Restauración de archivos individuales de forma interactiva (Solaris Volume Manager)” en la página 346
Para Solaris Volume Manager, restaurar el sistema de archivos root (/)	“Restauración del sistema de archivos root (/) (Solaris Volume Manager)” en la página 346  “Cómo restaurar un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager” en la página 349
Para Veritas Volume Manager, restaurar un sistema de archivos root (/)	“Restauración de un sistema de archivos root no encapsulado (/) (Veritas Volume Manager)” en la página 354
Para Veritas Volume Manager, restaurar un sistema de archivos root encapsulado (/)	“Restauración de un sistema de archivos root encapsulado (/) (Veritas Volume Manager)” en la página 356

▼ Restauración de archivos individuales de forma interactiva (Solaris Volume Manager)

Siga este procedimiento para restaurar al menos un archivo individual. Asegúrese de que el clúster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.admin` en el nodo del clúster que vaya a restaurar.
- 2 Detenga todos los servicios de datos que usen los archivos que va a restaurar.  
`# clresourcegroup offline resource-group`
- 3 Restaura los archivos.  
`# ufsrestore`

▼ Restauración del sistema de archivos root (/) (Solaris Volume Manager)

Siga este procedimiento para restaurar los sistemas de archivos root (/) en un disco nuevo, como después de reemplazar un disco root incorrecto. El nodo que se está restaurando no se debe arrancar. Asegúrese de que el clúster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

---

**Nota** – Como debe particionar el disco nuevo con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

`phys - schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en un nodo del clúster con acceso a los grupos de discos a los que también está vinculado el nodo que se va a restaurar.**

Use otro nodo *que no sea* el que va a restaurar.

- 2 Elimine el nombre de host del nodo que se va a restaurar desde todos los metaconjuntos.**

Ejecute este comando desde un nodo del metaconjunto que no sea el que va a eliminar. Debido a que el nodo de recuperación está fuera de línea, el sistema mostrará un error de RPC: `Rpcbind failure - RPC: Timed out`. Ignore este error y continúe con el siguiente paso.

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

-s                      Especifica el nombre del conjunto de discos.  
nombre\_conjunto

-f                      Elimina el último host del conjunto de discos.

-d                      Elimina del conjunto de discos.

-h lista\_nodos        Especifica el nombre del nodo que se va a eliminar del conjunto de discos.

- 3 Restaure los sistemas de archivos de raíz (/) y /usr.**

Para restaurar los sistemas de archivos de raíz y de /usr, siga el procedimiento del [Capítulo 26, “Restoring UFS Files and File Systems \(Tasks\)”](#) de *System Administration Guide: Devices and File Systems*. Omita el paso del procedimiento del sistema operativo Oracle Solaris para reiniciar el sistema.

---

**Nota** – Asegúrese de crear el sistema de archivos `/global/.devices/node@nodeid`.

---

- 4 Rearranque el nodo en modo multiusuario.**

```
reboot
```

**5 Sustituya el ID del dispositivo.**

```
cldevice repair rootdisk
```

**6 Use el comando `metadb(1M)` para recrear las réplicas de la base de datos de estado.**

```
metadb -c copies -af raw-disk-device
```

-c *copias* Especifica el número de réplicas que se van a crear.

-f *dispositivo\_disco\_básico* Dispositivo de disco básico en el que crear las réplicas.

-a Agrega réplicas.

**7 Desde un nodo del clúster que no sea el restaurado, agregue el nodo restaurado a todos los grupos de discos.**

```
phys-schost-2# metaset -s setname -a -h nodelist
```

-a Crea el host y lo agrega al conjunto de discos.

El nodo se reanuda en modo de clúster. El clúster está preparado para utilizarse.

## **Ejemplo 12-6 Restauración del sistema de archivos root (/) (Solaris Volume Manager)**

El ejemplo siguiente muestra el sistema de archivos root (/) restaurado en el nodo `phys-schost-1` desde el dispositivo de cinta `/dev/rmt/0`. El comando `metaset` se ejecuta desde otro nodo en el clúster, `phys-schost-2`, para eliminar y luego volver a agregar el nodo `phys-schost-1` al conjunto de discos `schost-1`. Todos los demás comandos se ejecutan desde `phys-schost-1`. Se crea un bloque de arranque nuevo en `/dev/rdsk/c0t0d0s0` y se vuelven a crear tres réplicas de la base de datos en `/dev/rdsk/c0t0d0s4`.

```
[Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on a cluster node
other than the node to be restored.]
[Remove the node from the metaset:]
phys-schost-2# metaset -s schost-1 -f -d -h phys-schost-1
[Replace the failed disk and boot the node:]
Restore the root (/) and /usr file system using the procedure in the Solaris system
administration documentation
[Reboot:]
reboot
[Replace the disk ID:]
cldevice repair /dev/dsk/c0t0d0
[Re-create state database replicas:]
metadb -c 3 -af /dev/rdsk/c0t0d0s4
[Add the node back to the metaset:]
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

## ▼ Cómo restaurar un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager

Siga este procedimiento para restaurar un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager cuando se llevaron a cabo las copias de seguridad. Realice este procedimiento en circunstancias como, por ejemplo, cuando un disco root está corrupto y se reemplaza por otro nuevo. El nodo que se está restaurando no se debe arrancar. Asegúrese de que el clúster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

---

**Nota** – Como debe particionar el disco nuevo con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

phys - schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

- 1 **Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify` en un nodo del clúster con acceso al conjunto de discos *que no sea* el nodo que va a restaurar.**

Use otro nodo *que no sea* el que va a restaurar.

- 2 **Elimine el nombre de host del nodo que se va a restaurar de todos los conjuntos de discos a los que esté conectado. Ejecute el comando siguiente una vez para cada conjunto de discos.**

```
metaset -s setname -d -h hostname
```

-s <i>nombre_conjunto</i>	Especifica el nombre del metaconjunto.
-f	Elimina el último host del conjunto de discos.
-d	Elimina del metaconjunto.
-h <i>lista_nodos</i>	Especifica el nombre del nodo que desea eliminar del metaconjunto.
-h <i>nombre_host</i>	Especifica el nombre del host.
-m <i>lista_host_mediador</i>	Especifica el nombre del host mediador para agregar o quitar desde el conjunto de discos.

- 3 Si el nodo es un host mediador de dos cadenas, elimine el mediador. Ejecute el comando siguiente una vez para cada conjunto de discos que esté conectado al nodo.

```
metaset -ssetname-d -m hostname
```

- 4 Sustituya el disco defectuoso del nodo en el que se restaurará el sistema de archivos root (/). Consulte los procesos de sustitución en la documentación que se suministra con el servidor.

- 5 Arranque el nodo que esté restaurando. El nodo reparado se inicia en modo de un solo usuario desde el CD-ROM, por lo que Solaris Volume Manager no se está ejecutando en el nodo.

- Si está usando el CD del sistema operativo Oracle Solaris, tenga en cuenta lo siguiente:

- SPARC: tipo:

```
ok boot cdrom -s
```

- x86: inserte el CD y arranque el sistema. Para ello, cierre el sistema y, a continuación, apáguelo y vuélvalo a encender. En la pantalla de los parámetros actuales de arranque, escriba b o i.

```
<<< 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 -s
```

- Si está usando un servidor JumpStart de Solaris, tenga en cuenta lo siguiente:

- SPARC: tipo:

```
ok boot net -s
```

- x86: inserte el CD y arranque el sistema. Para ello, cierre el sistema y, a continuación, apáguelo y vuélvalo a encender. En la pantalla de los parámetros actuales de arranque, escriba b o i.

```
<<< 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 -s
```

- 6 Cree todas las particiones e intercambie el espacio del disco root mediante el comando `format`.**  
Vuelva a crear el esquema de partición original que estaba en el disco defectuoso.

- 7 Cree el sistema de archivos root (/) y los pertinentes sistemas de archivos con el comando `newfs`.**

Vuelva a crear los sistemas de archivos originales que estaban en el disco defectuoso.

---

**Nota** – Compruebe que se cree el sistema de archivos `/global/.devices/node@nodeid`.

---

- 8 Monte el sistema de archivos root (/) en un punto de montaje temporal.**

```
mount device temp-mountpoint
```

- 9 Use los comandos siguientes para restaurar el sistema de archivos root (/).**

```
cd temp-mountpoint
ufsrestore rvf dump-device
rm restoresymtable
```

- 10 Instale un bloque de arranque nuevo en el disco nuevo.**

```
/usr/sbin/installboot /usr/platform/'uname -i'/lib/fs/ufs/bootblk
raw-disk-device
```

- 11 Elimine las líneas del archivo `/puntomontaje_temp/etc/system` para la información root de MDD.**

```
* Begin MDD root info (do not edit)
forceload: misc/md_trans
forceload: misc/md_raid
forceload: misc/md_mirror
forceload: misc/md_hotspares
forceload: misc/md_stripe
forceload: drv/pcipsy
forceload: drv/glm
forceload: drv/sd
rootdev:/pseudo/md@0:0,10,blk
* End MDD root info (do not edit)
```

- 12 Edite el archivo `/temp-mountpoint/etc/vfstab` para cambiar la entrada raíz de un volumen de Solaris Volume Manager a un segmento normal correspondiente para cada sistema de archivos del disco raíz perteneciente al metadispositivo o al volumen.**

Example:

Change from–

```
/dev/md/dsk/d10 /dev/md/rdisk/d10 / ufs 1 no -
```

Change to–

```
/dev/dsk/c0t0d0s0 /dev/rdisk/c0t0d0s0 / ufs 1 no -
```

**13 Desmonte el sistema de archivos temporales y compruebe el dispositivo de discos básico.**

```
cd /
umount temp-mountpoint
fsck raw-disk-device
```

**14 Rearranque el nodo en modo multiusuario.**

```
reboot
```

**15 Sustituya el ID del dispositivo.**

```
cldevice repair rootdisk
```

**16 Use el comando metadb para recrear las réplicas de la base de datos de estado.**

```
metadb -c copies -af raw-disk-device
```

-c *copias* Especifica el número de réplicas que se van a crear.

-af *dispositivo\_disco\_básico* Crea réplicas de base de datos de estado inicial en el denominado dispositivo de disco básico.

**17 Desde un nodo del clúster que no sea el restaurado, agregue el nodo restaurado a todos los grupos de discos.**

```
phys-schost-2# metaset -s setname -a -h nodelist
```

-a Agrega (crea) el metaconjunto.

Configure el volumen o la duplicación para la raíz (/) conforme a la documentación.

El nodo se rearranca en modo de clúster.

**18 Si el nodo era un host mediador de dos cadenas, vuelva a agregar el mediador.**

```
phys-schost-2# metaset -s setname -a -m hostname
```

**Ejemplo 12-7 Restauración de un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager**

El ejemplo siguiente muestra el sistema de archivos root (/) restaurado en el nodo phys-schost-1 desde el dispositivo de cinta /dev/rmt/0. El comando metaset se ejecuta desde otro nodo del clúster, phys-schost-2, para eliminar y luego volver a agregar el nodo phys-schost-1 al metaconjunto schost-1. Todos los demás comandos se ejecutan desde phys-schost-1. Se crea un bloque de arranque nuevo en /dev/rdisk/c0t0d0s0 y se vuelven a crear tres réplicas de la base de datos en /dev/rdisk/c0t0d0s4.

```
[Become superuser or assume a role that provides solaris.cluster.modify RBAC
 authorization on a cluster node with access to the metaset, other than the node to be restored.]
[Remove the node from the metaset:]
phys-schost-2# metaset -s schost-1 -d -h phys-schost-1
[Replace the failed disk and boot the node:]
```



Inicie el nodo desde el CD del sistema operativo Oracle Solaris:

- SPARC: tipo:

```
ok boot cdrom -s
```

- x86: inserte el CD y arranque el sistema. Para ello, cierre el sistema y, a continuación, apáguelo y vuélvalo a encender. En la pantalla de los parámetros actuales de arranque, escriba `b` o `i`.

```
<<< 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 -s
```

```
[Use format and newfs to recreate partitions and file systems
.]
[Mount the root file system on a temporary mount point:]
mount /dev/dsk/c0t0d0s0 /a
[Restore the root file system:]
cd /a
ufsrestore rvf /dev/rmt/0
rm restoresymtable
[Install a new boot block:]
/usr/sbin/installboot /usr/platform/'uname \
-i' /lib/fs/ufs/bootblk /dev/rdisk/c0t0d0s0

[Remove the lines in / temp-mountpoint/etc/system file for MDD root information:
]
* Begin MDD root info (do not edit)
forceload: misc/md_trans
forceload: misc/md_raid
forceload: misc/md_mirror
forceload: misc/md_hotspares
forceload: misc/md_stripe
forceload: drv/pcipsy
forceload: drv/glm
forceload: drv/sd
rootdev:/pseudo/md@0:0,10,blk
* End MDD root info (do not edit)
[Edit the /temp-mountpoint/etc/vfstab file]
Example:
Change from-
/dev/md/dsk/d10 /dev/md/rdsk/d10 / ufs 1 no -

Change to-
/dev/dsk/c0t0d0s0 /dev/rdsk/c0t0d0s0 /usr ufs 1 no -
[Unmount the temporary file system and check the raw disk device:]
cd /
umount /a
```

```
fsck /dev/rdisk/c0t0d0s0
[Reboot:]
reboot
[Replace the disk ID:]
cldevice repair /dev/rdisk/c0t0d0
[Re-create state database replicas:]
metadb -c 3 -af /dev/rdisk/c0t0d0s4
[Add the node back to the metaset:]
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

## ▼ Restauración de un sistema de archivos root no encapsulado (/) (Veritas Volume Manager)

Siga este procedimiento para restaurar un sistema de archivos root (/) no encapsulado en un nodo. El nodo que se está restaurando no se debe arrancar. Antes de efectuar el procedimiento de restauración, compruebe que el clúster se ejecute sin errores.

---

**Nota** – Como el disco nuevo debe particionarse con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

`phys-schost#` refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### 1 Sustituya el disco defectuoso en el nodo donde se restaurará el sistema de archivos root.

Consulte los procesos de sustitución en la documentación que se suministra con el servidor.

### 2 Arranque el nodo que esté restaurando.

- Si está usando el CD del sistema operativo Oracle Solaris, escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot cdrom -s
```

- Si está usando un servidor JumpStart de Solaris, escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot net -s
```

### 3 Cree todas las particiones e intercambie el disco root mediante el comando **format**.

Vuelva a crear el esquema de partición original que estaba en el disco defectuoso.

- 4 Cree el sistema de archivos root (/) y los pertinentes sistemas de archivos con el comando **newfs**.

Vuelva a crear los sistemas de archivos originales que estaban en el disco defectuoso.

---

**Nota** – Compruebe que se cree el sistema de archivos /global/.devices/node@nodeid.

---

- 5 Monte el sistema de archivos root (/) en un punto de montaje temporal.

```
mount device temp-mountpoint
```

- 6 Restaure el sistema de archivos root (/) desde la copia de seguridad; desmonte y compruebe el sistema.

```
cd temp-mountpoint
ufsrestore rvf dump-device
rm restoresymtable
cd /
umount temp-mountpoint
fsck raw-disk-device
```

El sistema de archivos ya está restaurado.

- 7 Instale un bloque de arranque nuevo en el disco nuevo.

```
/usr/sbin/installboot /usr/platform/'uname -i'/lib/fs/ufs/bootblk raw-disk-device
```

- 8 Rearranque el nodo en modo multiusuario.

```
reboot
```

- 9 Actualice el ID del dispositivo.

```
cldevice repair /dev/rdisk/disk-device
```

- 10 Pulse Control-D para reanudar en modo multiusuario.

El nodo se rearranca en modo de clúster. El clúster está preparado para utilizarse.

### Ejemplo 12–8 Restauración de un sistema de archivos root no encapsulado (/) (Veritas Volume Manager)

El ejemplo siguiente muestra un sistema de archivos root no encapsulado (/) restaurado en el nodo phys-schost-1 desde el dispositivo de cinta /dev/rmt/0.

[Replace the failed disk and boot the node:]

Inicie el nodo desde el CD del sistema operativo Oracle Solaris. Escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot cdrom -s
...
```

```
[Use format and newfs to create partitions and file systems]
[Mount the root file system on a temporary mount point:]
mount /dev/dsk/c0t0d0s0 /a
[Restore the root file system:]
cd /a
ufsrestore rvf /dev/rmt/0
rm restoresymtable
cd /
umount /a
fsck /dev/dsk/c0t0d0s0
[Install a new boot block:]
/usr/sbin/installboot /usr/platform/'uname \
-i'/lib/fs/ufs/bootblk /dev/dsk/c0t0d0s0

[Reboot:]
reboot
[Update the disk ID:]
cldevice repair /dev/dsk/c0t0d0
```

## ▼ Restauración de un sistema de archivos root encapsulado (/) (Veritas Volume Manager)

Siga este procedimiento para restaurar un sistema de archivos raíz encapsulado (/) en un nodo. El nodo que se está restaurando no se debe arrancar. Antes de efectuar el procedimiento de restauración, compruebe que el clúster se ejecute con errores.

---

**Nota** – Como el disco nuevo debe particionarse con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

phys -schost# refleja un indicador de clúster global. Siga este procedimiento en un clúster global.

Este procedimiento proporciona las formas largas de los comandos Oracle Solaris Cluster. La mayoría de los comandos también tienen una forma corta. A excepción de las formas de los nombres de comandos, los comandos son idénticos.

### 1 Sustituya el disco defectuoso en el nodo donde se restaurará el sistema de archivos root.

Consulte los procesos de sustitución en la documentación que se suministra con el servidor.

### 2 Arranque el nodo que esté restaurando.

- Si está usando el CD del sistema operativo Oracle Solaris, escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot cdrom -s
```

- Si está usando un servidor JumpStart de Solaris, escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot net -s
```

**3 Cree todas las particiones e intercambie el espacio del disco root mediante el comando `format`.**

Vuelva a crear el esquema de partición original que estaba en el disco defectuoso.

**4 Cree el sistema de archivos root (/) y los pertinentes sistemas de archivos con el comando `newfs`.**

Vuelva a crear los sistemas de archivos originales que estaban en el disco defectuoso.

---

**Nota** – Compruebe que se cree el sistema de archivos `/global/.devices/nodo@id_nodo`.

---

**5 Monte el sistema de archivos root (/) en un punto de montaje temporal.**

```
mount device temp-mountpoint
```

**6 Restaure el sistema de archivos root (/) desde la copia de seguridad.**

```
cd temp-mountpoint
ufsrestore rvf dump-device
rm restoresymtable
```

**7 Cree un archivo `install-db` vacío.**

Este archivo pone el nodo en modo de instalación VxVM en el próximo rearranque.

```
touch \
/temp-mountpoint/etc/vx/reconfig.d/state.d/install-db
```

**8 Elimine las entradas siguientes en el archivo `/punto_montaje_temp/etc/system`.**

```
* rootdev:/pseudo/vxio@0:0
* set vxio:vol_rootdev_is_volume=1
```

**9 Edite el dispositivo `/punto_montaje_temp/etc/vfstab` y sustituya todos los puntos de montaje de VxVM por los dispositivos de disco estándar para el disco root, como `/dev/dsk/c0t0d0s0`.**

Example:

Change from–

```
/dev/vx/dsk/rootdg/rootvol /dev/vx/rdisk/rootdg/rootvol / ufs 1 no -
```

Change to–

```
/dev/dsk/c0t0d0s0 /dev/rdsk/c0t0d0s0 / ufs 1 no -
```

**10 Desmonte el sistema de archivos temporales y compruebe el sistema de archivos.**

```
cd /
umount temp-mountpoint
fsck raw-disk-device
```

- 11 **Instale el bloque de arranque en el disco nuevo.**  
`# /usr/sbin/installboot /usr/platform/`uname -i`/lib/fs/ufs/bootblk raw-disk-device`
- 12 **Rearranque el nodo en modo multiusuario.**  
`# reboot`
- 13 **Actualice el ID del dispositivo con `scdidadm(1M)`.**  
`# cldevice repair /dev/rdisk/c0t0d0`
- 14 **Ejecute el comando de encapsulado `clvxdm` para encapsular el disco y reentrarlo.**
- 15 **Si se da un conflicto con cualquier otro sistema debido al número menor, desmonte los dispositivos globales y vuelva a duplicar el grupo de discos.**
  - Desmonte el sistema de archivos de dispositivo globales en el nodo del clúster.  
`# umount /global/.devices/node@nodeid`
  - Vuelva a duplicar el grupo de discos rootdg en el nodo del clúster.  
`# vxdg remminor rootdg 100`
- 16 **Cierre el nodo y reentrarlo en modo de clúster.**  
`# shutdown -g0 -i6 -y`

### **Ejemplo 12-9** Restauración de un sistema de archivos root encapsulado (/) (Veritas Volume Manager)

El ejemplo siguiente muestra un sistema de archivos root (/) restaurado en el nodo phys-schost-1 desde el dispositivo de cinta /dev/rmt/0.

[Replace the failed disk and boot the node:]

Inicie el nodo desde el CD del sistema operativo Oracle Solaris. Escriba el comando siguiente en el indicador ok de la PROM OpenBoot:

```
ok boot cdrom -s
...
[Use format and newfs to create partitions and file systems]
[Mount the root file system on a temporary mount point:]
mount /dev/dsk/c0t0d0s0 /a
[Restore the root file system:]
cd /a
ufsrestore rvf /dev/rmt/0
rm restoresymtable
[Create an empty install-db file:]
touch /a/etc/vx/reconfig.d/state.d/install-db
[Edit /etc/system on the temporary file system and
remove or comment out the following entries:]
rootdev:/pseudo/vxio@0:0
```

```

set vxio:vol_rootdev_is_volume=1
[Edit /etc/vfstab on the temporary file system:]
Example:
Change from-
/dev/vx/dsk/rootdg/rootvol /dev/vx/rdisk/rootdg/rootvol / ufs 1 no-

Change to-
/dev/dsk/c0t0d0s0 /dev/rdsk/c0t0d0s0 / ufs 1 no -
[Unmount the temporary file system, then check the file system:]
cd /
umount /a
fsck /dev/rdsk/c0t0d0s0
[Install a new boot block:]
/usr/sbin/installboot /usr/platform/'uname \
-i'/lib/fs/ufs/bootblk /dev/rdsk/c0t0d0s0
[Reboot:]
reboot
[Update the disk ID:]
cldevice repair /dev/rdsk/c0t0d0
[Encapsulate the disk:]
vxinstall
Choose to encapsulate the root disk.
[If a conflict in minor number occurs, remminor the rootdg disk group:]
umount /global/.devices/node@nodeid
vxdg remminor rootdg 100
shutdown -g0 -i6 -y

```

**Véase también** Para obtener instrucciones sobre cómo duplicar el disco raíz encapsulado, consulte [Guía de instalación del software Oracle Solaris Cluster](#).





# Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario

---

Este capítulo proporciona las descripciones de las herramientas de interfaz gráfica de usuario de Oracle Solaris Cluster Manager y Sun Management Center, que podrá utilizar para administrar muchos aspectos de un clúster. También contiene procedimientos para configurar e iniciar Oracle Solaris Cluster Manager. La ayuda en línea incluida con la GUI de Oracle Solaris Cluster Manager contiene instrucciones para realizar distintas tareas administrativas de Oracle Solaris Cluster.

Este capítulo incluye lo siguiente:

- “Información general sobre Oracle Solaris Cluster Manager” en la página 361
- “SPARC: Información general sobre Sun Management Center” en la página 362
- “Configuración de Oracle Solaris Cluster Manager” en la página 363
- “Inicio del software Oracle Solaris Cluster Manager” en la página 366

## Información general sobre Oracle Solaris Cluster Manager

Oracle Solaris Cluster Manager es una GUI que permite visualizar información del clúster, comprobar el estado de los componentes del clúster y supervisar los cambios aplicados a la configuración. Oracle Solaris Cluster Manager también permite realizar numerosas tareas administrativas para los componentes siguientes de Oracle Solaris Cluster:

- Adaptadores
- Cables
- Servicios de datos
- Dispositivos globales
- Interconexiones
- Uniones
- Límites de carga de nodos
- Dispositivos NAS
- Nodos
- Dispositivos de quórum

- Grupos de recursos
- Recursos

En las ubicaciones siguientes encontrará información sobre la instalación y el uso de Oracle Solaris Cluster Manager:

- **Instalación de Oracle Solaris Cluster Manager:** consulte *Guía de instalación del software Oracle Solaris Cluster*.
- **Inicio de Oracle Solaris Cluster Manager:** consulte “Inicio del software Oracle Solaris Cluster Manager” en la página 366.
- **Configuración de números de puerto, direcciones de servidor, certificados de seguridad y usuarios:** consulte “Configuración de Oracle Solaris Cluster Manager” en la página 363.
- **Instalación y administración de aspectos del clúster mediante Oracle Solaris Cluster Manager:** consulte la ayuda en línea que se proporciona con Oracle Solaris Cluster Manager.
- **Regeneración de las claves de seguridad de Oracle Solaris Cluster Manager:** consulte “Regeneración de las claves de seguridad de Common Agent Container” en la página 365.

---

**Nota** – Sin embargo, Oracle Solaris Cluster Manager no puede realizar de momento todas las tareas administrativas de Oracle Solaris Cluster. En algunas operaciones debe utilizarse la interfaz de línea de comandos.

---

## SPARC: Información general sobre Sun Management Center

El módulo de Oracle Solaris Cluster para la consola GUI de Sun Management Center (anteriormente Sun Enterprise SyMON) permite visualizar los recursos, los tipos y los grupos de recursos del clúster. También permite supervisar los cambios en la configuración y comprobar el estado de los componentes del clúster. Sin embargo, el módulo de Oracle Solaris Cluster para Sun Management Center no puede realizar las tareas de configuración de Oracle Solaris Cluster. Para efectuar las operaciones de configuración debe usar la interfaz de línea de comandos. Consulte "Interfaz de línea de comandos" en el Capítulo 1 para obtener más información.

Para obtener más información sobre cómo instalar e iniciar el módulo de Oracle Solaris Cluster para Sun Management Center, consulte el [Capítulo 8, “Instalación del módulo de Oracle Solaris Cluster en Sun Management Center”](#) de *Guía de instalación del software Oracle Solaris Cluster*.

El módulo Oracle Solaris Cluster de Sun Management Center es compatible con el protocolo SNMP. Oracle Solaris Cluster ha creado una MIB válida como definición de datos por parte de las estaciones de administración de otros proveedores basadas en SNMP.

El archivo de MIB de Oracle Solaris Cluster se ubica en `/opt/SUNWsymon/modules/cfg/sun-cluster-mib.mib` en todos los nodos del clúster.

El archivo de MIB de Oracle Solaris Cluster es una especificación modelada de ASN.1 de los datos de Oracle Solaris Cluster. Es la misma especificación que utilizan todas las MIB de Sun Management Center. Para usar la MIB de Oracle Solaris Cluster, consulte las instrucciones para utilizar otras MIB de Sun Management Center en *"SNMP MIBs for Sun Management Centre Modules" en la Sun Management Center 3.6 User's Guide*.

## Configuración de Oracle Solaris Cluster Manager

Oracle Solaris Cluster Manager es una GUI válida para administrar y visualizar el estado de todos los aspectos de los dispositivos de quórum, los grupos de IPMP, los componentes de interconexión y los dispositivos globales. La GUI puede utilizarse en lugar de muchos comandos de la interfaz de línea de comandos de Oracle Solaris Cluster.

El procedimiento para instalar Oracle Solaris Cluster Manager en el clúster se explica en [Guía de instalación del software Oracle Solaris Cluster](#). La ayuda en línea de Oracle Solaris Cluster Manager proporciona instrucciones para realizar varias tareas con la GUI.

Esta sección contiene los procedimientos siguientes para volver a configurar Oracle Solaris Cluster Manager tras la instalación inicial.

- ["Configuración de funciones de RBAC" en la página 363](#)
- ["Cambio de la dirección del servidor para Oracle Solaris Cluster Manager" en la página 365](#)
- ["Regeneración de las claves de seguridad de Common Agent Container" en la página 365](#)

## Configuración de funciones de RBAC

Oracle Solaris Cluster Manager utiliza RBAC para determinar quién tiene derechos para administrar el clúster. Algunos perfiles de derechos de RBAC están incluidos en el software Oracle Solaris Cluster. Puede asignar estos perfiles de derechos a los usuarios o a las funciones para proporcionar a los usuarios niveles diferentes de acceso a Oracle Solaris Cluster. Para obtener más información sobre cómo configurar y administrar RBAC para el software Oracle Solaris Cluster, consulte el [Capítulo 2, "Oracle Solaris Cluster y RBAC"](#).

## ▼ Uso de Common Agent Container para cambiar los números de puerto de los agentes de servicios o de administración

Si los números de puerto predeterminados para los servicios de Common Agent Container presentan conflictos con otros procesos de ejecución, puede usar el comando `cacaoadm` en todos los nodos del clúster para cambiar el número de puerto del agente de servicio de administración causante del conflicto.

- 1 Detenga el daemon de administración Common Agent Container en todos los nodos del clúster.

```
/opt/bin/cacaoadm stop
```

- 2 Detenga Sun Java Web Console.

```
/usr/sbin/smcwebserver stop
```

- 3 Recupere el número de puerto que el servicio de Common Agent Container usa con el subcomando `get-param`.

```
/opt/bin/cacaoadm get-param parameterName
```

Puede usar el comando `cacaoadm` para cambiar los números de puerto de los siguientes servicios de Common Agent Container. La lista siguiente proporciona algunos ejemplos de servicios y agentes que Common Agent Container puede administrar, junto con los nombres de parámetros correspondientes.

Puerto del conector JMX	<code>jmxmp-connector-port</code>
Puerto SNMP	<code>snmp-adapter-port</code>
Puerto de captura de SNMP	<code>snmp-adapter-trap-port</code>
Puerto de secuencia de comandos	<code>commandstream-adapter-port</code>

- 4 Cambie un número de puerto.

```
/opt/bin/cacaoadm set-param parameterName=parameterValue
=parameterValue
```

- 5 Repita el [Paso 4](#) en todos los nodos del clúster.

- 6 Reinicie Sun Java Web Console.

```
/usr/sbin/smcwebserver start
```

- 7 Reinicie el daemon de administración de Common Agent Container en todos los nodos del clúster.

```
/opt/bin/cacaoadm start
```

## ▼ Cambio de la dirección del servidor para Oracle Solaris Cluster Manager

Si cambia el nombre de host de un nodo del clúster, debe cambiar la dirección desde la que se ejecuta Oracle Solaris Cluster Manager. El certificado de seguridad predeterminado se genera conforme al nombre de host del nodo al instalarse Oracle Solaris Cluster Manager. Para restablecer el nombre de host del nodo, elimine el archivo de certificado, `keystore`, y reinicie Oracle Solaris Cluster Manager. Oracle Solaris Cluster Manager crea de forma automática un archivo de certificado con el nombre de host nuevo. Aplique este procedimiento en todos los nodos cuyo nombre de host se haya cambiado.

- 1 Elimine el archivo de certificado, `keystore`, ubicado en `/etc/opt/webconsole`.

```
cd /etc/opt/webconsole
pkgrm keystore
```

- 2 Reinicie Oracle Solaris Cluster Manager.

```
/usr/sbin/smcwebserver restart
```

## ▼ Regeneración de las claves de seguridad de Common Agent Container

Oracle Solaris Cluster Manager usa potentes técnicas de cifrado para asegurar una comunicación segura entre el servidor web de Oracle Solaris Cluster Manager y todos los nodos del clúster.

Las claves utilizadas por Oracle Solaris Cluster Manager se almacenan en el directorio `/etc/opt/SUNWcacao/security` en todos los nodos. Deben ser idénticas en todos los nodos del clúster.

Durante el funcionamiento normal, estas claves se pueden quedar en la configuración predeterminada. Si cambia el nombre de host de un nodo del clúster, debe regenerar las claves de seguridad de Common Agent Container. También podría tener que regenerar las claves si hubiese un posible peligro para la clave, por ejemplo un peligro en la raíz o root del equipo. Para regenerar las claves de seguridad, siga este siguiente.

- 1 Detenga el daemon de administración Common Agent Container en todos los nodos del clúster.

```
/opt/bin/cacaoadm stop
```

- 2 Regenera las claves de seguridad en un nodo del clúster.

```
phys-schost-1# /opt/bin/cacaoadm create-keys --force
```

- 3 **Reinicie el daemon de administración de Common Agent Container en el nodo en que regeneró las claves de seguridad.**

```
phys-schost-1# /opt/bin/cacaoadm start
```

- 4 **Cree un archivo tar del directorio `/etc/cacao/instances/default`.**

```
phys-schost-1# cd /etc/cacao/instances/default
phys-schost-1# tar cf /tmp/SECURITY.tar security
```

- 5 **Copie el archivo `/tmp/Security.tar` en todos los nodos del clúster.**

- 6 **Extraiga los archivos de seguridad de todos los nodos en los que haya copiado el archivo `/tmp/SECURITY.tar`.**

Se sobrescriben todos los archivos de seguridad que ya existan en el directorio `/etc/opt/SUNWcacao/`.

```
phys-schost-2# cd /etc/cacao/instances/default
phys-schost-2# tar xf /tmp/SECURITY.tar
```

- 7 **Elimine el archivo `/tmp/SECURITY.tar` de todos los nodos del clúster.**

Es preciso eliminar todas las copias del archivo tar para evitar riesgos de seguridad.

```
phys-schost-1# rm /tmp/SECURITY.tar
```

```
phys-schost-2# rm /tmp/SECURITY.tar
```

- 8 **Reinicie el daemon de administración Common Agent Container en todos los nodos.**

```
phys-schost-1# /opt/bin/cacaoadm start
```

- 9 **Reinicie Oracle Solaris Cluster Manager.**

```
/usr/sbin/smcwebserver restart
```

## Inicio del software Oracle Solaris Cluster Manager

La GUI de Oracle Solaris Cluster Manager permite administrar de forma sencilla los diferentes aspectos del software Oracle Solaris Cluster. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Sun Java Web Console y Common Agent Container se inician de forma automática al arrancar el clúster. Si necesita comprobar que Sun Java Web Console y Common Agent Container estén en funcionamiento, consulte la sección Resolución de problemas inmediatamente después de este procedimiento.

## ▼ Inicio de Oracle Solaris Cluster Manager

Este procedimiento muestra cómo iniciar Oracle Solaris Cluster Manager en el clúster.

- 1 **Determine si tiene intención de acceder a Oracle Solaris Cluster Manager mediante el nombre y la contraseña del usuario root del nodo del clúster, o de configurar un nombre y una contraseña de usuario diferentes.**
  - Si va a acceder a Oracle Solaris Cluster Manager con el nombre de usuario root del nodo del clúster, vaya al [Paso 5](#).
  - Si pretende configurar un nombre y una contraseña de usuario diferentes, vaya al [Paso 3](#) para establecer las cuentas de usuario de Oracle Solaris Cluster Manager.

- 2 **Conviértase en superusuario en un nodo de clúster.**

- 3 **Cree una cuenta de usuario para acceder al clúster mediante Oracle Solaris Cluster Manager.**

Use el comando `useradd(1M)` para agregar una cuenta de usuario al sistema. Debe configurarse al menos una cuenta de usuario para acceder a Oracle Solaris Cluster Manager si no usa la cuenta del sistema root. Las cuentas de usuario de Oracle Solaris Cluster Manager sólo las utiliza Oracle Solaris Cluster Manager. Estas cuentas no se corresponden con ninguna cuenta de usuario del sistema del sistema operativo Oracle Solaris. En “[Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster](#)” en la [página 53](#) se describe con más detalle cómo crear y asignar una función de RBAC a una cuenta de usuario.

---

**Nota** – Los usuarios que no tengan establecida una cuenta de usuario en un nodo determinado no pueden acceder al clúster mediante Oracle Solaris Cluster Manager desde ese nodo; tampoco pueden administrar el nodo a través de otro nodo del clúster al que los usuarios tengan acceso.

---

- 4 **(Opcional) Repita el [Paso 3](#) para configurar más cuentas de usuario.**
- 5 **Inicie un navegador desde la consola de administración u otro equipo fuera del clúster.**
- 6 **Compruebe que los tamaños del disco del navegador y de la memoria caché estén configurados con un valor superior a 0.**
- 7 **Compruebe que Java y Javascript estén habilitados en el navegador.**
- 8 **Conéctese al puerto de Oracle Solaris Cluster Manager en un nodo del clúster desde el navegador.**

El número de puerto predeterminado es 6789.

`https://node:6789/`

**9 Acepte los certificados que presente el navegador web.**

Aparece la página de inicio de sesión de Java Web Console.

**10 Indique el nombre y la contraseña del usuario que quiera que tenga acceso a Oracle Solaris Cluster Manager.**

**11 Haga clic en el botón Log In (Iniciar sesión).**

Aparece la página de inicio de la aplicación de Java Web Console.

**12 Haga clic en el vínculo de Oracle Solaris Cluster Manager en la categoría Systems (Sistemas).**

**13 Acepte cualquier otro certificado que presente el navegador web.**

**14 Si no puede conectarse a Oracle Solaris Cluster Manager, efectúe los subpasos siguientes para determinar si se ha elegido un perfil de red restringido durante la instalación de Solaris y para restaurar el acceso externo al servicio de Java Web Console.**

Si elige un perfil de red restringido durante la instalación de Oracle Solaris, el acceso externo al servicio de Sun Java Web Console está restringido. Esta red es necesaria para usar la GUI de Oracle Solaris Cluster Manager.

**a. Determine si el servicio de Java Web Console está restringido.**

```
svcprop /system/webconsole:console | grep tcp_listen
```

Si el valor de la propiedad de tcp\_listen no es true, el servicio de la consola web está restringido.

**b. Restaure el acceso externo al servicio de Java Web Console.**

```
svccfg
svc:> select system/webconsole
svc:/system/webconsole> setprop options/tcp_listen=true
svc:/system/webconsole> quit
/usr/sbin/smcwebserver restart
```

**c. Compruebe que el servicio esté disponible.**

```
netstat -a | grep 6789
```

Si el servicio está disponible, la salida del comando devuelve una entrada de 6789, que es el número de puerto que se utiliza para conectarse con Java Web Console.

**Errores más frecuentes**

Si después de realizar este proceso no se puede conectar con Oracle Solaris Cluster Manager, determine si Sun Java Web Console está en funcionamiento mediante el comando `/usr/sbin/smcwebserver status`. Si Sun Java Web Console no está en funcionamiento, haga que se inicie de forma manual con el comando `/usr/sbin/smcwebserver start`. Si todavía no



puede conectarse con Oracle Solaris Cluster Manager, determine si Common Agent Container se está ejecutando mediante `usr/sbin/cacaoadm status`. Si Common Agent Container no se está ejecutando, inícielo de forma manual mediante `/usr/sbin/cacaoadm start`.



## Ejemplo

---

### Configuración de replicación de datos basada en host con el software Sun StorageTek Availability Suite

Este es una alternativa a la replicación basada en host que no utiliza Oracle Solaris Cluster Cluster Geographic Edition. Oracle recomienda el uso de Oracle Solaris Cluster Geographic Edition para la replicación basada en host, con el fin de simplificar la configuración y el funcionamiento de la replicación basada en host en un clúster. Consulte [“Comprensión de la replicación de datos” en la página 86](#).

El ejemplo de este muestra cómo configurar una replicación de datos basada en host entre clústeres con el software Sun StorageTek Availability Suite 4.0. El ejemplo muestra una configuración de clúster completa para una aplicación NFS que proporciona información detallada sobre cómo realizar tareas individuales. Todas las tareas deben llevarse a cabo en el nodo de votación de clúster global. El ejemplo no incluye todos los pasos necesarios para otras aplicaciones ni las configuraciones de otros clústers.

Si utiliza el control de acceso basado en funciones (RBAC) en lugar de ser superusuario para acceder a los nodos del clúster, asegúrese de poder asumir una función de RBAC que proporcione autorización para todos los comandos de Oracle Solaris Cluster. Esta serie de procedimientos de replicación de datos requiere las siguientes autorizaciones de RBAC de Oracle Solaris Cluster si el usuario no es un superusuario:

- `solaris.cluster.modify`
- `solaris.cluster.admin`
- `solaris.cluster.read`

Consulte [System Administration Guide: Security Services](#) para obtener más información sobre el uso de las funciones de RBAC. Consulte las páginas de comando `man` de Oracle Solaris Cluster para saber la autorización de RBAC que requiere cada subcomando de Oracle Solaris Cluster.

## Comprensión del software Sun StorageTek Availability Suite en un clúster

Esta sección presenta la tolerancia ante problemas graves y describe los métodos de replicación de datos que utiliza el software &AvailSuite.

La tolerancia ante problemas graves es la capacidad de un sistema para recuperar una aplicación en un clúster alternativo cuando falla el clúster primario. La tolerancia ante problemas graves se basa en la *replicación de datos* y la *migración tras error*. La migración tras error es la reubicación automática de un grupo de recursos o de dispositivos de un clúster primario en un clúster secundario. Si falla el clúster primario, la aplicación y los datos quedan inmediatamente disponibles en el clúster secundario.

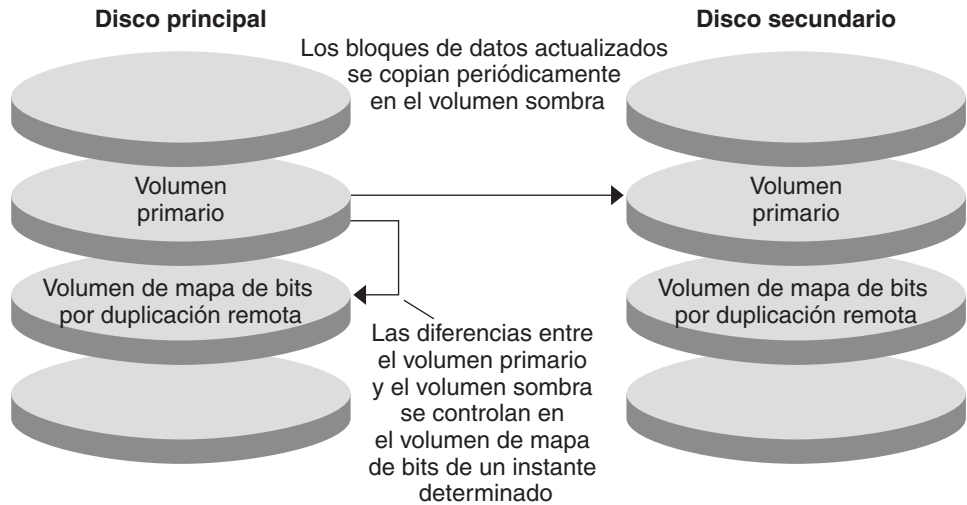
### Métodos de replicación de datos utilizados por Sun StorageTek Availability Suite

Esta sección describe los métodos de replicación por duplicación remota y de instantánea de un momento determinado utilizados por Sun StorageTek Availability Suite. Este software utiliza los comandos `sndradm(1RPC)` e `iiaadm(1II)` para replicar datos. Si desea más información sobre estos comandos, consulte uno de los manuales siguientes:

#### Replicación por duplicación remota

La replicación por duplicación remota se ilustra en la [Figura A–1](#). Los datos del volumen maestro del disco primario se duplican en el volumen maestro del disco secundario mediante una conexión TCP/IP. Un mapa de bits de duplicación remota controla las diferencias entre el volumen maestro del disco primario y el del secundario.

FIGURA A-1 Replicación por duplicación remota



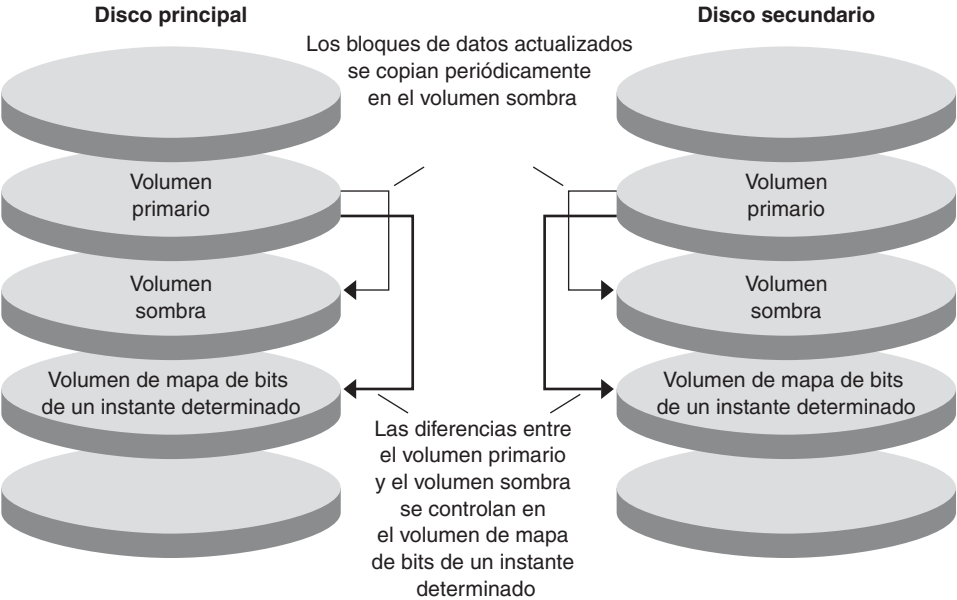
La replicación por duplicación remota se puede efectuar de manera sincrónica en tiempo real o de manera asincrónica. Cada volumen definido en cada clúster se puede configurar de manera individual, para replicación sincrónica o asincrónica.

- En la replicación sincrónica de datos, no se confirma la finalización de la operación de escritura hasta que el volumen remoto se haya actualizado.
- En la replicación asincrónica de datos, se confirma la finalización de la operación de escritura antes de que se actualice el volumen remoto. La replicación asincrónica de datos proporciona una mayor flexibilidad en largas distancias y poco ancho de banda.

## Instantánea de un momento determinado

La [Figura A-2](#) muestra instantánea de un momento determinado. Los datos del volumen primario de cada disco se copian en el volumen sombra del mismo disco. El mapa de bits instantáneo controla y detecta las diferencias entre el volumen primario y el volumen sombra. Cuando los datos se copian en el volumen sombra, el mapa de bits de un momento determinado se vuelve a configurar.

FIGURA A-2 Instantánea de un momento determinado



### Replicación en la configuración de ejemplo

La [Figura A-3](#) ilustra la forma en que se usan la replicación por duplicación remota y la instantánea de un momento determinado en este ejemplo de configuración.

FIGURA A-3 Replicación en la configuración de ejemplo



## Directrices para la configuración de replicación de datos basada en host entre clústers

Esta sección proporciona directrices para la configuración de replicación de datos entre clústers. Asimismo, se proporcionan consejos para configurar los grupos de recursos de replications y los de recursos de aplicaciones. Use estas directrices cuando esté configurando la replicación de datos en el clúster.

En esta sección se analizan los aspectos siguientes:

- “Configuración de grupos de recursos de replications” en la página 376
- “Configuración de grupos de recursos de aplicaciones” en la página 376
  - “Configuración de grupos de recursos en una aplicación de migración tras error” en la página 377
  - “Configuración de grupos de recursos en una aplicación escalable” en la página 378
- “Directrices para administrar migraciones tras error” en la página 379

## Configuración de grupos de recursos de replicaciones

Los grupos de recursos de replicaciones sitúan al grupo de dispositivos bajo el control del software &AvailSuite con el recurso de nombre de host lógico. Un grupo de recursos de replicaciones debe tener las características siguientes:

- Ser un grupo de recursos de migración tras error.

Un recurso de migración tras error sólo puede ejecutarse en un nodo a la vez. Cuando se produce la migración tras error, los recursos de recuperación participan en ella.

- Debe haber un recurso de nombres de host lógico.

El clúster primario debe alojar el nombre de host lógico. Después de una migración tras error, el clúster secundario debe alojar el nombre de host lógico. El sistema DNS se utiliza para asociar el nombre de host lógico a un clúster.

- Debe haber un recurso de HASToragePlus.

El recurso de HASToragePlus fuerza la migración tras error del grupo de dispositivos si el grupo de recursos de replicaciones se conmuta o migra tras error. Oracle Solaris Cluster también fuerza la migración tras error del grupo de recursos de replicaciones si el grupo de dispositivos se conmuta. De este modo, es siempre el mismo nodo el que sitúa o controla el grupo de recursos de replicaciones y el de dispositivos.

Las propiedades de extensión siguientes se deben definir en el recurso de &HASToragePlus:

- *GlobalDevicePaths*. La propiedad de esta extensión define el grupo de dispositivos al que pertenece un volumen.
- *AffinityOn property* = True. La propiedad de esta extensión provoca que el grupo de dispositivos se conmute o migre tras error si el grupo de recursos de replicaciones también lo hace. Esta función se denomina *conmutación de afinidad*.
- *ZPoolsSearchDir*. La propiedad de esta extensión es necesaria para utilizar el sistema de archivos ZFS.

Para obtener más información sobre HASToragePlus, consulte la página de comando [man SUNW.HASToragePlus\(5\)](#).

- Recibir el nombre del grupo de dispositivos con el que se coloca, seguido de -stor-rg  
Por ejemplo, devgrp-stor-rg.
- Estar en línea en el clúster primario y en el secundario

## Configuración de grupos de recursos de aplicaciones

Para que esté totalmente disponible, una aplicación se debe administrar como un recurso en un grupo de recursos de aplicaciones. Un grupo de recursos de aplicaciones se puede configurar en una aplicación de migración tras error o en una aplicación escalable.

Los recursos de aplicaciones y los grupos de recursos de aplicaciones configurados en el clúster primario también se deben configurar en el clúster secundario. Asimismo, se deben replicar en el clúster secundario los datos a los que tiene acceso el recurso de aplicaciones.



Esta sección proporciona directrices para la configuración de grupos de recursos de aplicaciones siguientes:

- “Configuración de grupos de recursos en una aplicación de migración tras error” en la página 377
- “Configuración de grupos de recursos en una aplicación escalable” en la página 378

## Configuración de grupos de recursos en una aplicación de migración tras error

En una aplicación de migración tras error, una aplicación se ejecuta en un solo nodo a la vez. Si éste falla, la aplicación migra a otro nodo del mismo clúster. Un grupo de recursos de una aplicación de migración tras error debe tener las características siguientes:

- Tener un recurso de HASToragePlus para forzar la migración tras error del grupo de dispositivos cuando el grupo de recursos de aplicaciones se conmute o migre tras error.

El grupo de dispositivos se coloca con el grupo de recursos de replications y el grupo de recursos de aplicaciones. Por este motivo, la migración tras error del grupo de recursos de aplicaciones fuerza la migración de los grupos de dispositivos y de recursos de replications. El mismo nodo controla los grupos de recursos de aplicaciones y de recursos de replications, y el grupo de dispositivos.

Sin embargo, que una migración tras error del grupo de dispositivos o del grupo de recursos de replications no desencadena una migración tras error en el grupo de recursos de aplicaciones.

- Si los datos de la aplicación están montados de manera global, la presencia de un recurso de HASToragePlus en el grupo de recursos de aplicaciones no es obligatoria, aunque sí aconsejable.

- Si los datos de la aplicación se montan de manera local, la presencia de un recurso de HASToragePlus en el grupo de recursos de aplicaciones es obligatoria.

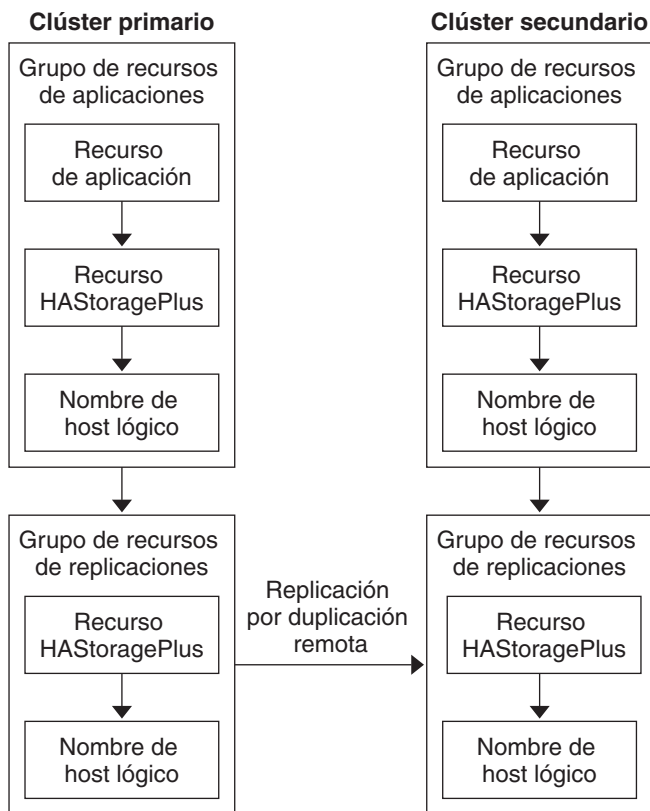
Sin un recurso de HASToragePlus, la migración tras error del grupo de recursos de aplicaciones no activaría la migración del grupo de recursos de replications y del grupo de dispositivos. Después de una migración tras error, el mismo nodo no controlaría el grupo de recursos de aplicaciones, el grupo de recursos de replications y el grupo de dispositivos.

Para obtener más información sobre HASToragePlus, consulte la página de comando `man SUNW.HASToragePlus(5)`.

- Debe estar en línea en el clúster primario y fuera de línea en el secundario.

El grupo de recursos de aplicaciones debe estar en línea en el clúster secundario cuando éste hace las funciones de clúster primario.

La [Figura A-4](#) ilustra la configuración de grupos de recursos de aplicaciones y de recursos de replications en una aplicación de migración tras error.

**FIGURA A-4** Configuración de grupos de recursos en una aplicación de migración tras error

## Configuración de grupos de recursos en una aplicación escalable

En una aplicación escalable, una aplicación se ejecuta en varios nodos con el fin de crear un único servicio lógico. Si se produce un error en un nodo que ejecuta una aplicación escalable, no habrá migración tras error. La aplicación continúa ejecutándose en los otros nodos.

Cuando una aplicación escalable se administra como recurso en un grupo de recursos de aplicaciones, no es necesario acoplar el grupo de recursos de aplicaciones con el grupo de dispositivos. Por este motivo, no es necesario crear un recurso de HAStoragePlus para el grupo de recursos de aplicaciones.

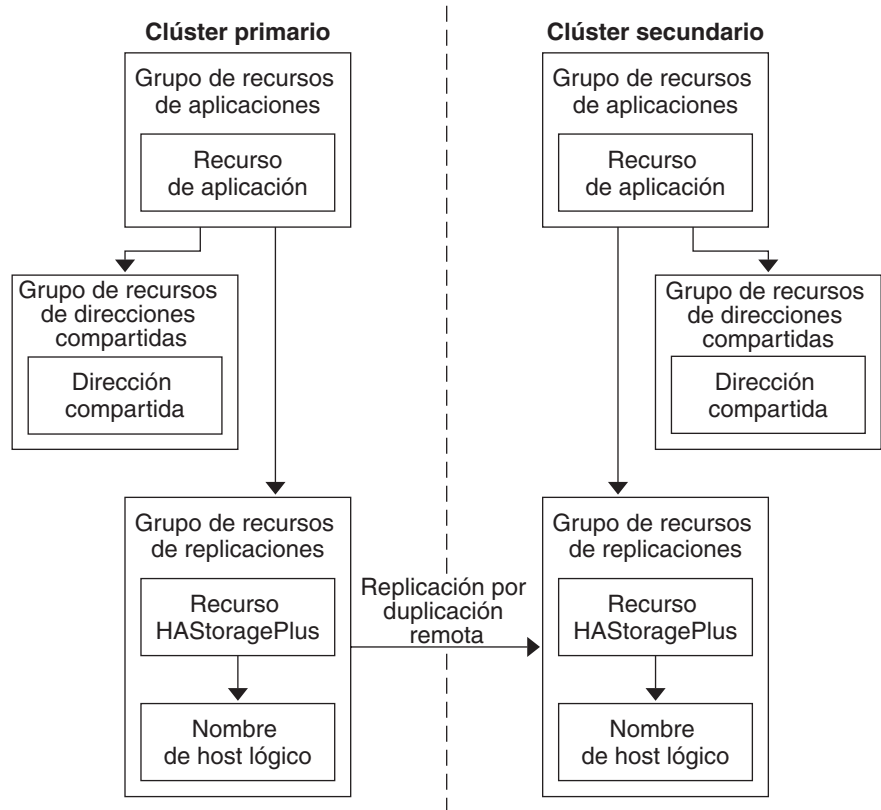
Un grupo de recursos de una aplicación escalable debe tener las características siguientes:

- Tener una dependencia en el grupo de recursos de direcciones compartidas.  
Los nodos que ejecutan la aplicación escalable utilizan la dirección compartida para distribuir datos entrantes.

- Estar en línea en el clúster primario y fuera de línea en el secundario.

La [Figura A-5](#) ilustra la configuración de grupos de recursos en una aplicación escalable.

**FIGURA A-5** Configuración de grupos de recursos en una aplicación escalable

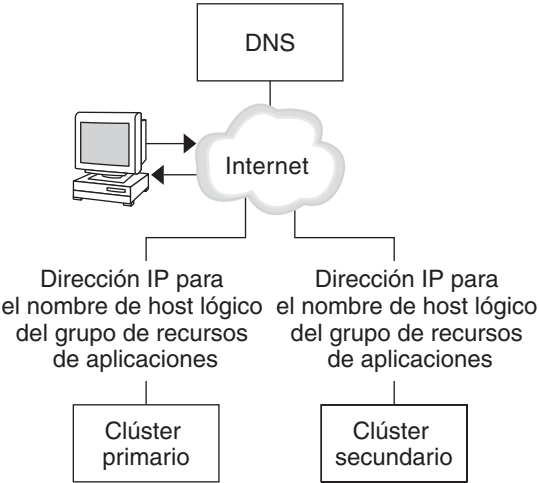


## Directrices para administrar migraciones tras error

Si se produce un error en el clúster primario, la aplicación se debe conmutar con el clúster secundario en cuanto sea posible. Para que el clúster secundario pueda realizar las funciones del principal, se debe actualizar el DNS.

El DNS asocia un cliente con el nombre de host lógico de una aplicación. Después de una migración tras error, se debe suprimir la asignación del DNS al clúster primario y se debe crear una asignación del DNS al clúster secundario. La [Figura A-6](#) muestra cómo se asigna el DNS a un cliente en un clúster.

FIGURA A-6 Asignación del DNS de un cliente a un clúster



Si desea actualizar el DNS, utilice el comando `nsupdate`. Para obtener más información, consulte la página de comando `man nsupdate(1M)`. Si desea ver un ejemplo de ocuparse de una migración tras error, consulte [“Ejemplo de cómo controlar una migración tras error” en la página 406](#).

Después de la reparación, el clúster primario puede volver a estar en línea. Para volver al clúster primario original, siga estos pasos:

1. Sincronice el clúster primario con el secundario para asegurarse de que el volumen principal esté actualizado.
2. Actualice el DNS de modo que los clientes tengan acceso a la aplicación en el clúster primario.

## Mapa de tareas: ejemplo de configuración de una replicación de datos

La [Tabla A-1](#) enumera las tareas de este ejemplo de configuración de replicación de datos para una aplicación de NFS mediante el software Sun StorageTek Availability Suite.

TABLA A-1 Mapa de tareas: ejemplo de configuración de una replicación de datos

Tarea	Instrucciones
1. Conectar e instalar los clústers	<a href="#">“Conexión e instalación de clústers” en la página 381</a>

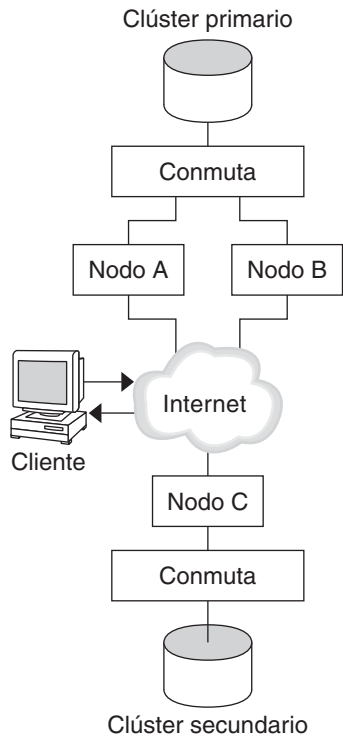
**TABLA A-1** Mapa de tareas: ejemplo de configuración de una replicación de datos (Continuación)

Tarea	Instrucciones
2. Configurar los grupos de dispositivos, los sistemas de archivos para la aplicación de NFS y los grupos de recursos de los clústers principal y secundario	“Ejemplo de configuración de grupos de dispositivos y de recursos” en la página 383
3. Habilitar la replicación de datos en el clúster primario y en el secundario	“Habilitación de la replicación en el clúster primario” en la página 398 “Habilitación de la replicación en el clúster secundario” en la página 400
4. Efectuar la replicación de datos	“Replicación por duplicación remota” en la página 401 “Instantánea de un momento determinado” en la página 402
5. Comprobar la configuración de la replicación de datos	“Procedimiento para comprobar la configuración de la replicación” en la página 403

## Conexión e instalación de clústers

La [Figura A-7](#) ilustra la configuración de clústers utilizada en la configuración de ejemplo. El clúster secundario de la configuración de ejemplo contiene un nodo, pero se pueden usar otras configuraciones de clúster.

FIGURA A-7 Ejemplo de configuración de clústers



La [Tabla A-2](#) resume el hardware y el software necesarios para la configuración de ejemplo. El sistema operativo Oracle Solaris, el software Oracle Solaris Cluster y el administrador de volúmenes se deben instalar en los nodos del clúster *antes* de instalar los parches y el software Sun StorageTek Availability Suite.

TABLA A-2 Hardware y software necesarios

Hardware o software	Requisito
Hardware de nodos	El software Sun StorageTek Availability Suite es compatible con todos los servidores que utilicen el sistema operativo Oracle Solaris.  Si desea información sobre el hardware que se debe usar, consulte <a href="#">Oracle Solaris Cluster 3.3 Hardware Administration Manual</a> .
Espacio en el disco	Aproximadamente 15 MB.

TABLA A-2 Hardware y software necesarios (Continuación)

Hardware o software	Requisito
Sistema operativo Oracle Solaris	Las versiones del sistema operativo Oracle Solaris compatibles con Oracle Solaris Cluster.   Todos los nodos deben utilizar la misma versión del sistema operativo Oracle Solaris.   Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software Oracle Solaris Cluster</a>
Software Oracle Solaris Cluster	Software Oracle Solaris Cluster 3.3.   Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software Oracle Solaris Cluster</a>.
Software del administrador de volúmenes	Software Solaris Volume Manager o Veritas Volume Manager (VxVM).   Todos los nodos deben usar la misma versión del administrador de volúmenes.   Para obtener más información sobre la instalación, consulte el <a href="#">Capítulo 4</a>, “Configuración del software Solaris Volume Manager” de <a href="#">Guía de instalación del software Oracle Solaris Cluster</a> y el <a href="#">Capítulo 5</a>, “Instalación y configuración de Veritas Volume Manager” de <a href="#">Guía de instalación del software Oracle Solaris Cluster</a>.
Software Sun StorageTek Availability Suite	Si desea más información sobre cómo instalar el software, consulte los manuales de instalación de su versión del software Sun StorageTek Availability Suite:      ■ Documentación de Sun StorageTek Availability Suite 4.0: Sun StorageTek Availability
Parches del software Sun StorageTek Availability Suite	Si desea información sobre los parches más actuales, consulte <a href="http://www.sunsolve.com">http://www.sunsolve.com</a>.

## Ejemplo de configuración de grupos de dispositivos y de recursos

Esta sección describe cómo se configuran los grupos de dispositivos y los de recursos en una aplicación NFS. Para obtener más información, consulte “[Configuración de grupos de recursos de replicaciones](#)” en la página 376 y “[Configuración de grupos de recursos de aplicaciones](#)” en la página 376.

Esta sección incluye los procedimientos siguientes:

- “[Configuración de un grupo de dispositivos en el clúster primario](#)” en la página 385
- “[Configuración de un grupo de dispositivos en el clúster secundario](#)” en la página 386
- “[Configuración de sistemas de archivos en el clúster primario para la aplicación NFS](#)” en la página 387

- “Configuración del sistema de archivos en el clúster secundario para la aplicación NFS” en la página 388
- “Creación de un grupo de recursos de replications en el clúster primario” en la página 390
- “Creación de un grupo de recursos de replications en el clúster secundario” en la página 391
- “Creación de un grupo de recursos de aplicaciones NFS en el clúster primario” en la página 393
- “Creación de un grupo de recursos de aplicaciones NFS en el clúster secundario” en la página 395
- “Procedimiento para comprobar la configuración de la replicación” en la página 403

La tabla siguiente enumera los nombres de los grupos y recursos creados para la configuración de ejemplo.

TABLA A-3 Resumen de los grupos y de los recursos en la configuración de ejemplo

Grupo o recurso	Nombre	Descripción
Grupo de dispositivos	devgrp	El grupo de dispositivos
El grupo de recursos de replications y los recursos	devgrp-stor-rg	El grupo de recursos de replications
	lhost-reprg-prim, lhost-reprg-sec	Los nombres de host lógicos para el grupo de recursos de replications en el clúster primario y el secundario
	devgrp-stor	El recurso de HAStoragePlus para el grupo de recursos de replications
El grupo de recursos de aplicaciones y los recursos	nfs-rg	El grupo de recursos de aplicaciones
	lhost-nfsrg-prim, lhost-nfsrg-sec	Los nombres de hosts lógicos para el grupo de recursos de aplicaciones en el clúster primario y el secundario
	nfs-dg-rs	El recurso de &HAStoragePlus para la aplicación
	nfs-rs	El recurso de NFS

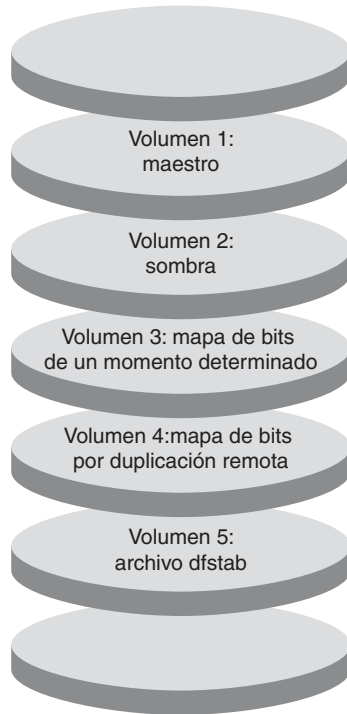
Con la excepción de devgrp-stor-rg, los nombres de los grupos y recursos son nombres de ejemplos que se pueden cambiar cuando sea necesario. El grupo de recursos de replications debe tener un nombre con el formato *nombre\_grupo\_dispositivos-stor-rg*.

Este ejemplo de configuración utiliza el software VxVM. Si desea más información sobre el software &SDSSVM, consulte el [Capítulo 4, “Configuración del software Solaris Volume Manager”](#) de *Guía de instalación del software Oracle Solaris Cluster*.

La figura siguiente ilustra los volúmenes creados en el grupo de dispositivos.



FIGURA A-8 Volúmenes del grupo de dispositivos



**Nota** – Los volúmenes definidos en este procedimiento no deben incluir las áreas privadas de etiquetas de discos, como el cilindro 0. El software VxVM administra automáticamente esta limitación.

## ▼ Configuración de un grupo de dispositivos en el clúster primario

### Antes de empezar

Asegúrese de realizar las tareas siguientes:

- Lea las directrices y los requisitos de las secciones siguientes:
  - “Comprensión del software Sun StorageTek Availability Suite en un clúster” en la página 372
  - “Directrices para la configuración de replicación de datos basada en host entre clústers” en la página 375
- Configure los clústers primario y secundario como se describe en “Conexión e instalación de clústers” en la página 381.

- 1 **Obtenga acceso al nodo nodeA como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`.**

El nodo nodeA es el primero del clúster primario. Si desea recordar qué nodo es nodeA, consulte la [Figura A-7](#).

- 2 **Cree un grupo de discos en el nodeA que contenga desde el volumen 1, vol01 hasta el volumen 4, vol04.**

Si desea información sobre cómo configurar un grupo de discos mediante VxVM, consulte el [Capítulo 5, “Instalación y configuración de Veritas Volume Manager” de Guía de instalación del software Oracle Solaris Cluster](#).

- 3 **Configure el grupo de discos para crear un grupo de dispositivos.**

```
nodeA# cldevicegroup create -t vxvm -n nodeA nodeB devgrp
```

El grupo de dispositivos se denomina devgrp.

- 4 **Cree el sistema de archivos para el grupo de dispositivos.**

```
nodeA# newfs /dev/vx/rdisk/devgrp/vol01 < /dev/null
nodeA# newfs /dev/vx/rdisk/devgrp/vol02 < /dev/null
```

No se necesita ningún sistema de archivos para el vol03 o el vol04, ya que se utilizan como volúmenes básicos.

**Pasos siguientes** Vaya a [“Configuración de un grupo de dispositivos en el clúster secundario” en la página 386](#).

## ▼ Configuración de un grupo de dispositivos en el clúster secundario

### Antes de empezar

Complete el procedimiento [“Configuración de un grupo de dispositivos en el clúster primario” en la página 385](#).

- 1 **Obtenga acceso al nodo nodeC como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`.**

- 2 **Cree un grupo de discos en el nodeC que contenga cuatro volúmenes: desde el volumen 1, vol01, hasta el volumen 4, vol04 .**

- 3 **Configure el grupo de discos para crear un grupo de dispositivos.**

```
nodeC# cldevicegroup create -t vxvm -n nodeC devgrp
```

El grupo de dispositivos se denomina devgrp.

- 4 **Cree el sistema de archivos para el grupo de dispositivos.**

```
nodeC# newfs /dev/vx/rdisk/devgrp/vol01 < /dev/null
nodeC# newfs /dev/vx/rdisk/devgrp/vol02 < /dev/null
```

No se necesita ningún sistema de archivos para el vol03 o el vol04, ya que se utilizan como volúmenes básicos.

**Pasos siguientes** Vaya a “Configuración de sistemas de archivos en el clúster primario para la aplicación NFS” en la página 387.

## ▼ Configuración de sistemas de archivos en el clúster primario para la aplicación NFS

**Antes de empezar** Complete el procedimiento “Configuración de un grupo de dispositivos en el clúster secundario” en la página 386.

- 1 **Conviértase en superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.admin` en el nodo `nodeA` y el nodo `nodeB`.**
- 2 **En el nodo `nodeA` y el nodo `nodeB`, cree un directorio de punto de montaje para el sistema de archivos NFS.**

Por ejemplo:

```
nodeA# mkdir /global/mountpoint
```

- 3 **En el nodo `nodeA` y el nodo `nodeB`, configure el volumen maestro para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo `nodeA` y el nodo `nodeB`. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol01 /dev/vx/rdisk/devgrp/vol01 \
/global/mountpoint ufs 3 no global,logging
```

Si desea recordar los nombres y los números de volúmenes utilizados en el grupo de dispositivos, consulte la [Figura A-8](#).

- 4 **En el nodo `nodeA`, cree un volumen para la información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.**

```
nodeA# vxassist -g devgrp make vol05 120m disk1
```

El volumen 5, `vol05`, contiene la información del sistema de archivos utilizada por el servicio de datos Oracle Solaris Cluster HA para NFS.

- 5 **En el nodo `nodeA`, vuelva a sincronizar el grupo de dispositivos con el software Oracle Solaris Cluster.**

```
nodeA# cldevicegroup sync devgrp
```

- 6 **En el nodo `nodeA`, cree el sistema de archivos para el `vol05`.**

```
nodeA# newfs /dev/vx/rdisk/devgrp/vol05
```

**7 En el nodo nodeA y el nodo nodeB, cree un punto de montaje para el vol05.**

El ejemplo siguiente crea el punto de montaje /global/etc.

```
nodeA# mkdir /global/etc
```

**8 En el nodo nodeA y el nodo nodeB, configure el vol05 para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo /etc/vfstab del nodo nodeA y el nodo nodeB. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol05 /dev/vx/rdisk/devgrp/vol05 \
/global/etc ufs 3 yes global,logging
```

**9 Monte el vol05 en el nodo nodeA.**

```
nodeA# mount /global/etc
```

**10 Haga que los sistemas remotos puedan tener acceso al vol05.****a. Cree un directorio con el nombre /global/etc/SUNW.nfs en el nodo nodeA.**

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

**b. Cree el archivo /global/etc/SUNW.nfs/dfstab.nfs-rs en el nodo nodeA.**

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

**c. Agregue la línea siguiente al archivo /global/etc/SUNW.nfs/dfstab.nfs-rs en el nodo nodeA.**

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

**Pasos siguientes** Vaya a [“Configuración del sistema de archivos en el clúster secundario para la aplicación NFS” en la página 388.](#)

## ▼ Configuración del sistema de archivos en el clúster secundario para la aplicación NFS

**Antes de empezar** Complete el procedimiento [“Configuración de sistemas de archivos en el clúster primario para la aplicación NFS” en la página 387.](#)

**1 Conviértase en superusuario o asuma una función que cuente con la autorización de RBAC solaris.cluster.admin en el nodo nodeC.****2 En el nodo nodeC, cree un directorio de punto de montaje para el sistema de archivos de NFS.**

Por ejemplo:

```
nodeC# mkdir /global/mountpoint
```

- 3 En el nodo nodeC, configure el volumen maestro para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo nodeC. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol01 /dev/vx/rdisk/devgrp/vol01 \
/global/mountpoint ufs 3 no global,logging
```

- 4 En el nodo nodeC, cree un volumen para la información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.**

```
nodeC# vxassist -g devgrp make vol05 120m disk1
```

El volumen 5, vol05, contiene la información del sistema de archivos utilizada por el servicio de datos Oracle Solaris Cluster HA para NFS.

- 5 En el nodo nodeC, vuelva a sincronizar el grupo de dispositivos con el software Oracle Solaris Cluster.**

```
nodeC# cldevicegroup sync devgrp
```

- 6 En el nodo nodeC, cree el sistema de archivos para el vol05.**

```
nodeC# newfs /dev/vx/rdisk/devgrp/vol05
```

- 7 En el nodo nodeC, cree un punto de montaje para el vol05.**

El ejemplo siguiente crea el punto de montaje `/global/etc`.

```
nodeC# mkdir /global/etc
```

- 8 En el nodo nodeC configure el vol05 para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo nodeC. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol05 /dev/vx/rdisk/devgrp/vol05 \
/global/etc ufs 3 yes global,logging
```

- 9 Monte el vol05 en el nodo nodeC.**

```
nodeC# mount /global/etc
```

- 10 Haga que los sistemas remotos puedan tener acceso al vol05.**

- a. Cree un directorio con el nombre `/global/etc/SUNW.nfs` en el nodo nodeC.**

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

- b. Cree el archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en el nodo nodeC.**

```
nodeC# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. Agregue la línea siguiente al archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en el nodo `nodeC`.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de replications en el clúster primario” en la página 390.](#)

## ▼ Creación de un grupo de recursos de replications en el clúster primario

**Antes de empezar** Complete el procedimiento [“Configuración del sistema de archivos en el clúster secundario para la aplicación NFS” en la página 388.](#)

- 1 **Obtenga acceso al nodo `nodeA` como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.**
- 2 **Registre el tipo de recurso de `SUNW.HAStoragePlus`.**  

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```
- 3 **Cree un grupo de recursos de replications para el grupo de dispositivos.**  

```
nodeA# clresourcegroup create -n nodeA,nodeB devgrp-stor-rg
```

`-n nodeA,nodeB` Especifica que los nodos del clúster `nodeA` y `nodeB` pueden controlar el grupo de recursos de replications.

`devgrp-stor-rg` El nombre del grupo de recursos de replications. En este nombre, `devgrp` especifica el nombre del grupo de dispositivos.
- 4 **Agregue un recurso `SUNW.HAStoragePlus` al grupo de recursos de replications.**  

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HAStoragePlus \
-p GlobalDevicePaths=devgrp \
-p AffinityOn=True \
devgrp-stor
```

`-g` Especifica el grupo de recursos en el que se agrega el recurso.

`-p GlobalDevicePaths=` Especifica la propiedad de la extensión en que se basa el software Sun StorageTek Availability Suite.

`-p AffinityOn=True` Especifica que el recurso `SUNW.HAStoragePlus` debe realizar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos definidos en `-x GlobalDevicePaths=`. Por ese motivo, si el grupo de recursos de replications migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

Para obtener más información sobre estas propiedades de extensión, consulte la página de comando `man SUNW.HAStoragePlus(5)`.

**5 Agregue un recurso de nombre de host lógico al grupo de recursos de replicaciones.**

```
nodeA# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-prim
```

El nombre de host lógico para el grupo de recursos de replicaciones en el clúster primario es `lhost-reprg-prim`.

**6 Habilite los recursos, administre el grupo de recursos y póngalo en línea.**

```
nodeA# clresourcegroup online -e -M -n nodeA devgrp-stor-rg
```

-e      Habilita los recursos asociados.

-M      Administra el grupo de recursos.

-n      Especifica el nodo en el que poner el grupo de recursos en línea.

**7 Compruebe que el grupo de recursos esté en línea.**

```
nodeA# clresourcegroup status devgrp-stor-rg
```

Examine el campo de estado del grupo de recursos para confirmar que el grupo de recursos de replicaciones esté en línea para el nodo `nodeA`.

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de replicaciones en el clúster secundario” en la página 391](#).

## ▼ Creación de un grupo de recursos de replicaciones en el clúster secundario

**Antes de empezar** Complete el procedimiento [“Creación de un grupo de recursos de replicaciones en el clúster primario” en la página 390](#).

**1 Obtenga acceso al nodo `nodeC` como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.**

**2 Registre `SUNW.HAStoragePlus` como tipo de recurso.**

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

**3 Cree un grupo de recursos de replicaciones para el grupo de dispositivos.**

```
nodeC# clresourcegroup create -n nodeC devgrp-stor-rg
```

`create`                      Crea el grupo de recursos.

-n                          Especifica la lista de nodos para el grupo de recursos.

devgrp El nombre del grupo de dispositivos.

devgrp-stor-rg El nombre del grupo de recursos de replicaciones.

#### 4 Agregue un recurso SUNW.HAStoragePlus al grupo de recursos de replicaciones.

```
nodeC# clresource create \
-t SUNW.HAStoragePlus \
-p GlobalDevicePaths=devgrp \
-p AffinityOn=True \
devgrp-stor
```

create Crea el recurso.

-t Especifica el tipo de recurso.

-p GlobalDevicePaths= Especifica la propiedad de la extensión en la que se basa Sun StorageTek Availability Suite.

-p AffinityOn=True Especifica que el recurso SUNW.HAStoragePlus debe realizar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos definidos en -x GlobalDevicePaths=. Por ese motivo, si el grupo de recursos de replicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

devgrp-stor El recurso de HAStoragePlus para el grupo de recursos de replicaciones.

Para obtener más información sobre estas propiedades de extensión, consulte la página de comando man [SUNW.HAStoragePlus\(5\)](#).

#### 5 Agregue un recurso de nombre de host lógico al grupo de recursos de replicaciones.

```
nodeC# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-sec
```

El nombre de host lógico para el grupo de recursos de replicaciones en el clúster primario es lhost-reprg-sec.

#### 6 Habilite los recursos, administre el grupo de recursos y póngalo en línea.

```
nodeC# clresourcegroup online -e -M -n nodeC devgrp-stor-rg
```

online Lo pone en línea.

-e Habilita los recursos asociados.

-M Administra el grupo de recursos.

-n Especifica el nodo en el que poner el grupo de recursos en línea.

#### 7 Compruebe que el grupo de recursos esté en línea.

```
nodeC# clresourcegroup status devgrp-stor-rg
```



Examine el campo de estado del grupo de recursos para confirmar que el grupo de recursos de replicaciones esté en línea para el nodo nodeC.

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de aplicaciones NFS en el clúster primario” en la página 393.](#)

## ▼ Creación de un grupo de recursos de aplicaciones NFS en el clúster primario

Este procedimiento describe la creación de los grupos de recursos de aplicaciones para NFS. Es un procedimiento específico de esta aplicación: no se puede usar para otro tipo de aplicación.

**Antes de empezar** Complete el procedimiento [“Creación de un grupo de recursos de replicaciones en el clúster secundario” en la página 391.](#)

- 1 **Obtenga acceso al nodo nodeA como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.**

- 2 **Registre `SUNW.nfs` como tipo de recurso.**

```
nodeA# clresourcetype register SUNW.nfs
```

- 3 **Si `SUNW.HAStoragePlus` no se ha registrado como tipo de recurso, hágalo.**

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

- 4 **Cree un grupo de recursos de aplicaciones para el grupo de dispositivos `devgrp`.**

```
nodeA# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_dependencies=devgrp-stor-rg \
nfs-rg
```

```
Pathprefix=/global/etc
```

Especifica el directorio en el que los recursos del grupo pueden guardar los archivos de administración.

```
Auto_start_on_new_cluster=False
```

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

```
RG_dependencies=devgrp-stor-rg
```

Especifica el grupo de recursos del que depende el grupo de recursos de aplicaciones. En este ejemplo, el grupo de recursos de aplicaciones depende del grupo de recursos de replicaciones `devgrp-stor-rg`.

Si el grupo de recursos de aplicaciones se conmuta a un nodo primario nuevo, el grupo de recursos de replications se conmuta automáticamente. Sin embargo, si el grupo de recursos de replications se conmuta a un nodo primario nuevo, el grupo de recursos de aplicaciones se debe conmutar de forma manual.

nfs-rg

El nombre del grupo de recursos de aplicaciones.

## 5 Agregue un recurso de SUNW.HAStoragePlus al grupo de recursos de aplicaciones.

```
nodeA# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
-p FileSystemMountPoints=/global/mountpoint \
-p AffinityOn=True \
nfs-dg-rs
```

create

Crea el recurso.

-g

Especifica el grupo de recursos en el que se agrega el recurso.

-t SUNW.HAStoragePlus

Especifica que el recurso es del tipo SUNW.HAStoragePlus.

-p FileSystemMountPoints=/global/

Especifica que el punto de montaje del sistema de archivos es global.

-p AffinityOn=True

Especifica que el recurso de aplicaciones debe efectuar una conmutación de afinidad en los dispositivos globales y en los sistemas de archivos del clúster definidos por -p GlobalDevicePaths=. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

nfs-dg-rs

El nombre del recurso de &HAStoragePlus para la aplicación NFS.

Para obtener más información sobre estas propiedades de extensión, consulte la página de comando [man SUNW.HAStoragePlus\(5\)](#).

## 6 Agregue un recurso de nombre de host lógico al grupo de recursos de aplicaciones.

```
nodeA# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-prim
```

El nombre de host lógico del grupo de recursos de aplicaciones del clúster primario es lhost-nfsrg-prim.

**7 Habilitar los recursos, administre el grupo de recursos de aplicaciones y deje en línea el grupo de recursos de aplicaciones.**

**a. Habilite el recurso de HASToragePlus para la aplicación NFS.**

```
nodeA# clresource enable nfs-rs
```

**b. Ponga en línea el grupo de recursos de aplicaciones en el nodo nodeA.**

```
nodeA# clresourcegroup online -e -M -n nodeA nfs-rg
```

`online`      Pone el grupo de recursos en línea.

`-e`            Habilita los recursos asociados.

`-M`            Administra el grupo de recursos.

`-n`            Especifica el nodo en el que poner el grupo de recursos en línea.

`nfs-rg`      El nombre del grupo de recursos.

**8 Compruebe que el grupo de recursos de aplicaciones esté en línea.**

```
nodeA# clresourcegroup status
```

Examine el campo de estado del grupo de recursos para determinar si el grupo de recursos de aplicaciones está en línea para el nodo nodeA y el nodo nodeB.

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de aplicaciones NFS en el clúster secundario” en la página 395.](#)

## ▼ Creación de un grupo de recursos de aplicaciones NFS en el clúster secundario

**Antes de empezar** Siga los pasos de [“Creación de un grupo de recursos de aplicaciones NFS en el clúster primario” en la página 393.](#)

**1 Obtenga acceso al nodo nodeC como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.**

**2 Registre `SUNW.nfs` como tipo de recurso.**

```
nodeC# clresourcetype register SUNW.nfs
```

**3 Si `SUNW.HASToragePlus` no se ha registrado como tipo de recurso, hágalo.**

```
nodeC# clresourcetype register SUNW.HASToragePlus
```

**4 Cree un grupo de recursos de replicaciones para el grupo de dispositivos.**

```
nodeC# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_dependencies=devgrp-stor-rg \
nfs-rg
```

create

Crea el grupo de recursos.

-p

Especifica una propiedad del grupo de recursos.

Pathprefix=/global/etc

Especifica un directorio en el que los recursos del grupo pueden guardar los archivos de administración.

Auto\_start\_on\_new\_cluster=False

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

RG\_dependencies=devgrp-stor-rg

Especifica los grupos de recursos de los que depende el grupo de recursos de aplicaciones. En este ejemplo, el grupo de recursos de aplicaciones depende del grupo de recursos de replicaciones.

Si el grupo de recursos de aplicaciones se conmuta a un nodo primario nuevo, el grupo de recursos de replicaciones se conmuta automáticamente. Sin embargo, si el grupo de recursos de replicaciones se conmuta a un nodo primario nuevo, el grupo de recursos de aplicaciones se debe conmutar de forma manual.

nfs-rg

El nombre del grupo de recursos de aplicaciones.

**5 Agregue un recurso de SUNW.HAStoragePlus al grupo de recursos de aplicaciones.**

```
nodeC# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
-p FileSystemMountPoints=/global/mountpoint \
-p AffinityOn=True \
nfs-dg-rs
```

create

Crea el recurso.

-g

Especifica el grupo de recursos en el que se agrega el recurso.

-t SUNW.HAStoragePlus

Especifica que el recurso es de tipo SUNW.HAStoragePlus.

-p

Especifica una propiedad del recurso.

```
FileSystemMountPoints=/global/
```

Especifica que el punto de montaje del sistema de archivos es global.

```
AffinityOn=True
```

Especifica que el recurso de aplicaciones debe efectuar una conmutación de afinidad en los dispositivos globales y en los sistemas de archivos del clúster definidos por `-x GlobalDevicePaths=`. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

```
nfs-dg-rs
```

El nombre del recurso de `&HASToragePlus` para la aplicación NFS.

Para obtener más información sobre estas propiedades de extensión, consulte la página de comando [man SUNW.HASToragePlus\(5\)](#).

## 6 Agregue un recurso de nombre de host lógico al grupo de recursos de aplicaciones.

```
nodeC# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-sec
```

El nombre de host lógico del grupo de recursos de aplicaciones del clúster secundario es `lhost-nfsrg-sec`.

## 7 Agregue un recurso de NSF al grupo de recursos de aplicaciones.

```
nodeC# clresource create -g nfs-rg \
-t SUNW.nfs -p Resource_dependencies=nfs-dg-rs nfs-rg
```

## 8 Compruebe que el grupo de recursos de aplicaciones no esté en línea en el nodo nodeC.

```
nodeC# clresource disable -n nodeC nfs-rs
nodeC# clresource disable -n nodeC nfs-dg-rs
nodeC# clresource disable -n nodeC lhost-nfsrg-sec
nodeC# clresourcegroup online -n "" nfs-rg
```

El grupo de recursos queda fuera de línea después de un re arranque, debido a `Auto_start_on_new_cluster=False`.

## 9 Si el volumen global se monta en el clúster primario, desmonte el volumen global del clúster secundario.

```
nodeC# umount /global/mountpoint
```

Si el volumen está montado en un clúster secundario, se da un error de sincronización.

**Pasos siguientes** Vaya a “Ejemplo de cómo habilitar la replicación de datos” en la página 397.

# Ejemplo de cómo habilitar la replicación de datos

Esta sección describe cómo habilitar la replicación de datos en la configuración de ejemplo. Esta sección utiliza los comandos `sndradm` e `iiadm` del software Sun StorageTek Availability Suite.

Para obtener más información sobre estos comandos, consulte la documentación de Sun StorageTek Availability *Sun Cluster 3.0 and Sun StorEdge Software Integration Guide*.

Esta sección incluye los procedimientos siguientes:

- “Habilitación de la replicación en el clúster primario” en la página 398
- “Habilitación de la replicación en el clúster secundario” en la página 400

## ▼ **Habilitación de la replicación en el clúster primario**

- 1 Acceda al nodo nodeA como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.read`.**

- 2 Purgue todas las transacciones.**

```
nodeA# lockfs -a -f
```

- 3 Compruebe que los nombres de hosts lógicos `lhost-reprg-prim` y `lhost-reprg-sec` estén en línea.**

```
nodeA# clresourcegroup status
nodeC# clresourcegroup status
```

Examine el campo de estado del grupo de recursos.

- 4 Habilite la duplicación por duplicación remota del clúster primario al secundario.**

Este paso habilita la replicación del volumen maestro del clúster primario en el volumen maestro del clúster secundario. Además, este paso permite la replicación en el mapa de bits duplicado remoto en el `vol04`.

- Si los clústeres primario y secundario no están sincronizados, ejecute este comando para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

- Si el clúster primario y el secundario están sincronizados, ejecute este comando para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -E lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

- 5 Habilite la sincronización automática.**

Ejecute este comando para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -a on lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
```

```
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Este paso habilita la sincronización automática. Si el estado activo de la sincronización automática es on, los conjuntos de volúmenes se vuelven a sincronizar cuando el sistema reinicie o cuando haya un error.

## 6 Compruebe que el clúster esté en modo de registro.

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

## 7 Habilite la instantánea de un momento determinado.

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/iadm -e ind \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
nodeA# /usr/sbin/iadm -w \
/dev/vx/rdisk/devgrp/vol02
```

Este paso habilita la copia del volumen maestro del clúster primario en el volumen sombra del mismo clúster. El volumen maestro, el volumen sombra y el volumen de mapa de bits de un momento determinado deben estar en el mismo grupo de dispositivos. En este ejemplo, el volumen maestro es vol01, el volumen sombra es vol02 y el volumen de mapa de bits de un momento determinado es vol03.

## 8 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -I a \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
```

Este paso asocia la instantánea de un momento determinado con el conjunto duplicado remoto de volúmenes. El software Sun StorageTek Availability Suite comprueba que se tome una instantánea de un momento determinado antes de que pueda haber una replicación por duplicación remota.

**Pasos siguientes** Vaya a [“Habilitación de la replicación en el clúster secundario”](#) en la página 400.

**▼** **Habilitación de la replicación en el clúster secundario**  
**Antes de empezar** Complete el procedimiento [“Habilitación de la replicación en el clúster primario”](#) en la página 398.

**1 Acceda al nodo nodeC como superusuario.**

**2 Purgue todas las transacciones.**

```
nodeC# lockfs -a -f
```

**3 Habilite la duplicación por duplicación remota del clúster primario al secundario.**

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeC# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

El clúster primario detecta la presencia del clúster secundario y comienza la sincronización. Si desea información sobre el estado de los clústeres, consulte el archivo de registro del sistema `/var/adm` de Sun StorageTek Availability Suite.

**4 Habilite la instantánea de un momento determinado independiente.**

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeC# /usr/sbin/iidm -e ind \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
nodeC# /usr/sbin/iidm -w \
/dev/vx/rdisk/devgrp/vol02
```

**5 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.**

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeC# /usr/sbin/sndradm -I a \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
```

**Pasos siguientes** Vaya a [“Ejemplo de cómo efectuar una replicación de datos”](#) en la página 400.

## Ejemplo de cómo efectuar una replicación de datos

Esta sección describe la ejecución de la replicación de datos en la configuración de ejemplo. Esta sección utiliza los comandos `sndradm` e `iidm` del software Sun StorageTek Availability Suite. Si



desea más información sobre estos comandos, consulte la documentación de Sun StorageTek Availability Suite *Sun Cluster 3.0 and Sun StorEdge Software Integration Guide*.

Esta sección incluye los procedimientos siguientes:

- “Replicación por duplicación remota” en la página 401
- “Instantánea de un momento determinado” en la página 402
- “Procedimiento para comprobar la configuración de la replicación” en la página 403

## ▼ Replicación por duplicación remota

En este procedimiento, el volumen maestro del disco primario se replica en el volumen maestro del disco secundario. El volumen maestro es vol01 y el de mapa de bits de duplicación remota es vol04.

### 1 Obtenga acceso al nodo nodeA como superusuario.

### 2 Compruebe que el clúster esté en modo de registro.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

### 3 Purgue todas las transacciones.

```
nodeA# lockfs -a -f
```

### 4 Repita los pasos del Paso 1 al Paso 3 en el nodo nodeC.

### 5 Copie el volumen principal del nodo nodeA en el volumen principal del nodo nodeC.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -m lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**6 Espere hasta que la termine replicación y los volúmenes se sincronicen.**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**7 Confirme que el clúster esté en modo de replicación.**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe en el volumen principal, el software Sun StorageTek Availability Suite, actualiza el volumen secundario.

**Pasos siguientes** Vaya a [“Instantánea de un momento determinado” en la página 402.](#)

**▼ Instantánea de un momento determinado**

En este procedimiento, la instantánea de un momento determinado se ha utilizado para sincronizar el volumen sombra del clúster primario con el volumen maestro del clúster primario. El volumen maestro es vol01, el de mapa de bits es vol04 y el de sombra es vol02.

**Antes de empezar** Realice el procedimiento [“Replicación por duplicación remota” en la página 401.](#)

**1 Obtenga acceso al nodo nodeA como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.modify` y `solaris.cluster.admin`.****2 Inhabilite el recurso que se esté ejecutando en el nodo nodeA.**

```
nodeA# clresource disable -n nodeA nfs-rs
```

**3 Cambie el clúster primario a modo de registro.**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

#### 4 Sincronice el volumen sombra del clúster primario con el volumen maestro del clúster primario.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/iiadm -u s /dev/vx/rdisk/devgrp/vol02
nodeA# /usr/sbin/iiadm -w /dev/vx/rdisk/devgrp/vol02
```

#### 5 Sincronice el volumen sombra del clúster secundario con el volumen maestro del clúster secundario.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeC# /usr/sbin/iiadm -u s /dev/vx/rdisk/devgrp/vol02
nodeC# /usr/sbin/iiadm -w /dev/vx/rdisk/devgrp/vol02
```

#### 6 Reinicie la aplicación en el nodo nodeA.

```
nodeA# clresource enable -n nodeA nfs-rs
```

#### 7 Vuelva a sincronizar el volumen secundario con el volumen principal.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**Pasos siguientes** Vaya a [“Procedimiento para comprobar la configuración de la replicación” en la página 403.](#)

## ▼ Procedimiento para comprobar la configuración de la replicación

### Antes de empezar

Complete el procedimiento [“Instantánea de un momento determinado” en la página 402.](#)

- 1 Obtenga acceso al nodo nodeA y al nodo nodeC como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.admin`.
- 2 Compruebe que el clúster primario esté en modo de replicación, con la sincronización automática activada.

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe en el volumen principal, el software Sun StorageTek Availability Suite, actualiza el volumen secundario.

### 3 Si el clúster primario no está en modo de replicación, póngalo en ese modo.

Utilice el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

### 4 Cree un directorio en un equipo cliente.

#### a. Regístrese en un equipo cliente como superusuario.

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

#### b. Cree un directorio en el equipo cliente.

```
client-machine# mkdir /dir
```

### 5 Monte el directorio de la aplicación en el clúster primario y visualice el directorio montado.

#### a. Monte el directorio de la aplicación en el clúster primario.

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

#### b. Visualice el directorio montado.

```
client-machine# ls /dir
```

### 6 Monte el directorio de la aplicación en el clúster secundario y visualice el directorio montado.

#### a. Desmonte el directorio de la aplicación en el clúster primario.

```
client-machine# umount /dir
```

#### b. Ponga el grupo de recursos de aplicaciones fuera de línea en el clúster primario.

```
nodeA# clresource disable -n nodeA nfs-rs
nodeA# clresource disable -n nodeA nfs-dg-rs
nodeA# clresource disable -n nodeA lhost-nfsrg-prim
nodeA# clresourcegroup online -n "" nfs-rg
```

#### c. Cambie el clúster primario a modo de registro.

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
```

```
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

**d. Asegúrese de que el directorio PathPrefix esté disponible.**

```
nodeC# mount | grep /global/etc
```

**e. Ponga en línea el grupo de recursos de aplicaciones en el clúster secundario.**

```
nodeC# clresourcegroup online -n nodeC nfs-rg
```

**f. Obtenga acceso al equipo cliente como superusuario.**

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

**g. Monte en la aplicación el directorio creado en el [Paso 4](#) en el clúster secundario.**

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

**h. Visualice el directorio montado.**

```
client-machine# ls /dir
```

**7 Compruebe que el directorio mostrado en el [Paso 5](#) sea el mismo del [Paso 6](#).**

**8 Devuelva la aplicación del clúster primario al directorio montado.**

**a. Ponga fuera de línea el grupo de recursos de aplicaciones del clúster secundario.**

```
nodeC# clresource disable -n nodeC nfs-rs
nodeC# clresource disable -n nodeC nfs-dg-rs
nodeC# clresource disable -n nodeC lhost-nfsrg-sec
nodeC# clresourcegroup online -n "" nfs-rg
```

**b. Compruebe que el volumen global esté desmontado desde el clúster secundario.**

```
nodeC# umount /global/mountpoint
```

**c. Ponga en línea el grupo de recursos de aplicaciones en el clúster primario.**

```
nodeA# clresourcegroup online -n nodeA nfs-rg
```

**d. Ponga el clúster primario en modo de replicación.**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe en el volumen principal, el software Sun StorageTek Availability Suite, actualiza el volumen secundario.

**Véase también** [“Ejemplo de cómo controlar una migración tras error” en la página 406](#)

## Ejemplo de cómo controlar una migración tras error

Esta sección describe cómo provocar una migración tras error y cómo transferir la aplicación al clúster secundario. Después de una migración tras error, actualice las entradas de DNS. Para obtener más información, consulte [“Directrices para administrar migraciones tras error” en la página 379](#).

Esta sección incluye los procedimientos siguientes:

- [“Procedimiento para propiciar una conmutación” en la página 406](#)
- [“Actualización de la entrada de DNS” en la página 407](#)

### ▼ Procedimiento para propiciar una conmutación

- 1 **Obtenga acceso al nodo nodeA y al nodo nodeC como superusuario o asuma una función que cuente con la autorización de RBAC `solaris.cluster.admin`.**

- 2 **Cambie el clúster primario a modo de registro.**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe en el volumen de datos del disco, se actualiza el volumen de mapa de bits del mismo grupo de dispositivos. No se produce ninguna replicación.

- 3 **Confirme que el clúster primario y el secundario estén en modo de registro, con la sincronización automática desactivada.**

- a. **En el nodo nodeA, confirme el modo y la configuración:**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync:off, max q writes:4194304,max q fbas:16384,mode:sync,ctag:
devgrp, state: logging
```

**b. En el nodo nodeC, confirme el modo y la configuración:**

Ejecute el comando siguiente para el software Sun StorageTek Availability Suite:

```
nodeC# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 <-
lhost-reprg-prim:/dev/vx/rdisk/devgrp/vol01
autosync:off, max q writes:4194304,max q fbas:16384,mode:sync,ctag:
devgrp, state: logging
```

En el nodo nodeA y el nodo nodeC, el estado debe ser logging y el estado activo de la sincronización automática debe ser off.

**4 Confirme que el clúster secundario esté listo para asumir las funciones del primario.**

```
nodeC# fsck -y /dev/vx/rdisk/devgrp/vol01
```

**5 Conmute al clúster secundario.**

```
nodeC# clresourcegroup switch -n nodeC nfs-rg
```

**Pasos siguientes** Vaya a [“Actualización de la entrada de DNS”](#) en la página 407.

## ▼ Actualización de la entrada de DNS

Si desea una ilustración de cómo asigna el DNS un cliente a un clúster, consulte la [Figura A-6](#).

### Antes de empezar

Realice el procedimiento [“Procedimiento para propiciar una conmutación”](#) en la página 406.

**1 Inicie el comando nsupdate.**

Para obtener más información, consulte la página de comando `man nsupdate(1M)`.

**2 Elimine en ambos clústers la asignación de DNS actual entre el nombre de host lógico del grupo de recursos de aplicaciones y la dirección IP del clúster.**

```
> 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
```

*dirección\_IP\_1\_inverso* Dirección IP del clúster primario en orden inverso.

*dirección\_IP\_2\_inverso* Dirección IP del clúster secundario en orden inverso.

*ttl* Tiempo de vida, en segundos. Un valor normal es 3600.

### 3 Cree la nueva asignación de DNS entre el nombre de host lógico del grupo de recursos de aplicaciones y la dirección IP del clúster en ambos clústers.

Asigne el nombre de host lógico principal a la dirección IP del clúster secundario y el nombre sistema lógico secundario a la dirección IP del clúster primario.

```
> 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
```

*dirección\_IP\_2* Dirección IP del clúster secundario hacia delante.

*dirección\_IP\_1* Dirección IP del clúster primario hacia delante.



# Índice

---

## A

- actualizar espacio de nombre global, 126
- actualizar manualmente información de DID, 184–185
- adaptadores de transporte, agregar, 226, 229
- administración, configuración de clúster
  - global, 259–301
- administración de energía, 259
- administradores de volúmenes, Veritas, 96
- administrar
  - clúster con herramienta de interfaz gráfica de usuario (GUI), 361–369
  - clústeres de zona, 18
  - clústeres globales, 18
  - clústers de zona, 281
  - dispositivos replicados basados en almacenamiento, 97–121
  - dispositivos replicados con Hitachi TrueCopy, 97–108
  - dispositivos replicados de EMC SRDF, 109–121
  - interconexiones de clústers y redes
    - públicas, 223–241
  - IPMP, 223–241
  - nodos sin votación de clústeres globales, 18
  - quórum, 189–215
  - sistema de archivos de clúster, 121
- agregar (*Continuación*)
  - dispositivos de quórum de disco compartido con conexión directa, 194
  - dispositivos de quórum de servidor de quórum, 200
  - dispositivos de quórum NAS de Sun, 195
  - dispositivos de quórum NAS de Sun Storage 7000 Unified Storage Systems, 195
  - funciones (RBAC), 53
  - funciones personalizadas (RBAC), 56
  - grupo de dispositivos, 132, 134–135
  - grupos de dispositivos de Solaris Volume Manager, 134
  - grupos de dispositivos de ZFS, 135
  - nodos, 243–249
    - nodos a un clúster de zona, 244
    - nodos a un clúster global, 244
    - nodos a un grupo de dispositivos, 153
  - nuevos volúmenes a grupos de dispositivos, 143
  - sistema de archivos de clúster, 174–178
  - sistemas SNMP, 294
  - usuarios de SNMP, 295
- anular registro
  - grupos de dispositivos, 152
  - grupos de dispositivos de Solaris Volume Manager, 138
- anular supervisión, rutas de disco, 183–184
- Aplicación de parches del software Oracle Solaris Cluster, 315–317
- aplicaciones de migración tras error para replicación de datos
  - administrar, 406–408
  - conmutación de afinidad, 376

aplicaciones de migración tras error para replicación de datos (*Continuación*)

- directrices

  - administración de migración tras error, 379

  - grupos de recursos, 377

aplicaciones escalables para replicación de datos, 378–379

aplicar

- parche que no sea de reinicio, 325

- parches, 317

aplicar parches, en nodos sin votación de clúster global, 320

archivo `/etc/inet/hosts`, configurar en zonas de direcciones IP exclusivas, 249

archivo `/etc/nsswitch.conf`, modificaciones en las zonas no globales, 248

archivo `/etc/vfstab`, 47

- agregar puntos de montaje, 175

- comprobar la configuración, 176

archivo `/var/adm/messages`, 82

archivo `hosts`, configurar en zonas de direcciones IP exclusivas, 249

archivo `lofi`, desinstalar, 289

archivo `md.tab`, 23

archivo `nsswitch.conf`, modificaciones en las zonas no globales, 248

archivo `ntp.conf.cluster`, 269

archivo `vfstab`

- agregar puntos de montaje, 175

- comprobar la configuración, 176

archivos

- `/etc/vfstab`, 47

- `md.conf`, 132

- `md.tab`, 23

- `ntp.conf.cluster`, 269

- restaurar de forma interactiva, 346

arrancar

- clúster de zona, 59–83

- clúster global, 59–83

- modo que no sea de clúster, 79

- nodos, 69–82

- nodos de clúster de zona, 69–82

- nodos de clúster global, 69–82

- zonas no globales, 69

arrancar en un modo que no sea de clúster, 79

Asistente para agregar funciones administrativas, descripción, 53

atributos, *Ver* propiedades

Availability Suite, usar para replicación de datos, 371

## B

buscar

- ID de nodo para clúster de zona, 262

- ID de nodo para clúster global, 262

- nombres de sistema de archivos, 334

## C

cables de transporte

- agregar, 226, 229

- habilitar, 232

- inhabilitar, 233

cambiar

- dirección de servidor de Oracle Solaris Cluster Manager, 365

- nombre de clúster, 260–261

- nombres de host privados, 267

- número de puerto, mediante Common Agent Container, 364

- protocolo de MIB de eventos de SNMP, 293

cambiar el nombre de los nodos

- en un clúster de zona, 272

- en un clúster global, 272

cerrar

- clúster de zona, 59–83

- clúster global, 59–83

- nodos, 69–82

- nodos de clúster de zona, 69–82

- nodos de clúster global, 69–82

- zonas no globales, 69

claves de seguridad, regenerar, 365

clúster

- aplicar parche de rearranque, 322

- autenticación de nodos, 262

- cambiar nombre, 260–261

- copia de seguridad, 23

- clúster (*Continuación*)
  - establecer la hora del día, 264
  - hacer copia de seguridad, 333–345
  - restaurar archivos, 345
- clúster de zona
  - administración, 259–301
  - apagar, 59–83
  - arrancar, 59–83
  - clonar, 281
  - definición, 18
  - eliminar un sistema de archivos, 281
  - estado de componentes, 32
  - mover una ruta de zona, 281
  - prepararlo para aplicaciones, 281
  - rearrancar, 65
  - validar configuración, 45
  - ver configuración, 36
- clúster global
  - administración, 259–301
  - arrancar, 59–83
  - cerrar, 59–83
  - definición, 18
  - eliminar nodos, 251
  - estado de componentes, 32
  - rearrancar, 65
  - validar configuración, 45
  - ver configuración, 36
- clústers de campus
  - recuperación con replicación de datos basada en almacenamiento, 91
  - replicación de datos basada en almacenamiento, 88–92
- clzonecluster
  - boot, 63–65
  - descripción, 27
  - halt, 59–69
- comando `/usr/cluster/bin/cluster check`,
  - comprobación del archivo `vfstab`, 176
- comando `boot`, 63–65
- comando `cconsole`, 22, 25
- comando `ccp`, 21, 25
- comando `claccess`, 21
- comando `cldevice`, 21
- comando `cldevicegroup`, 21
- comando `clinterconnect`, 21
- comando `clnasdevice`, 21
- comando `clnode check`, 21
- comando `clquorum`, 21
- comando `clreslogicalhostname`, 21
- comando `clresource`, 21
- comando `clresourcegroup`, 21
- comando `clresourcetype`, 21
- comando `clressharedaddress`, 21
- comando `clsetup`, 21
- comando `clsnmp host`, 21
- comando `clsnmp mib`, 21
- comando `clinterconnect`, 21
- comando `clnasdevice`, 21
- comando `clnode`, 278, 279–281
- comando `clnode check`, 21
- comando `clquorum`, 21
- comando `clreslogicalhostname`, 21
- comando `clresource`, 21
  - eliminar recursos y grupos de recursos, 282
- comando `clresourcegroup`, 21, 279–281
- comando `clresourcetype`, 21
- comando `clressharedaddress`, 21
- comando `clsnmp host`, 21
- comando `clsnmp mib`, 21
- comando `clsnmp user`, 21
- comando `cltelem attribute`, 21
- comando `cluster check`, 21
  - comprobación del archivo `vfstab`, 176
- comando `cluster shutdown`, 59–69
- comando `clzonecluster`, 21
- comando `crlogin`, 25
- comando `cssh`, 25
- comando `ctelnet`, 25
- comando `metaset`, 93–96
- comando `netcon`, 22
- comando `showrev -p`, 28, 29
- comandos
  - `boot`, 63–65
  - `cconsole`, 22, 25
  - `ccp`, 21, 25
  - `claccess`, 21
  - `cldevice`, 21
  - `cldevicegroup`, 21
  - `clinterconnect`, 21
  - `clnasdevice`, 21
  - `clnode check`, 21
  - `clquorum`, 21
  - `clreslogicalhostname`, 21
  - `clresource`, 21
  - `clresourcegroup`, 21
  - `clresourcetype`, 21
  - `clressharedaddress`, 21
  - `clsetup`, 21
  - `clsnmp host`, 21
  - `clsnmp mib`, 21

comandos (*Continuación*)

- clsnmpuser, 21
- cltelemetryattribute, 21
- cluster check, 21, 24, 45, 47
- cluster scshutdown, 59–69
- clzonecluster, 21, 59–69
- clzonecluster boot, 63–65
- clzonecluster verify, 45
- crlogin, 25
- cssh, 25
- ctelnet, 25
- metaset, 93–96
- netcon, 22

Common Agent Container, regenerar claves de seguridad, 365

Common Agent Container?, cambiar el número de puerto, 364

comprobar

- configuración de replicación de datos, 403–406
- configuración de vfstab, 176
- estado de interconexión de clúster, 225
- puntos de montaje globales, 47, 180
- SME, 247

conexiones seguras con consolas del clúster, 27

configuraciones de ejemplo (clúster de campus), replicación basada en almacenamiento de dos salas, 88–92

configurar

- dispositivos replicados con Hitachi TrueCopy, 98–100
- funciones (RBAC), 51
- número menor de un grupo de dispositivos, 144
- replicación de datos, 371–408

configurar límites de carga, en nodos, 279–281

conjunto de procesadores dedicado, configurar, 310

conmutación de afinidad, configurar para replicación de datos, 390

conmutación para replicación de datos

- conmutación de afinidad, 376
- realizar, 406–408

conmutadores de transporte, agregar, 226, 229

conmutar, nodo primario de grupo de dispositivos, 165–166

conmutar a nodo primario de grupo de dispositivos, 165–166

consola de administración, 22

consolas

- conectar a, 25
- conexiones seguras, 27

contenedores Oracle Solaris

- modificaciones en el archivo nsswitch.conf, 248
- zonas de direcciones IP compartidas, 247
- zonas de direcciones IP exclusivas
- configurar el archivo hosts, 249
- configurar grupos IPMP, 248

contenedores Solaris, propiedad autoboot, 247

control de acceso basado en funciones, *Ver* RBAC

convención de asignación de nombre, grupos de recursos de replicaciones, 376

convención de asignación de nombres, dispositivos de disco básicos, 175

copia de seguridad, clúster, 23

crear, nuevos grupos de discos, 137

**D**

desinstalar

- archivo de dispositivo de lofi, 289
- software de Oracle Solaris Cluster, 287

detener

- clúster de zona, 65
- clúster global, 65
- nodos, 69–82
- nodos de clúster de zona, 69–82
- nodos de clúster global, 69–82

direcciones IP, añadir a un servicio de nombres para zonas de direcciones IP exclusivas, 249

dispositivos

- globales, 93–188
- reconfiguración dinámica, 94–96

dispositivos de disco básicos, convenciones de asignación de nombres, 175

dispositivos de quórum

- agregar, 193
- dispositivos de quórum de almacenamiento conectados a red, 197

dispositivos de quórum, agregar (*Continuación*)  
 dispositivos de quórum de disco compartido con  
 conexión directa, 194  
 dispositivos de quórum de servidor de  
 quórum, 200  
 dispositivos de quórum NAS de Sun, 195  
 dispositivos de quórum NAS de Sun Storage 7000  
 Unified Storage Systems, 195  
 eliminar, 191  
 enumerar configuración en una lista, 213  
 estado de mantenimiento, poner un dispositivo  
 en, 210  
 estado de mantenimiento, sacar a un  
 dispositivo, 211  
 modificar listas de nodos, 207  
 quitar, 204  
 quitar el último dispositivo de quórum, 205  
 reconfiguración dinámica de dispositivos, 191  
 reemplazar, 206–207  
 reparar, 214  
 y replicación basada en almacenamiento, 91  
 dispositivos de quórum de almacenamiento conectados  
 a red, agregar e instalar, 197  
 dispositivos de quórum de disco compartido con  
 conexión directa, agregar, 194  
 dispositivos de quórum de servidor de quórum  
 agregar, 200  
 requisitos para instalar, 200  
 resolución de problemas para quitar  
 dispositivos, 205  
 dispositivos de quórum NAS de Sun, agregar, 195  
 dispositivos de quórum NAS de Sun Storage 7000  
 Unified Storage Systems, agregar, 195  
 dispositivos replicados basados en almacenamiento,  
 administrar, 97–121  
 Domain Name System (DNS)  
 actualizar en replicación de datos, 407–408  
 directrices para actualizar, 379  
 DR, *Ver* reconfiguración dinámica  
 duplicación local, *Ver* replicación basada en  
 almacenamiento  
 duplicación remota  
*Ver* replicación basada en almacenamiento  
 realizar, 401–402

duplicaciones, copia de seguridad en línea, 338

## E

ejemplos, crear un sistema de archivos de clúster, 177  
 ejemplos de configuración (clúster de campus),  
 replicación basada en almacenamiento de dos  
 salas, 88–92  
 eliminar  
 de un clúster de zona, 250  
 dispositivos de quórum, 191  
 grupos de dispositivos, 152  
 grupos de dispositivos de Solaris Volume  
 Manager, 138  
 hosts de SNMP, 294  
 matrices de almacenamiento, 255  
 nodos, 249, 251  
 nodos de grupos de dispositivos, 155  
 nodos de todos los grupos de dispositivos, 138  
 nodos sin votación en un clúster global, 254  
 recursos y grupos de recursos de un clúster de  
 zona, 282  
 usuarios de SNMP, 296  
 volúmenes de un grupo de dispositivos, 151–152  
 EMC SRDF  
 Adaptive Copy, 89  
 administrar, 109–121  
 comprobar la configuración, 113  
 configurar dispositivos de DID, 111–112  
 configurar un grupo de replicaciones, 109–111  
 ejemplo de configuración, 114–121  
 modo Domino, 89  
 prácticas recomendadas, 91  
 recuperar tras error completo con migración de sala  
 primaria de clúster de campus, 118–121  
 requisito, 90  
 restricciones, 90  
 encapsular discos, 141  
 enumerar, configuración de grupo de dispositivos, 163  
 enumerar en una lista, configuración de quórum, 213  
 espacio de nombre  
 global, 93–96, 126  
 migrar, 128

- espacio de nombre de dispositivos globales, migrar, 128
- establecer la hora del clúster, 264
- estado
  - componente de clúster de zona, 32
  - componente de clúster global, 32
- estado de mantenimiento
  - nodos, 274
  - poner un dispositivo de mantenimiento en, 210
  - sacar un dispositivo de quórum de, 211

## F

- fuera de servicio, dispositivo de quórum, 210
- función
  - agregar funciones, 53
  - agregar funciones personalizadas, 56
  - configurar, 51
- funciones admitidas, VxFS, 121
- funciones admitidas por VxFS, 121

## G

- global
  - espacio de nombre, 93–96
  - puntos de montaje, comprobar, 180
- globales
  - dispositivos, 93–188
    - configurar permisos, 94
  - puntos de montaje, comprobar, 47
- grupo de dispositivos de discos básicos, agregar, 134–135
- grupos de discos
  - crear, 137
  - modificar, 144
  - registrar, 145
  - registrar cambios en configuración, 148
- grupos de dispositivos
  - agregar, 134
  - asignación de un nuevo número menor, 144
  - comprobar registro, 151
  - configurar para replicación de datos, 384

- grupos de dispositivos (*Continuación*)
  - disco básico
    - agregar, 134–135
    - eliminar y anular registro, 138, 152
    - enumerar configuración, 163
    - estado de mantenimiento, 166
    - información general de administración, 123
    - modificar propiedades, 158
    - propiedad primaria, 158
  - SVM
    - agregar, 132
- grupos de recursos
  - replicación de datos
    - configurar, 376
    - directrices para configurar, 375
    - función en migración tras error, 376
- grupos de recursos de aplicaciones
  - configurar para replicación de datos, 393–395
  - directrices, 376
- grupos de recursos de direcciones compartidas para replicación de datos, 378

## H

- habilitar cables de transporte, 232
- habilitar e inhabilitar una MIB de eventos de SNMP, 292
- hacer copia de seguridad, clúster, 333–345
- herramienta de administración, 20
- herramienta de administración de GUI, 361–369
  - Oracle Solaris Cluster Manager, 361
  - Sun Management Center, 362
- herramienta de administración de línea de comandos, 20
- herramienta User Accounts, descripción, 57
- Hitachi TrueCopy
  - administrar, 97–108
  - comprobar configuración, 102–103
  - configurar dispositivos de DID, 100–102
  - configurar grupo de replicaciones, 98–100
  - ejemplo de configuración, 103–108
  - modos de datos o estado, 89
  - prácticas recomendadas, 91
  - requisitos, 90

**Hitachi TrueCopy (Continuación)**

restricciones, 90

**hosts**

agregar y eliminar SNMP, 294

**I**

imprimir, rutas de disco erróneas, 184

información de DID, actualizar

manualmente, 184–185

información de versión, 28, 29

información general, quórum, 189–215

inhabilitar cables de transporte, 233

**iniciar**

clúster de zona, 63–65

clúster global, 63–65

nodos, 69–82

nodos de clúster global, 69–82

Oracle Solaris Cluster Manager, 366

inicio de sesión, remoto, 25

inicio de sesión remoto, 25

instantánea, captura, 373

instantánea de un momento determinado

definición, 373

ejecutar, 402–403

instantáneas, *Ver* replicación basada en  
almacenamiento

interconexiones de clústers

administrar, 223–241

comprobar estado, 225

reconfiguración dinámica, 225

**IPMP**

administración, 238

estado, 35

grupos en zonas de direcciones IP exclusivas

configurar, 248

**K**/kernel/drv/, archivo `md.conf`, 132**L**

límites de carga

configurar en nodos, 278, 279–281

propiedad `concentrate_load`, 278propiedad `preemption_mode`, 278**M**

mantener, dispositivo de quórum, 210

mapa de bits

instantánea de un momento determinado, 373

replicación por duplicación remota, 372

matrices de almacenamiento, eliminar, 255

mensajes de error

archivo `/var/adm/messages`, 82

eliminar nodos, 257–258

**MIB**

cambiar protocolo de eventos de SNMP, 293

habilitar e inhabilitar un evento de SNMP, 292

MIB de eventos

cambiar protocolo SNMP, 293

habilitar e inhabilitar SNMP, 292

migración tras error de afinidad, propiedad de

extensión para replicación de datos, 376

migrar, espacio de nombre de dispositivos

globales, 128

modificar

grupos de discos, 144

listas de nodos de dispositivo de quórum, 207

nodos primarios, 165–166

propiedad `numsecondaries`, 160

propiedades, 158

usuarios (RBAC), 57

mostrar recursos configurados, 31

**N**NAS, *Ver* dispositivos de quórum de almacenamiento  
conectados a redNetApp, *Ver* dispositivos de quórum de  
almacenamiento conectados a redNetwork Appliance, *Ver* dispositivos de quórum de  
almacenamiento conectados a red

Network File System (NFS), configurar sistemas de archivos de aplicaciones para replicación de datos, 387–388

## nodos

agregar, 243–249

agregar a grupos de dispositivos, 153

aplicar un parche de re arranque, 317

arrancar, 69–82

autenticación, 262

buscar ID, 262

cambiar nombres en un clúster de zona, 272

cambiar nombres en un clúster global, 272

cerrar, 69–82

conectar a, 25

configurar límites de carga, 279–281

eliminar

mensajes de error, 257–258

eliminar de grupos de dispositivos, 138, 155

eliminar de un clúster de zona, 250

eliminar nodos de un clúster global, 251

eliminar nodos sin votación de un clúster global, 254

poner en estado de mantenimiento, 274

primarios, 94–96, 158

secundarios, 158

## nodos de clúster de zona

cerrar, 69–82

iniciar, 69–82

reiniciar, 76–79

## nodos de clúster global

arrancar, 69–82

cerrar, 69–82

reiniciar, 76–79

## nodos de votación de clúster global

administrar sistema de archivos de clúster, 121

recursos compartidos de CPU, 305

## nodos sin votación de clúster global

agregar nombre de host privado, 270

aplicar parches, 320

cambiar nombre de host privado, 270

cerrar y re arranque, 69

nombre de host privado, eliminar, 272

recursos compartidos de CPU, 307, 310

nodos sin votación de clústeres globales, administración, 18

nodos sin votación del clúster global, administrar sistema de archivos de clúster, 121

## nombres de host privados

asignar a zonas, 247

cambiar, 267

cambiar en nodos sin votación de clúster global, 270

eliminar nodos sin votación de clúster global, 272

nodos sin votación de clúster global, 270

número de puerto, cambiar mediante Common Agent Container, 364

## O

opciones de montaje para sistemas de archivos de clúster, requisitos, 176

Oracle Solaris Cluster Manager, 20, 361

cambiar la dirección de servidor, 365

funciones de RBAC, configurar, 363

iniciar, 366

## P

panel de control del clúster (CCP), 22

## parches

aplicar a clúster y a firmware, 322

aplicar un parche de re arranque, 317

aplicar un parche que no sea de reinicio, 325

sugerencias, 316

perfiles, derechos de RBAC, 52–53

perfiles de derechos, RBAC, 52–53

permisos, dispositivo global, 94

planificador por reparto equitativo, configuración de recursos compartidos de CPU, 304

## prácticas recomendadas

EMC SRDF, 91

Hitachi TrueCopy, 91

replicación de datos basada en almacenamiento, 91

procesador de servicios del sistema (SSP, System Service Processor), 22

PROM OpenBoot (OBP), 266

propiedad autoboot, 247



- propiedad failback, 158
- propiedad numsecondaries, 160
- propiedad primaria de grupos de dispositivos, 158
- propiedades
  - failback, 158
  - numsecondaries, 160
  - preferenced, 158
- propiedades de extensión para replicación de datos
  - recurso de aplicación, 394, 396
  - recurso de replicación, 390, 392
- puntos de montaje
  - globales, 47
  - modificación del archivo `/etc/vfstab`, 175

## Q

- quitar
  - cables, adaptadores y conmutadores de transporte, 229
  - dispositivos de quórum, 204
  - sistema de archivos de clúster, 178–180
  - último dispositivo de quórum, 205
- quórum
  - administración, 189–215
  - información general, 189–215

## R

- RBAC, 51–58
  - Oracle Solaris Cluster Manager, 363
  - para nodos de votación de clúster global, 52
  - para nodos sin votación, 52
  - perfiles de derechos (descripción), 52–53
  - tareas
    - agregar funciones, 53
    - agregar funciones personalizadas, 56
    - configurar, 51
    - modificar usuarios, 57
    - usar, 51
- realizar copia de seguridad
  - duplicaciones en línea, 338
  - sistema de archivos, 334
  - sistemas de archivos root, 335

- realizar copia de seguridad (*Continuación*)
  - volúmenes en línea, 341
- rearrancar, clúster de zona, 65
- reconfiguración dinámica, 94–96
  - dispositivos de quórum, 191
  - interconexiones de clústers, 225
  - interfaces de red pública, 240
- recuperación, clústers con replicación de datos basada en almacenamiento, 91
- recurso de nombres de host lógicos, función en migración tras error de replicación de datos, 376
- recursos
  - eliminar, 282
  - mostrar tipos configurados, 31
- recursos compartidos de CPU
  - configurar, 303
  - controlar, 303
  - nodos de votación de clúster global, 305
  - nodos sin votación de clúster global, 307
  - nodos sin votación de clúster global, conjunto de procesadores dedicado, 310
- red pública
  - administración, 223–241
  - reconfiguración dinámica, 240
- reemplazar dispositivos de quórum, 206–207
- regenerar, claves de seguridad, 365
- registrar
  - grupos de discos como grupos de dispositivos, 145
  - modificaciones en configuración de grupo de discos, 148
- reiniciar
  - clúster global, 65
  - nodos de clúster de zona, 76–79
  - nodos de clúster global, 76–79
- reparar, dispositivo de quórum, 214
- reparar sistema de archivos completo
  - `/var/adm/messages`, 82
- replicación, *Ver* replicación de datos
- replicación, basada en almacenamiento, 88–92
- replicación de datos, 85–92
  - actualizar una entrada DNS, 407–408
  - administrar una migración tras error, 406–408
  - asincrónica, 373
  - basada en almacenamiento, 86, 88–92

replicación de datos (*Continuación*)

- basada en host, 86
  - comprobar configuración, 403–406
  - configuración
    - conmutación de afinidad, 376
  - configuración de ejemplo, 380
  - configurar
    - conmutación de afinidad, 390
    - grupos de dispositivos, 384
    - grupos de recursos de aplicaciones NFS, 393–395
    - sistemas de archivos para una aplicación NFS, 387–388
  - definición, 86–87
  - directrices
    - administrar conmutación, 379
    - administrar migración tras error, 379
    - configurar grupos de recursos, 375
  - duplicación remota, 372, 401–402
  - ejemplo, 400–406
  - grupos de recursos
    - aplicación, 376
    - aplicaciones de migración tras error, 377
    - aplicaciones escalables, 378–379
    - configurar, 376
    - convención de asignación de nombre, 376
    - crear, 390–391
    - dirección compartida, 378
  - habilitar, 397–400
  - hardware y software necesarios, 382
  - instantánea de un momento determinado, 373, 402–403
  - presentación, 372
  - sincrónica, 373
- replicación de datos asincrónica, 89, 373
- replicación de datos basada en almacenamiento, 88–92
- definición, 86
  - prácticas recomendadas, 91
  - recuperación, 91
  - requisitos, 90
  - y dispositivos de quórum, 91
- replicación de datos basada en host
- definición, 86
  - ejemplo, 371–408

- replicación de datos sincrónica, 89, 373
- replicación por duplicación remota
- Ver* replicación basada en almacenamiento
  - definición, 372
- restaurar
- archivos de forma interactiva, 346
  - archivos del clúster, 345
  - sistema de archivos raíz
    - desde el metadispositivo, 349
    - desde el volumen, 349
  - sistema de archivos raíz encapsulados, 356
  - sistema de archivos root, 346
  - sistema de archivos root no encapsulado, 354
- ruta de disco
- anular supervisión, 183–184
  - resolver error de estado, 184–185
  - supervisar, 93–188
    - imprimir rutas de disco erróneas, 184
- ruta de zona, mover, 281
- rutas de disco compartido
- habilitar rearranque automático, 187–188
  - inhabilitar reinicio automático, 188
  - supervisar, 180–188

**S**

- SATA, 192, 194
- secundarios
- establecer número, 160
  - número predeterminado, 158
- servicio de nombres, añadir asignaciones de direcciones IP para zonas de direcciones IP exclusivas, 249
- servicios multiusuario, comprobar, 247
- servidores de quórum, *Ver* dispositivos de quórum de servidor de quórum
- shell seguro, 26, 27
- sistema de archivos
- aplicación NFS
    - configurar para replicación de datos, 387–388
  - buscar nombres, 334
  - quitarlo en un clúster de zona, 281
  - realizar copia de seguridad, 334
  - restaurar la raíz encapsulada, 356

- sistema de archivos (*Continuación*)
    - restaurar raíz
      - descripción, 346
      - desde el metadispositivo, 349
      - desde el volumen, 349
    - restaurar root no encapsulado, 354
  - sistema de archivos de clúster, 93–188
    - administración, 121
    - agregar, 174–178
    - comprobar la configuración, 176
    - nodos de votación de clúster global, 121
    - nodos sin votación de clúster global, 121
    - opciones de montaje, 176
    - quitar, 178–180
  - Sistema de archivos de Veritas (VxFS)
    - administrar, 177
    - montaje de sistemas de archivos de clúster, 177
  - sistema operativo Oracle Solaris
    - comando svcadm, 267
    - control CPU, 303
    - definición de clúster de zona, 17
    - definición de clúster global, 17
    - instrucciones especiales para iniciar nodos, 73–75
    - instrucciones especiales para reiniciar un nodo, 76–79
    - replicación basada en host, 87
    - SMF, 247
    - tareas administrativas de clúster global, 18
  - sistema operativo Solaris
    - Ver también* sistema operativo Oracle Solaris
  - sistemas de archivos globales, *Ver* sistemas de archivos de clúster
  - SMF, comprobar servicios en línea, 247
  - SNMP
    - agregar usuarios, 295
    - cambiar protocolo, 293
    - eliminar usuarios, 296
    - habilitar e inhabilitar una MIB de eventos, 292
    - habilitar hosts, 294
    - inhabilitar hosts, 294
  - Solaris Volume Manager, nombres de dispositivo de disco básico, 175
  - SRDF
    - Ver* EMC SRDF
  - ssh, 27
  - Sun Management Center
    - descripción, 20
    - información general, 362
    - instalación, 22
  - Sun StorageTek Availability Suite, usar para replicación de datos, 371
  - SunMC, *Ver* Sun Management Center
  - supervisar
    - rutas de disco, 181–183
    - rutas de disco compartido, 187–188
  - switchback, directrices para realizar replicación de datos, 380
- T**
- tolerancia ante problemas graves, definición, 372
  - transporte, adaptadores, 229
  - transporte, cables, 229
  - transporte, conmutadores, 229
  - TrueCopy
    - Ver* Hitachi TrueCopy
- U**
- último dispositivo de quórum, quitar, 205
  - usar, funciones (RBAC), 51
  - /usr/cluster/bin/clresource, eliminar grupos de recursos, 282
  - usuarios
    - agregar SNMP, 295
    - eliminar SNMP, 296
    - modificar propiedades, 57
  - utilidad clsetup, 20, 21, 27
- V**
- validar
    - configuración de clúster de zona, 45
    - configuración de clúster global, 45
  - ver
    - configuración de clúster de zona, 36

ver (*Continuación*)

configuración de clúster global, 36

Veritas

administración, 96

copias de seguridad en línea, 341

restaurar sistema de archivos raíz encapsulados, 356

restaurar sistema de archivos root no  
encapsulado, 354

Veritas Volume Manager (VxVM), nombres de

dispositivos de disco básico, 175

volumen, *Ver* replicación basada en almacenamiento

volúmenes

agregar a grupos de dispositivos, 143

eliminar de grupos de dispositivos, 151–152

realizar una copia de seguridad en línea, 341

VxVM, 96

## **Z**

ZFS

agregar grupos de dispositivos, 135

eliminar sistema de archivos, 283–285

replicación, 135

restricciones para sistemas de archivos  
root, 121–122

zonas de direcciones IP compartidas, *Ver* contenedores

Oracle Solaris

zonas de direcciones IP exclusivas, *Ver* contenedores

Oracle Solaris