

Guía de administración del sistema de Oracle® Solaris Cluster

Copyright © 2000, 2012, Oracle y/o sus filiales. 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 END USERS:

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

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.

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. AMD, Opteron, el logotipo de AMD y el logotipo de AMD Opteron son marcas comerciales o marcas comerciales registradas de Advanced Micro Devices. UNIX es una marca comercial registrada de The Open Group.

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	15
Descripción general de la administración de Oracle Solaris Cluster	16
Trabajo con un clúster de zona	16
Restricciones de las funciones del sistema operativo Oracle Solaris	17
Herramientas de administración	18
Interfaz de línea de comandos	18
Preparación para administrar el clúster	20
Documentación de las configuraciones de hardware de Oracle Solaris Cluster	20
Uso de la consola de administración	20
Copia de seguridad del clúster	20
Procedimiento para empezar a administrar el clúster	21
Inicio de sesión de manera remota en el clúster	22
Conexión segura a las consolas del clúster	23
▼ Obtención de acceso a las utilidades de configuración del clúster	23
▼ Visualización de la información de versión de Oracle Solaris Cluster	24
▼ Visualización de tipos de recursos configurados, grupos de recursos y recursos	26
▼ Comprobación del estado de los componentes del clúster	27
▼ Comprobación del estado de la red pública	29
▼ Visualización de la configuración del clúster	30
▼ Validación de una configuración básica de clúster	38
▼ Comprobación de los puntos de montaje globales	44
▼ Visualización del contenido de los registros de comandos de Oracle Solaris Cluster	46
2 Oracle Solaris Cluster y RBAC	49
Instalación y uso de RBAC con Oracle Solaris Cluster	49

Perfiles de derechos de RBAC en Oracle Solaris Cluster	50
Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster	51
▼ Creación de una función desde la línea de comandos	51
Modificación de las propiedades de RBAC de un usuario	53
▼ Modificación de las propiedades de RBAC de un usuario desde la línea de comandos	53
3 Cierre y arranque de un clúster	55
Descripción general sobre el cierre y el arranque de un clúster	55
▼ Cierre de un clúster	57
▼ Arranque de un clúster	59
▼ Rearranque de un clúster	61
Cierre y arranque de un solo nodo de un clúster	65
▼ Cierre de un nodo	66
▼ Arranque de un nodo	69
▼ Rearranque de un nodo	71
▼ Rearranque de un nodo en un modo que no sea de clúster	74
Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad	77
▼ Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad	77
4 Métodos de replicación de datos	79
Comprensión de la replicación de datos	80
Métodos admitidos de replicación de datos	80
5 Administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de clúster	83
Información general sobre la administración de dispositivos globales y el espacio de nombre global	83
Permisos de dispositivos globales en Solaris Volume Manager	84
Reconfiguración dinámica con dispositivos globales	84
Información general sobre la administración de sistemas de archivos de clústeres	85
Restricciones del sistema de archivos de clúster	85
Administración de grupos de dispositivos	86
▼ Actualización del espacio de nombre de dispositivos globales	88
▼ Cómo cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres	

de dispositivos globales	89
Migración del espacio de nombre de dispositivos globales	90
▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de <code>lofi</code>	91
▼ Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada	92
Adición y registro de grupos de dispositivos	93
▼ Adición y registro de un grupo de dispositivos (Solaris Volume Manager)	93
▼ Adición y registro de un grupo de dispositivos (disco básico)	96
▼ Adición y registro de un grupo de dispositivos replicado (ZFS)	96
Mantenimiento de grupos de dispositivos	98
Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)	98
▼ Eliminación de un nodo de todos los grupos de dispositivos	98
▼ Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)	99
▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos	101
▼ Cambio de propiedades de los grupos de dispositivos	103
▼ Establecimiento del número de secundarios para un grupo de dispositivos	105
▼ Enumeración de la configuración de grupos de dispositivos	107
▼ Conmutación al nodo primario de un grupo de dispositivos	109
▼ Colocación de un grupo de dispositivos en estado de mantenimiento	110
Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento	112
▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento	112
▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento	113
▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento	114
▼ Modificación del protocolo de protección en un solo dispositivo de almacenamiento	115
Administración de sistemas de archivos de clúster	117
▼ Adición de un sistema de archivos de clúster	117
▼ Eliminación de un sistema de archivos de clúster	121
▼ Comprobación de montajes globales en un clúster	123
Administración de la supervisión de rutas de disco	123
▼ Supervisión de una ruta de disco	124
▼ Anulación de la supervisión de una ruta de disco	126
▼ Impresión de rutas de disco erróneas	126

▼ Resolución de un error de estado de ruta de disco	127
▼ Supervisión de rutas de disco desde un archivo	127
▼ Habilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	129
▼ Inhabilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	130
6 Administración de quórum	131
Administración de dispositivos de quórum	131
Reconfiguración dinámica con dispositivos de quórum	133
Adición de un dispositivo de quórum	134
Eliminación o sustitución de un dispositivo de quórum	142
Mantenimiento de dispositivos de quórum	145
Cambio del tiempo de espera predeterminado del quórum	153
Administración de servidores de quórum de Oracle Solaris Cluster	153
Inicio y detención del software del servidor del quórum	154
▼ Inicio de un servidor de quórum	154
▼ Detención de un servidor de quórum	155
Visualización de información sobre el servidor de quórum	156
Limpieza de la información caducada sobre clústeres del servidor de quórum	157
7 Administración de interconexiones de clústeres y redes públicas	161
Administración de interconexiones de clústeres	161
Reconfiguración dinámica con interconexiones de clústeres	163
▼ Comprobación del estado de la interconexión de clúster	163
▼ Adición de dispositivos de cable de transporte de clúster, adaptadores o conmutadores de transporte	164
▼ Eliminación de cable de transporte de clúster, adaptadores de transporte y conmutadores de transporte	167
▼ Habilitación de un cable de transporte de clúster	169
▼ Inhabilitación de un cable de transporte de clúster	171
▼ Determinación del número de instancia de un adaptador de transporte	172
▼ Modificación de la dirección de red privada o del intervalo de direcciones de un clúster	173
Administración de redes públicas	176
Administración de grupos de varias rutas de red IP en un clúster	176
Reconfiguración dinámica con interfaces de red pública	177

8 Adición y eliminación de un nodo	179
Adición de nodos a un clúster	179
▼ Cómo agregar un nodo a un clúster existente	180
Eliminación de nodos de un clúster	182
▼ Eliminación de un nodo de un clúster de zona	184
▼ Eliminación de un nodo de la configuración de software del clúster	184
▼ Eliminación de la conectividad entre una matriz y un único nodo, en un clúster con conectividad superior a dos nodos	187
▼ Corrección de mensajes de error	189
9 Administración del clúster	191
Información general sobre la administración del clúster	191
▼ Cambio del nombre del clúster	192
▼ Asignación de un ID de nodo a un nombre de nodo	194
▼ Uso de la autenticación del nodo del clúster nuevo	194
▼ Restablecimiento de la hora del día en un clúster	196
▼ SPARC: Visualización de la PROM OpenBoot en un nodo	198
▼ Cómo cambiar un nombre de host privado de nodo	199
▼ Cómo cambiar el nombre de un nodo	202
▼ Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes	203
▼ Cómo poner un nodo en estado de mantenimiento	204
▼ Cómo sacar un nodo del estado de mantenimiento	206
▼ Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster	208
Resolución de problemas de desinstalación de nodos	211
Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster	212
Configuración de límites de carga	218
Cambio de números de puerto para servicios o agentes de gestión	220
Cómo realizar tareas administrativas del clúster de zona	222
▼ Cómo eliminar un clúster de zona	223
▼ Cómo eliminar un sistema de archivos de un clúster de zona	224
▼ Cómo eliminar un dispositivo de almacenamiento de un clúster de zona	226
Resolución de problemas	228
Ejecución de una aplicación fuera del clúster global	228
Restaure un conjunto de discos dañado	230

10	Configuración del control del uso de la CPU	235
	Introducción al control de la CPU	235
	Elección de una situación hipotética	235
	Planificador por reparto equitativo	236
	Configuración del control de la CPU	236
	▼ Control del uso de la CPU en el nodo de votación de un clúster global	236
11	Actualización de software	239
	Descripción general de actualización de software de Oracle Solaris Cluster	239
	Actualización del software de Oracle Solaris Cluster	240
	Actualización del clúster a una nueva versión	240
	Actualización a un paquete específico	241
	Actualización de un servidor de quórum o servidor de instalación AI	242
	Desinstalación de un paquete	243
	▼ Cómo desinstalar un paquete	243
	▼ Cómo desinstalar paquetes de servidor de quórum o servidor de instalación AI	243
	Consejos de actualización	243
12	Copias de seguridad y restauraciones de clústeres	245
	Copia de seguridad de un clúster	245
	▼ Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)	245
	▼ Copias de seguridad de la configuración del clúster	247
	Restauración de los archivos del clúster	248
	▼ Cómo restaurar el sistema de archivos raíz ZFS (/) (Solaris Volume Manager)	248
A	Ejemplo	251
	Configuración de replicación de datos basada en host con el software StorageTek Availability Suite	251
	Comprensión del software StorageTek Availability Suite en un clúster	252
	Directrices para la configuración de replicación de datos basada en host entre clústeres	255
	Mapa de tareas: ejemplo de configuración de una replicación de datos	260
	Conexión e instalación de clústeres	261
	Ejemplo de configuración de grupos de dispositivos y de recursos	263
	Ejemplo de cómo habilitar la replicación de datos	276
	Ejemplo de cómo efectuar una replicación de datos	279

Ejemplo de cómo gestionar una recuperación	285
Índice	287

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 usan arquitecturas de las familias de procesadores SPARC y x86. En este documento, x86 hace referencia a la familia más amplia de productos compatibles con x86. 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 del software y hardware de Oracle. Este documento no se puede usar como una guía de planificación ni de preventas.

Las instrucciones incluidas en este manual presuponen un conocimiento del sistema operativo Oracle Solaris y el dominio del software Volume Manager que se utiliza con Oracle Solaris Cluster.

Bash es el shell predeterminado para Oracle Solaris 11. Los nombres de máquinas que se muestran con la solicitud del shell Bash se indican con fines de aclaración.

Uso de los comandos de UNIX

Este documento contiene información sobre los comandos específicos para la instalación y la configuración de los servicios de datos de Oracle Solaris Cluster. Este documento *no* contiene información exhaustiva acerca de los comandos y los procedimientos básicos de UNIX como el cierre o el arranque del sistema, o la configuración de los dispositivos. Puede encontrar información sobre los comandos y procedimientos básicos de UNIX en las fuentes siguientes:

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

Convenciones tipográficas

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

TABLA P-1 Convenciones tipográficas

Tipos de letra	Descripción	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> <code>Password:</code>
<i>aabbcc123</i>	Marcador de posición: debe sustituirse por un valor o nombre real.	El comando para eliminar un archivo es <code>rm nombre_archivo</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>antememoria</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	<code>machine_name%</code>

TABLA P-2 Indicadores del shell (Continuación)

Shell	Indicador
Shell C para superusuario	machine_name#

Documentación relacionada

Puede encontrar información sobre temas referentes a Oracle Solaris Cluster en la documentación enumerada en la tabla siguiente. Toda la documentación de Oracle Solaris Cluster está disponible en <http://www.oracle.com/technetwork/indexes/documentation/index.html>.

Tema	Documentación
Administración e instalación de software	<i>Oracle Solaris Cluster Hardware Administration Manual</i> Guías de administración de hardware individual
Conceptos	<i>Oracle Solaris Cluster Concepts Guide</i>
Instalación de software	<i>Guía de instalación del software de Oracle Solaris Cluster</i>
Administración e instalación de servicio de datos	<i>Oracle Solaris Cluster Data Services Planning and Administration Guide</i> y guías de servicio de datos individuales
Desarrollo de servicios 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> <i>Oracle Solaris Cluster Quick Reference</i>
Actualización de software	<i>Oracle Solaris Cluster Upgrade Guide</i>
Mensajes de error	<i>Oracle Solaris Cluster Error Messages Guide</i>
Referencias de comandos y funciones	<i>Oracle Solaris Cluster Reference Manual</i> <i>Oracle Solaris Cluster Data Services Reference Manual</i> <i>Oracle Solaris Cluster Geographic Edition Reference Manual</i> <i>Oracle Solaris Cluster Quorum Server Reference Manual</i>

Acceso a Oracle Support

Los clientes de Oracle tienen acceso a soporte electrónico por medio de My Oracle Support. Para obtener más información, visite <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=info> o, si tiene alguna discapacidad auditiva, visite <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=trs>.

Obtención de ayuda

Si tiene problemas al instalar o utilizar Oracle Solaris Cluster, póngase en contacto con su proveedor de servicios y proporcione la siguiente información.

- Su nombre y dirección de correo electrónico (si estuviera disponible)
- El nombre, dirección y número de teléfono de su empresa
- Los modelos y números de serie de sus sistemas
- El número de versión del entorno operativo (por ejemplo, Oracle Solaris 11)
- El número de versión de Oracle Solaris Cluster (por ejemplo, Oracle Solaris Cluster 4.0)

Use los comandos siguientes para reunir información sobre el 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>pkg list</code>	Indica los paquetes instalados.
<code>prtdiag -v</code>	Muestra información de diagnóstico del sistema.
<code>/usr/cluster/bin/clnode show -rev</code>	Muestra información de la versión del paquete y la versión de Oracle Solaris Cluster para cada nodo.

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:

- “Descripción general de la administración de Oracle Solaris Cluster” en la página 16
- “Restricciones de las funciones del sistema operativo Oracle Solaris” en la página 17
- “Herramientas de administración” en la página 18
- “Preparación para administrar el clúster” en la página 20
- “Procedimiento para empezar a administrar el clúster” en la página 21

Todos los procedimientos indicados en esta guía son para su uso en el sistema operativo Oracle Solaris 11.

Un clúster global se compone de solamente uno o más nodos de votación de clúster global. Un clúster global también incluye la marca `solaris`, zonas no globales que no son nodos, pero sí contenedores de alta disponibilidad (como recursos) configurados con HA para el servicio de datos de zonas. Un clúster de zona precisa de un clúster global. Para obtener información general sobre los clústeres de zona, consulte la [Oracle Solaris Cluster Concepts Guide](#).

Un clúster de zona se compone de uno o más nodos de votación de marca `solaris`. Todos los nodos de un clúster de zona se configuran como zonas no globales de la marca `solaris` que se establecen con el atributo `cluster`. No se permite ningún otro tipo de marca en un clúster de zona. Puede ejecutar servicios compatibles en el clúster de zona del mismo modo que en un clúster global, con el aislamiento proporcionado por las zonas de Oracle Solaris. Un clúster de zona depende de, y por tanto requiere, un clúster global. Un clúster global no contiene un clúster de zona. Un clúster de zona tiene, como máximo, un nodo 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.

Descripción general de 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 el [Capítulo 1, “Planificación de la configuración de Oracle Solaris Cluster” de *Guía de instalación del software de Oracle Solaris Cluster*](#) y en la [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 de 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 el [Oracle Solaris Cluster Hardware Administration Manual](#).

En general, puede realizar tareas administrativas de Oracle Solaris Cluster mientras el clúster está en funcionamiento. 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 clústeres 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 clústeres de zona del clúster global. Los comandos de clústeres de zona también afectan a todos los nodos del clúster de zona, aunque un nodo esté inactivo al ejecutarse el comando.

Los clústeres de zona admiten la administración delegada de los recursos situados jerárquicamente bajo el control del administrador de grupos de recursos. Así, los administradores de clústeres de zona pueden ver las dependencias de dichos clústeres de zona que superan los límites de los clústeres 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 clústeres 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 de 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 184](#).
- Visualización de la configuración de un clúster de zona: consulte [“Visualización de la configuración del clúster” en la página 30](#).
- Validación de la configuración de un clúster de zona: consulte [“Validación de una configuración básica de clúster” en la página 38](#).
- 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 Servicios de Oracle Solaris Cluster

Servicios de Oracle Solaris Cluster	FMRI
<code>pnm</code>	<code>svc:/system/cluster/pnm:default</code>
<code>cl_event</code>	<code>svc:/system/cluster/cl_event:default</code>
<code>cl_eventlog</code>	<code>svc:/system/cluster/cl_eventlog:default</code>
<code>rpc_pmf</code>	<code>svc:/system/cluster/rpc_pmf:default</code>
<code>rpc_fed</code>	<code>svc:/system/cluster/rpc_fed:default</code>

TABLA 1-1 Servicios de Oracle Solaris Cluster (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 mediante la línea de comandos. La sección siguiente ofrece información general sobre las herramientas de línea de comandos.

Interfaz de línea de comandos

La mayoría de las tareas de administración de Oracle Solaris Cluster se pueden efectuar de manera interactiva con la utilidad `clsetup`. 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 siguiente lista se muestran otros comandos que se pueden utilizar para administrar una configuración de Oracle Solaris Cluster. Consulte las páginas de comando man si desea obtener información más detallada.

<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>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 gestor de volumen específico que el clúster utiliza.

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

Puede usar una estación de trabajo dedicada o una estación de trabajo conectada mediante una red de gestión como la *consola de administración* para administrar el clúster activo.

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.

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.

TABLA 1–2 Herramientas de administración de Oracle Solaris Cluster

Tarea	Herramienta	Instrucciones
Iniciar sesión en el clúster de forma remota	Use la utilidad <code>pconsole</code> de Oracle Solaris de la línea de comandos para iniciar sesión en el clúster de manera remota.	“Inicio de sesión de manera remota en el clúster” en la página 22 “Conexión segura a las consolas del clúster” en la página 23
Configurar el clúster de forma interactiva	Utilice el comando <code>clzonecluster</code> o la utilidad <code>clsetup</code> .	“Obtención de acceso a las utilidades de configuración del clúster” en la página 23
Mostrar información sobre el número y la versión de Oracle Solaris Cluster	Utilice el comando <code>clnode</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 24
Mostrar los recursos, grupos de recursos y tipos de recursos instalados	Utilice los comandos siguientes para mostrar la información sobre recursos: ■ <code>clresource</code> ■ <code>clresourcegroup</code> ■ <code>clresourcetype</code>	“Visualización de tipos de recursos configurados, grupos de recursos y recursos” en la página 26
Comprobar el estado de los componentes del clúster	Utilice el comando <code>cluster</code> con el subcomando <code>status</code> .	“Comprobación del estado de los componentes del clúster” en la página 27

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 status</code> con la opción <code>-m</code> . Para un clúster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	“Comprobación del estado de la red pública” en la página 29
Ver la configuración del clúster	Para un clúster global, utilice el comando <code>cluster</code> con el subcomando <code>show</code> . Para un clúster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	“Visualización de la configuración del clúster” en la página 30
Ver y mostrar los dispositivos NAS configurados	Para un clúster global o un clúster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	<code>clnasdevice(1CL)</code>
Comprobar los puntos de montaje globales o verificar la configuración del clúster	Para un clúster global, utilice el comando <code>cluster</code> con el subcomando <code>check</code> . En un clúster de zona, utilice el comando <code>clzonecluster verify</code> .	“Validación de una configuración básica de clúster” en la página 38
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 46
Examinar los mensajes del sistema del Oracle Solaris Cluster	Examine el archivo <code>/var/adm/messages</code> .	“Viewing System Messages” de <i>Oracle Solaris Administration: Common Tasks</i>
Supervisar el estado de Solaris Volume Manager	Utilice el comando <code>metastat</code> .	Solaris Volume Manager Administration Guide

Inicio de sesión de manera remota en el clúster

Puede utilizar la utilidad de acceso a la consola paralela (`pconsole`) de la línea de comandos para iniciar sesión en el clúster de manera remota. La utilidad `pconsole` es parte del paquete `terminal/pconsole` de Oracle Solaris. Instale el paquete ejecutando `pkg install terminal/pconsole`. La utilidad `pconsole` crea una ventana de terminal de host para cada host remoto que especifica en la línea de comandos. La utilidad también abre una ventana de consola central o principal que propaga lo que introduce allí a cada una de las conexiones que abre.

La utilidad `pconsole` se puede ejecutar desde Windows X o en el modo de consola. Instale `pconsole` en el equipo que utilizará como la consola de administración para el clúster. Si tiene un servidor de terminal que le permite conectarse a números de puertos específicos en la dirección IP del servidor, puede especificar el número de puerto además del nombre de host o dirección IP como `terminal-server:portnumber`.

Consulte la página del comando `man pconsole(1)` para obtener más información.

Conexión segura a las consolas del clúster

Si el concentrador de terminales o controlador del sistema admite `ssh`, puede utilizar la utilidad `pconsole` para conectarse a las consolas de esos sistemas. La utilidad `pconsole` es parte del paquete `terminal/pconsole` de Oracle Solaris y se instala cuando se instala el paquete. La utilidad `pconsole` crea una ventana de terminal de host para cada host remoto que especifica en la línea de comandos. La utilidad también abre una ventana de consola central o principal que propaga lo que introduce allí a cada una de las conexiones que abre. Consulte la página del comando `man pconsole(1)` para obtener más información.

▼ 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 clústeres 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 sczone
```

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

```
clzc:sczone> ?
```

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 de Oracle Solaris Cluster*.

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

▼ 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 1–1 Visualización de la información de versión de Oracle Solaris Cluster

En el siguiente ejemplo se muestra la información de versión del clúster y de versión de los paquetes que se enviaron con Oracle Solaris Cluster 4.0.

```
phys-schost# clnode show-rev
4.0
```

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

```
Oracle Solaris Cluster 4.0 for Solaris 11 sparc
ha-cluster/data-service/apache           :4.0.0-0.21
ha-cluster/data-service/dhcp             :4.0.0-0.21
ha-cluster/data-service/dns              :4.0.0-0.21
ha-cluster/data-service/ha-l1dom         :4.0.0-0.21
ha-cluster/data-service/ha-zones         :4.0.0-0.21
ha-cluster/data-service/nfs              :4.0.0-0.21
ha-cluster/data-service/oracle-database :4.0.0-0.21
ha-cluster/data-service/tomcat           :4.0.0-0.21
ha-cluster/data-service/weblogic         :4.0.0-0.21
ha-cluster/developer/agent-builder       :4.0.0-0.21
ha-cluster/developer/api                 :4.0.0-0.21
ha-cluster/geo/geo-framework             :4.0.0-0.21
ha-cluster/geo/manual                    :4.0.0-0.21
ha-cluster/geo/replication/availability-suite :4.0.0-0.21
ha-cluster/geo/replication/data-guard    :4.0.0-0.21
ha-cluster/geo/replication/sbp           :4.0.0-0.21
ha-cluster/geo/replication/srdf          :4.0.0-0.21
ha-cluster/group-package/ha-cluster-data-services-full :4.0.0-0.21
ha-cluster/group-package/ha-cluster-framework-full :4.0.0-0.21
ha-cluster/group-package/ha-cluster-framework-l10n :4.0.0-0.21
ha-cluster/group-package/ha-cluster-framework-minimal :4.0.0-0.21
ha-cluster/group-package/ha-cluster-framework-scm :4.0.0-0.21
ha-cluster/group-package/ha-cluster-framework-slm :4.0.0-0.21
ha-cluster/group-package/ha-cluster-full :4.0.0-0.21
ha-cluster/group-package/ha-cluster-geo-full :4.0.0-0.21
ha-cluster/group-package/ha-cluster-geo-incorporation :4.0.0-0.21
ha-cluster/group-package/ha-cluster-incorporation :4.0.0-0.21
ha-cluster/group-package/ha-cluster-minimal :4.0.0-0.21
ha-cluster/group-package/ha-cluster-quorum-server-full :4.0.0-0.21
ha-cluster/group-package/ha-cluster-quorum-server-l10n :4.0.0-0.21
ha-cluster/ha-service/derby              :4.0.0-0.21
ha-cluster/ha-service/gds                 :4.0.0-0.21
ha-cluster/ha-service/logical-hostname   :4.0.0-0.21
ha-cluster/ha-service/smf-proxy          :4.0.0-0.21
ha-cluster/ha-service/telemetry          :4.0.0-0.21
ha-cluster/library/cacao                 :4.0.0-0.21
ha-cluster/library/ucmm                   :4.0.0-0.21
ha-cluster/locale                         :4.0.0-0.21
ha-cluster/release/name                   :4.0.0-0.21
ha-cluster/service/management            :4.0.0-0.21
ha-cluster/service/management/slm        :4.0.0-0.21
ha-cluster/service/quorum-server         :4.0.0-0.21
ha-cluster/service/quorum-server/locale  :4.0.0-0.21
ha-cluster/service/quorum-server/manual/locale :4.0.0-0.21
ha-cluster/storage/svm-mediator           :4.0.0-0.21
ha-cluster/system/cfgchk                  :4.0.0-0.21
ha-cluster/system/core                    :4.0.0-0.21
ha-cluster/system/dsconfig-wizard        :4.0.0-0.21
ha-cluster/system/install                 :4.0.0-0.21
ha-cluster/system/manual                  :4.0.0-0.21
ha-cluster/system/manual/data-services   :4.0.0-0.21
ha-cluster/system/manual/locale           :4.0.0-0.21
```

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

`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.**

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

Siga todos los pasos de este procedimiento desde un nodo del clúster global. 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 1-2 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.sctelemetry
RT_description:	sctelemetry service for Oracle Solaris Cluster
RT_version:	1
API_version:	7
RT_basedir:	/usr/cluster/lib/rgm/rt/sctelemetry
Single_instance:	True
Proxy:	False
Init_nodes:	All potential masters
Installed_nodes:	<All>
Failover:	False
Pkglist:	<NULL>
RT_system:	True
Global_zone:	True

=== Resource Groups and Resources ===

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

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

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

▼ Comprobación del estado de los componentes del clúster

El comando `cluster status` 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.

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

Siga todos los pasos de este procedimiento desde un nodo del clúster global.

Ejemplo 1-3 Comprobación del estado de los componentes del clúster

En el siguiente ejemplo se muestra parte de la información de estado de los componentes del clúster devueltos por el comando `cluster 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:nge1	phys-schost-4:nge1	Path online
phys-schost-1:e1000g1	phys-schost-4:e1000g1	Path online

```
=== Cluster Quorum ===
```

```
--- Quorum Votes Summary ---
```

Needed	Present	Possible
3	3	4

```
--- Quorum Votes by Node ---
```

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

```
--- Quorum Votes by Device ---
```

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

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

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

Device Group Name	Primary	Secondary	Status
schost-2	phys-schost-2	-	Degraded

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

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

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

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

=== Cluster Resources ===

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

Device Instance	Node	Status
/dev/did/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

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 múltiples rutas de red IP, utilice el comando con el comando `clnode status`.

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

- **Compruebe el estado de los componentes del clúster.**

`phys-schost# clnode status -m`

Siga todos los pasos de este procedimiento desde un nodo del clúster global.

Ejemplo 1–4 Comprobación del estado de la red pública

En el siguiente ejemplo se muestra parte de la información de estado de los componentes del clúster devueltos 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  nge2     Online
phys-schost-2  test-rg         Online  nge3     Online
```

▼ **Visualización de la configuración 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.

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

- **Visualice la configuración de un clúster global o un clúster de zona.**

`% cluster show`

Siga todos los pasos de este procedimiento desde un nodo del clúster global.

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 clústeres de zona si es que están configurados.

También puede usar el comando `clzonecluster show` para visualizar la información de configuración sólo de los clústeres 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 del comando `man cluster(1CL)`. Consulte la página del comando `man clzonecluster(1CL)` si desea obtener más información sobre `clzonecluster`.

Ejemplo 1–5 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
clusterid:                   0x4DA2C888
installmode:                 disabled
heartbeat_timeout:           10000
heartbeat_quantum:           1000
private_netaddr:             172.11.0.0
private_netmask:             255.255.248.0
max_nodes:                   64
max_privatenets:             10
num_zoneclusters:            12
udp_session_timeout:         480
concentrate_load:            False
global_fencing:              prefer3
Node List:                   phys-schost-1
Node Zones:                  phys_schost-2:za
```

```
=== Host Access Control ===
```

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

```
=== Cluster Nodes ===
```

```
Node Name:                   phys-schost-1
Node ID:                     1
Enabled:                     yes
privatehostname:             clusternode1-priv
```

```
reboot_on_path_failure:      disabled
globalzoneshares:           3
defaulttpsetmin:             1
quorum_vote:                 1
quorum_defaultvote:          1
quorum_resv_key:             0x43CB1E1800000001
Transport Adapter List:      net1, net3
```

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

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

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

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

```
SNMP MIB Name:               Event
State:                       Disabled
Protocol:                     SNMPv2
```

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

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

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

```
Node Name:                    phys-schost-2:za
Node ID:                      2
Type:                          cluster
Enabled:                       yes
privatehostname:               clusternode2-priv
```

```

reboot_on_path_failure:      disabled
globalzoneshares:          1
defaultpsetmin:             2
quorum_vote:                1
quorum_defaultvote:         1
quorum_resv_key:            0x43CB1E1800000002
Transport Adapter List:     e1000g1, nge1

```

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

```

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

```

```

Transport Adapter:          nge1
Adapter State:              Enabled
Adapter Transport Type:     dlpi
Adapter Property(device_name): nge
Adapter Property(device_instance): 3
Adapter Property(lazy_free): 1
Adapter Property(dlpi_heartbeat_timeout): 10000
Adapter Property(dlpi_heartbeat_quantum): 1000
Adapter Property(nw_bandwidth): 80
Adapter Property(bandwidth): 10
Adapter Property(ip_address): 172.16.1.2
Adapter Property(netmask): 255.255.255.128
Adapter Port Names:         0
Adapter Port State(0):      Enabled

```

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

```

SNMP MIB Name:              Event
State:                       Disabled
Protocol:                    SNMPv2

```

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

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

=== Transport Cables ===

```

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

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

```

```
Cable Endpoint2:          switch1@1
Cable State:              Enabled

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

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

=== Transport Switches ===

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

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

=== Quorum Devices ===

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

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

=== Device Groups ===

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

=== Registered Resource Types ===

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

```
API_version: 7
RT_basedir: /usr/cluster/lib/rgm/rt/sctelemetry
Single_instance: True
Proxy: False
Init_nodes: All potential masters
Installed_nodes: <All>
Failover: False
Pkglist: <NULL>
RT_system: True
Global_zone: True
```

=== Resource Groups and Resources ===

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

--- Resources for Group HA_RG ---

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

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

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

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

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

```

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

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

=== DID Device Instances ===

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

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

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

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

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

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

=== NAS Devices ===

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

```

Ejemplo 1–6 Visualización de la información del clúster de zona

El siguiente ejemplo lista las propiedades de la configuración del clúster de zona con RAC.

```
% clzonecluster show
=== Zone Clusters ===

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

--- Solaris Resources for sczone ---

Resource Name:                                    net
  address:                                       172.16.0.1
  physical:                                     auto

Resource Name:                                    net
  address:                                       172.16.0.2
  physical:                                     auto

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

--- 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`. Para obtener más información, consulte la página del comando `man clnasdevice(1CL)`.

▼ Validación de una configuración básica de clúster

El comando `cluster` utiliza el subcomando `check` para validar la configuración básica que se requiere para que un clúster global funcione correctamente. Si ninguna comprobación falla, `cluster check` vuelve 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.

Además de las comprobaciones básicas, que se ejecutan sin la interacción del usuario, el comando también puede ejecutar comprobaciones interactivas y funcionales. Las comprobaciones básicas se ejecutan cuando la opción `-k palabra_clave` no se especifica.

- Las comprobaciones interactivas requieren información del usuario que las comprobaciones no pueden determinar. La comprobación solicita al usuario la información necesaria, por ejemplo, el número de versión del firmware. Utilice la palabra clave `-k interactive` para especificar una o más comprobaciones interactivas.
- Las comprobaciones funcionales ejercen una función o un comportamiento determinados del clúster. La comprobación solicita la entrada de datos del usuario, como, por ejemplo, qué nodo debe utilizar para la conmutación por error o la confirmación para iniciar o continuar con la comprobación. Utilice la palabra clave `-k funcional id_comprobación` para especificar una comprobación funcional. Realice sólo una comprobación funcional cada vez.

Nota – Dado que algunas comprobaciones funcionales implican la interrupción del servicio del clúster, no inicie ninguna comprobación funcional hasta que haya leído la descripción detallada de la comprobación y haya determinado si es necesario retirar antes el clúster de la producción. Para mostrar esta información, utilice el comando siguiente:

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

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 del 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 clústeres de zona.

1 Conviértase en superusuario de un nodo de miembro activo en un clúster global.

```
phys-schost# su
```

Siga todos los pasos de este procedimiento desde un nodo del clúster global.

2 Asegúrese de que dispone de las comprobaciones más actuales.

a. Vaya a la ficha Patches & Updates (Parches y actualizaciones) de [My Oracle Support](#).

b. En la búsqueda avanzada, seleccione **Solaris Cluster** como el producto y escriba **comprobación** en el campo de descripción.

La búsqueda localiza actualizaciones del software de Oracle Solaris Cluster que contienen comprobaciones.

c. Aplique las actualizaciones del software que no aún no estén instaladas en el clúster.

3 Ejecute las comprobaciones de validación básicas.

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

-v Modo detallado.

-o directorio_salida Redirige la salida al subdirectorio *dir_salida*.

El comando ejecuta todas las comprobaciones básicas disponibles. No se ve afectada ninguna función del clúster.

4 Ejecute las comprobaciones de validación interactivas.

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

-k interactive Especifica comprobaciones de validación interactivas en ejecución

El comando ejecuta todas las comprobaciones de validación interactivas disponibles y le solicita información necesaria sobre el clúster. No se ve afectada ninguna función del clúster.

5 Ejecute las comprobaciones de validación funcionales.

a. Enumere todas las comprobaciones funcionales disponibles en el modo detallado.

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

b. Determine qué comprobaciones funcionales realizan acciones que puedan afectar a la disponibilidad o los servicios del clúster en un entorno de producción.

Por ejemplo, una comprobación funcional puede desencadenar que el nodo genere avisos graves o una conmutación por error a otro nodo.

```
# cluster list-checks -v -C check-ID
```

-C ID_comprobación Especifica una comprobación específica.

c. Si hay peligro de que la comprobación funcional que desea efectuar interrumpa el funcionamiento del clúster, asegúrese de que el clúster no esté en producción.

d. Inicie la comprobación funcional.

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

-k functional Especifica comprobaciones de validación funcionales en ejecución.

Responda a las peticiones de la comprobación para confirmar la ejecución de la comprobación y para cualquier información o acciones que deba realizar.

e. Repita el Paso c y el Paso d para cada comprobación funcional que quede por ejecutar.

Nota – para fines de registro, especifique un único nombre de subdirectorio *dir_salida* para cada comprobación que se ejecuta. Si vuelve a utilizar un nombre *dir_salida*, la salida para la nueva comprobación sobrescribe el contenido existente del subdirectorio *dir_salida* reutilizado.

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

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

7 Grabe la configuración del clúster para poder realizar tareas de diagnóstico en el futuro.

Consulte “Cómo registrar los datos de diagnóstico de la configuración del clúster” de *Guía de instalación del software de Oracle Solaris Cluster*.

Ejemplo 1–7 Comprobación de la configuración del clúster global con resultado correcto en todas las comprobaciones básicas

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 1–8 Listado de comprobaciones de validación interactivas

En el siguiente ejemplo se enumeran todas las comprobaciones interactivas que están disponibles para ejecutarse en el clúster. En la salida del ejemplo aparece un muestreo de posibles comprobaciones; las comprobaciones disponibles reales varían en cada configuración.

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

```
Some checks might take a few moments to run (use -v to see progress)...
```

```
I6994574 : (Moderate) Fix for GLDv3 interfaces on cluster transport vulnerability applied?
```

Ejemplo 1–9 Ejecución de una comprobación de validación funcional

El siguiente ejemplo muestra primero el listado detallado de comprobaciones funcionales. La descripción detallada aparece en una lista para la comprobación F6968101, que indica que la comprobación podría alterar los servicios del clúster. El clúster se elimina de la producción. La comprobación funcional se ejecuta con salida detallada registrada en el subdirectorio `funct.test.F6968101.12Jan2011`. En la salida de ejemplo aparece un muestreo de posibles comprobaciones; las comprobaciones disponibles reales varían en cada configuración.

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

```
F6968101 : (Critical) Perform resource group switchover
```

```
F6984120 : (Critical) Induce cluster transport network failure - single adapter.
```

```
F6984121 : (Critical) Perform cluster shutdown
```

```
F6984140 : (Critical) Induce node panic
```

```
...
```

```
# cluster list-checks -v -C F6968101
```

```
F6968101: (Critical) Perform resource group switchover
```

```
Keywords: SolarisCluster3.x, functional
```

```
Applicability: Applicable if multi-node cluster running live.
```

```
Check Logic: Select a resource group and destination node. Perform  
'/usr/cluster/bin/clresourcegroup switch' on specified resource group  
either to specified node or to all nodes in succession.
```

```
Version: 1.2
```

```
Revision Date: 12/10/10
```

Take the cluster out of production

```
# cluster check -k functional -C F6968101 -o funct.test.F6968101.12Jan2011
```

```
F6968101
```

```
initializing...
```

```
initializing xml output...
```

```
loading auxiliary data...
```

```
starting check run...
```

```
pschost1, pschost2, pschost3, pschost4: F6968101.... starting:
```

```
Perform resource group switchover
```

```
=====
```

```
>>> Functional Check <<<
```

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

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

Select 'continue' for more details on this check.

- 1) continue
- 2) exit

choice: 1

```
=====
```

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

```
...
```

Follow onscreen directions

Ejemplo 1-10 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 4.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
```

▼ Comprobación de los puntos de montaje globales

El comando `cluster` 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. Para obtener más información, consulte la página del comando `man cluster(1CL)`.

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 1–11 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 4.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.phys-schost-1.txt
...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 1398
SEVERITY : HIGH
FAILURE : An unsupported server is being used as an Oracle Solaris Cluster 4.x node.
ANALYSIS : This server may not been qualified to be used as an Oracle Solaris Cluster 4.x node.
Only servers that have been qualified with Oracle Solaris Cluster 4.0 are supported as
Oracle Solaris Cluster 4.x nodes.
RECOMMEND: Because the list of supported servers is always being updated, check with
your Oracle representative to get the latest information on what servers
are currently supported and only use a server that is supported with Oracle Solaris Cluster 4.x.
...
#

```

▼ 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 del 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 1–12 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 49
- “Perfiles de derechos de RBAC en Oracle Solaris Cluster” en la página 50
- “Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster” en la página 51
- “Modificación de las propiedades de RBAC de un usuario” en la página 53

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>Oracle Solaris Administration: Security Services</i>
Instalar, administrar elementos y utilizar RBAC	Capítulo 9, “Using Role-Based Access Control (Tasks)” de <i>Oracle Solaris Administration: Security Services</i>
Más información sobre elementos y herramientas de RBAC	Capítulo 10, “Security Attributes in Oracle Solaris (Reference)” de <i>Oracle Solaris Administration: Security Services</i>

Perfiles de derechos de RBAC en Oracle Solaris Cluster

Los comandos y las opciones seleccionados de Oracle Solaris Cluster que se introducen en la línea de comandos usan RBAC para obtener autorización. Los comandos y las opciones de Oracle Solaris Cluster que requieren autorización RBAC necesitan uno o más de los siguientes niveles de autorización. 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 en un 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 del 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 rol
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 scha_control scha_resource_get scha_resource_setstatus scha_resourcegroup_get scha_resourcetype_get

Perfil de derechos	Incluye autorizaciones	Permiso de identidad de rol
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>	Realice las operaciones de enumeración, visualización y otras operaciones de lectura para los comandos de Oracle Solaris Cluster.
Funcionamiento 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>	Realice las operaciones de enumeración, visualización, exportación, estado y otras operaciones de lectura. Cambie 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.	Realice las mismas operaciones que la identidad de función de administración del clúster, además de otras operaciones de administración del sistema.
Gestión 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>	Realice 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 un rol de RBAC nuevo con un perfil de derechos de gestión de Oracle Solaris Cluster y para asignar a los usuarios este rol nuevo.

▼ 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 los roles en el ámbito local, use el comando `roleadd` para especificar un rol local nuevo y sus atributos. Para obtener más información, consulte la página del comando [man roleadd\(1M\)](#).

- Asimismo, para los roles en el ámbito local, edite el archivo `user_attr` para agregar un usuario con `type=role`. Para obtener más información, consulte la página del comando `man user_attr(4)`.

Use este método sólo en caso de emergencia.

- Para los roles en un servicio de nombres, use los comandos `roleadd` y `rolemod` para especificar el rol nuevo y sus atributos. Para obtener más información, consulte las páginas del comando `man roleadd(1M)` y `rolemod(1M)`.

Este comando requiere autenticación por parte del usuario o una función que, a su vez, pueda crear otras. Puede aplicar el comando `roleadd` a todos los servicios de nombres.

3 Inicie y detenga el daemon de antememoria del servicio de nombres.

Las funciones nuevas no se activan hasta que se reinicia el daemon de antememoria 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 RBAC de un usuario, consulte [“Modificación de las propiedades de RBAC de un usuario desde la línea de comandos” en la página 53.](#)

▼ 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 propiedades de usuario que están asignadas a un usuario definido en el ámbito local o en un depósito LDAP, use el comando `usermod`. Para obtener más información, consulte la página del comando `man usermod(1M)`.
 - Otra posibilidad es editar el archivo `user_attr`.
Use este método sólo en caso de emergencia.
 - Para gestionar roles localmente o en un servicio de nombres, como un depósito LDAP, use los comandos `roleadd` o `rolemod`. Para obtener más información, consulte las páginas del comando `man roleadd(1M)` o `rolemod(1M)`.
Estos comandos requieren autenticación como superusuario o como un rol que pueda modificar archivos de usuario. Puede aplicar estos comandos a todos los servicios de nombres. Consulte [“Command-Line Tools for User and Group Account Management” de Oracle Solaris Administration: Common Tasks.](#)

Los perfiles de derechos de detención y privilegio forzado que se proporcionan con Oracle Solaris 11 no se pueden modificar.

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.

- “Descripción general sobre el cierre y el arranque de un clúster” en la página 55
- “Cierre y arranque de un solo nodo de un clúster” en la página 65
- “Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad” en la página 77

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 74 y la Tabla 3–2.

Descripción general sobre el cierre y el arranque de un clúster

El comando `cluster shutdown` de Oracle Solaris Cluster detiene los servicios del clúster global de manera correcta y cierra 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 clústeres de zona. Para obtener más información, consulte la página del comando `man cluster(1CL)`.

En los procedimientos tratados en este capítulo, `phys - s chost#` 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` evacua para cerrar nodos independientes. Para obtener más información, consulte [“Cierre de un clúster” en la página 57](#), [“Cierre y arranque de un solo nodo de un clúster” en la página 65](#) o la página del comando `man clnode(1CL)`.

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. Para obtener más información, consulte [“Booting an x86 Based System Interactively” de *Booting and Shutting Down Oracle Solaris on x86 Platforms*](#). 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 74](#).

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

Tarea	Instrucciones
Detener el clúster.	“Cierre de un clúster” en la página 57
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 59
Rearrancar el clúster.	“Rearranque de un clúster” en la página 61

▼ Cierre de un clúster

Puede cerrar un clúster global, uno de zona o todos los clústeres 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 clúster de zona ejecutan Oracle Real Application Clusters (RAC), cierre todas las instancias de la base de datos del cluster que está cerrando.**
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 clústeres de zona.**
 - **Cierre el clúster global. Esta acción cierra también todos los clústeres de zona.**
`phys-schost# cluster shutdown -g0 -y`
 - **Cierre un clúster de zona concreto.**
`phys-schost# clzonecluster halt zoneclustername`
 - **Cierre todos los clústeres de zona.**
`phys-schost# clzonecluster halt +`
El comando `cluster shutdown` también puede usarse dentro de un clúster de zona para cerrar todos los clústeres 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 siguiente ejemplo se cierra un clúster de zona denominado *zona_sc*.

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

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 59](#) 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 de Oracle Solaris Cluster](#).

1 Inicie cada nodo en el 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.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System Interactively”](#) de *Booting and Shutting Down Oracle Solaris on x86 Platforms*.

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 clústeres 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 `cluster status` informa sobre el estado de los nodos del clúster global.

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

Si ejecuta el comando de estado `clzonecluster status` desde un nodo del clúster global, el 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 77. Para obtener más información, consulte la página del comando `man clzonecluster(1CL)`.

Ejemplo 3–4 SPARC: Arranque de un clúster global

El ejemplo siguiente muestra la salida de la consola cuando se arranca 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.

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System Interactively” de Booting and Shutting Down Oracle Solaris on x86 Platforms](#).

- 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 clústeres 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 77](#).

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

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 o un nodo del clúster de zona. Esta sección proporciona instrucciones para cerrar nodos del clúster global y nodos de clústeres 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.

Para obtener más información, consulte las páginas del comando `man 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 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 un nodo del clúster global, use los comandos <code>clnode evacuate</code> y <code>shutdown</code> . Para un nodo del clúster de zona, utilice el comando <code>clzonecluster halt</code> .	“Cierre de un nodo” en la página 66
Iniciar un nodo. El nodo debe disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembro de este último.	Para un nodo del clúster global, utilice el comando <code>boot</code> o <code>b</code> . Para un nodo del clúster de zona, utilice el comando <code>clzonecluster boot</code> .	“Arranque de un nodo” en la página 69
Detener y reiniciar (rearrancar) un nodo de un clúster. El nodo debe disponer de una conexión operativa con la interconexión de clúster para poder convertirse en miembro de este último.	Para un nodo del clúster global, utilice los comandos <code>clnode evacuate</code> y <code>shutdown</code> , seguidos de <code>boot</code> o <code>b</code> . Para un nodo del clúster de zona, utilice el comando <code>clzonecluster reboot</code> .	“Rearranque de un nodo” en la página 71

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

Tarea	Herramienta	Instrucciones
Arrancar un nodo de manera que no participe como miembro en el clúster.	Para un nodo del clúster global, use los comandos <code>clnode evacuate</code> y <code>shutdown</code> , 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 74

▼ 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 clústeres 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 grupos de dispositivos 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-szone' died.
phys-host#
phys-host# 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	Offline	Installed
	schost-4	sczone-4	Online	Running

```
phys-schost#
```

Véase también Consulte [“Arranque de un nodo” en la página 69](#) 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 reorganizan.

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.
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 clústeres 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 77](#).

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.

`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 finaliza el tiempo de espera de un método para un recurso y no se puede eliminar, el nodo se reiniciará sólo si la propiedad `Failover_mode` del recurso se establece en `HARD`. Si la propiedad `Failover_mode` se establece en cualquier otro valor, el nodo no se reiniciará.

- 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 clústeres 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 del nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos de las zonas globales del nodo especificado a las zonas globales de otros nodos que se sitúen a continuación en el orden de preferencia.

Nota – Para cerrar un único nodo, utilice el comando `shutdown -g0 -y -i6`. Para cerrar varios nodos al mismo tiempo, utilice el comando `shutdown -g0 -y -i0` para detenerlos. Después de detener todos los nodos, utilice el comando `boot` en todos ellos para volver a arrancarlos en el clúster.

- En un sistema basado en SPARC, ejecute los comandos siguientes para reiniciar un nodo único.

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

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

- En un sistema basado en x86, ejecute los comandos siguientes para reiniciar un nodo único.

```
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.

- 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 rearranca el nodo `phys-schost-1`. Los mensajes correspondientes a este nodo, como las notificaciones de cierre y arranque, 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 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-sczone
Waiting for zone reboot commands to complete on all the nodes of the zone cluster
"sparse-sczone"...
Sep  5 19:40:59 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster
'sparse-sczone' died.
phys-schost# Sep  5 19:41:27 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster
'sparse-sczone' joined.
```

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

=== Zone Clusters ===

--- Zone Cluster Status ---
Name           Node Name      Zone HostName   Status   Zone Status
-----
sparse-szone    schost-1       szone-1         Online   Running
                schost-2       szone-2         Online   Running
                schost-3       szone-3         Online   Running
                schost-4       szone-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 software del clúster o en ciertos procedimientos administrativos, como la actualización de 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 o el nodo del clúster global.**
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 de las zonas globales del nodo especificado a las zonas globales de otros nodos que se sitúen a continuación en el orden de preferencia.

- **Cierre un nodo del clúster global específico.**

```
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.**

Se muestra el menú de GRUB.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System Interactively”](#) de *Booting and Shutting Down Oracle Solaris on x86 Platforms*.

- 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.**

Se muestra la pantalla de parámetros de inicio de GRUB.

- 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/kernel/$ISADIR/unix -B $ZFS-BOOTFS -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.

- e. Escriba **b** para iniciar el nodo en el modo sin clúster.

Nota – Este cambio en el comando del parámetro de inicio del núcleo no se conserva tras el inicio del sistema. La siguiente vez que reinicie el nodo, se iniciará en el 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 3–13 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 software de Oracle Solaris Cluster escriben mensajes de error en el archivo `/var/adm/messages`, el cual puede llenar el sistema de archivos `/var` con el tiempo. Si el sistema de archivos `/var` del nodo del clúster alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda iniciarse en ese nodo en el próximo inicio. 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. Consulte [“Viewing System Messages” de Oracle Solaris Administration: Common Tasks](#) para obtener más información.

- 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 Solaris Volume Manager, ZFS o discos sin procesar, consulte [“Administración de grupos de dispositivos” en la página 86](#) en el capítulo 5.

En este capítulo, se incluye la siguiente sección:

- [“Comprensión de la replicación de datos” en la página 80](#)

Comprensión de la replicación de datos

Oracle Solaris Cluster 4.0 admite la replicación de datos basada en host.

La replicación de datos basada en host usa el software para replicar en tiempo real volúmenes de discos entre clústeres 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 encuentra en un clúster) es StorageTek Availability Suite.

La replicación de datos basada en host es una solución de replicación de datos económica porque usa recursos de 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 que ejecutan el sistema operativo Oracle Solaris escriban datos en un volumen compartido (por ejemplo, Oracle RAC). Para obtener más información sobre el uso de la replicación de datos basada en host entre dos clústeres, consulte la [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Solaris Availability Suite](#). Para ver un ejemplo de replicación basada en host que no use Oracle Solaris Cluster Geographic Edition, consulte el apéndice A, “Configuración de replicación de datos basada en host con el software StorageTek Availability Suite” en la página 251.

Métodos admitidos de replicación de datos

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

1. Replicación entre clústeres: para la recuperación después de un desastre, puede usar una replicación basada en host con el objeto de realizar una replicación de datos entre clústeres. Puede administrar la replicación con el software Oracle Solaris Cluster Geographic Edition.

- Replicación basada en host
 - StorageTek Availability Suite

Si desea utilizar la replicación basada en host sin el software Oracle Solaris Cluster Geographic Edition, consulte las instrucciones en el [Apéndice A, “Ejemplo”, “Configuración de replicación de datos basada en host con el software StorageTek Availability Suite”](#) en la página 251.

2. Replicación dentro de un clúster: este método se utiliza como sustitución para la duplicación basada en host.

3. Replicación basada en aplicaciones: Oracle Data Guard es un ejemplo de software de replicación basada en aplicaciones. Este tipo de software se utiliza solamente para la recuperación después de un desastre con el fin de replicar una base de datos RAC o de una sola instancia. Para obtener más información, consulte la [*Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard*](#).

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 83
- “Información general sobre la administración de sistemas de archivos de clústeres” en la página 85
- “Administración de grupos de dispositivos” en la página 86
- “Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento” en la página 112
- “Administración de sistemas de archivos de clúster” en la página 117
- “Administración de la supervisión de rutas de disco” en la página 123

Si desea ver una descripción detallada de los procedimientos relacionados con el presente capítulo, consulte la [Tabla 5–2](#).

Para obtener información conceptual relacionada con los dispositivos globales, los espacios de nombres globales, los grupos de dispositivos, la supervisión de las rutas del disco y el sistema de archivos del clúster, consulte la [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 gestor de volumen instalado en el clúster. Solaris Volume Manager “reconoce clústeres” para que pueda agregar, registrar y eliminar grupos de dispositivos mediante el comando `metaset` de Solaris Volume Manager. Para obtener más información, consulte la página del comando `man metaset(1M)`.

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
```

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 133](#).



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	Para obtener 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 109
2. Efectuar la operación de eliminación de DR en el dispositivo que se va a eliminar.	Consulte la documentación que se incluye con el sistema.

Información general sobre la administración de sistemas de archivos de clústeres

Ningún comando de Oracle Solaris Cluster especial es necesario para la administración de sistemas 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 utilizan UFS y también se pueden montar automáticamente durante el inicio. Los sistemas de archivos de clúster sólo son visibles desde el nodo de votación de un clúster global.

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` no se admite en los directorios que no están vacíos. Para obtener más información, consulte la página del comando `man unlink(1M)`.
- 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.

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 setname -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 5-2 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>	“Actualización del espacio de nombre de dispositivos globales” en la página 88
Cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales	“Cómo cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales” en la página 89

TABLA 5-2 Mapa de tareas: administrar grupos de dispositivos (Continuación)

Tarea	Instrucciones
Desplazar un espacio de nombre de dispositivos globales	“Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de <code>lofi</code>” en la página 91 “Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada” en la página 92
Agregar conjuntos de discos de Solaris Volume Manager y registrarlos como grupos de dispositivos mediante el comando <code>metaset</code>	“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 93
Agregar y registrar un grupo de dispositivos de discos básicos con el comando <code>cldevicegroup</code>	“Adición y registro de un grupo de dispositivos (disco básico)” en la página 96
Agregar un grupo de dispositivos con nombre para ZFS con el comando <code>cldevicegroup</code>	“Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 96
Eliminar grupos de dispositivos de Solaris Volume Manager de la configuración con los comandos <code>metaset</code> y <code>metaclear</code>	“Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 98
Eliminar un nodo de todos los grupos de dispositivos con los comandos <code>cldevicegroup</code> , <code>metaset</code> y <code>clsetup</code>	“Eliminación de un nodo de todos los grupos de dispositivos” en la página 98
Eliminar un nodo de un grupo de dispositivos de Solaris Volume Manager con el comando <code>metaset</code>	“Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)” en la página 99
Eliminar un nodo de un grupo de dispositivos de discos básicos con el comando <code>cldevicegroup</code>	“Eliminación de un nodo de un grupo de dispositivos de discos básicos” en la página 101
Modificar las propiedades del grupo de dispositivos con <code>clsetup</code> para generar <code>cldevicegroup</code>	“Cambio de propiedades de los grupos de dispositivos” en la página 103
Mostrar las propiedades y los grupos de dispositivos con el comando <code>cldevicegroup show</code>	“Enumeración de la configuración de grupos de dispositivos” en la página 107
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 105

TABLA 5-2 Mapa de tareas: administrar grupos de dispositivos (Continuación)

Tarea	Instrucciones
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 109
Poner un grupo de dispositivos en estado de mantenimiento con el comando <code>metaset o vx dg</code>	“Colocación de un grupo de dispositivos en estado de mantenimiento” en la página 110

▼ 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` en cada nodo del clúster.**

Este comando puede ejecutarse simultáneamente en todos los nodos del clúster. Para obtener más información, consulte la página del comando `man devfsadm(1M)`.
- 3 **Reconfigure el espacio de nombres.**
`# 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 manera remota en todos los nodos, incluso al ejecutarse en 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 5-1 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 **Evacue los servicios fuera del nodo y reinicie 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 74.](#)
- 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.

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 nombres de dispositivos globales puede coexistir con un sistema de archivos ZFS para el sistema de archivos raíz (`/`) y otros sistemas de archivos raíz, 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 91
- “Migración del espacio de nombre de dispositivos globales desde un dispositivo de `lofi` hasta una partición dedicada” en la página 92

▼ 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.
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 74.](#)
- 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@ID_nodo`.

- 6 Desmonte la partición de dispositivos globales `/global/.devices/node@ID_nodo`.

- 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 Evacue los servicios fuera del nodo y reinicie 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 74](#).
- 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 -n node- 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 Inicie en modo de clúster y complete el espacio de nombres de dispositivos globales.****a. Desde un nodo del clúster, complete los espacios de nombres de dispositivos globales.**

```
# /usr/cluster/bin/cldevice populate
```

b. 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 cldevice 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 o disco sin procesar.

▼ 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 o disco sin procesar) 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 varios 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` Agrega 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 105](#) 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 setname` Especifica el nombre del conjunto de discos, idéntico al del grupo de dispositivos.

`-a` Agrega la unidad al conjunto de discos.

Nota – No utilice el nombre de dispositivo de nivel inferior (`cNtXdY`) cuando agregue 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 5–2 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 `cldevicegroup`.



Precaución – Si crea un grupo de dispositivos en dispositivos replicados, el nombre del grupo de dispositivos que crea (Solaris Volume Manager o disco sin procesar) 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 o disco sin procesar) debe ser idéntico al del grupo de dispositivos replicado.

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 add dsk/d2 dsk/d13
```

2 Cree un grupo de dispositivos con nombre cuyos DID se correspondan con los del grupo de dispositivos eliminado en el [Paso 1](#).

```
# cldevicegroup create -n pnode1,pnode2 -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](#) estableciendo la propiedad `globaldevicepaths` en un grupo de dispositivos de disco sin procesar.

Creó este dispositivo en el [Paso 2](#).

```
# clrs create -t HASStoragePlus -x globaldevicepaths=mypool -g \
migrate_truecopydg-rg hasp2migrate_mypool
```

6 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 `HASStoragePlus` (`hasp-rs`) para el `zpool` creado en el [Paso 3](#) en el grupo de recursos creado en el [Paso 4](#) o el [Paso 6](#).

Configure la propiedad `resource_dependencies` para el recurso `hasp-rs` creado en el [Paso 5](#).

```
# clrs create -g sybase-rg -t HASStoragePlus -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.

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 99](#) para todos los grupos de dispositivos de dicho tipo.
- 4 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
```
- 5 Si alguno de los grupos de dispositivos que figuran en la lista del [Paso 4](#) pertenece a los tipos de grupos `Disk` o `Local_Disk`, siga los pasos descritos en [“Eliminación de un nodo de un grupo de dispositivos de discos básicos” en la página 101](#) para todos estos grupos de dispositivos.
- 6 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 setname -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 setname` Especifica el nombre del grupo de dispositivos.

`-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 5-3 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
```

▼ 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 del 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 5-4 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/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 `preferenced`, el atributo `failback` también se inhabilita automáticamente. Sin embargo, si intenta habilitar o volver a habilitar el atributo `preferenced`, debe elegir entre habilitar o inhabilitar el atributo `failback`.

Si el atributo `preferenced` se habilita o inhabilita, se solicitará que reestablecer el orden de los nodos en la lista de preferencias de propiedades primarias.

En este procedimiento se emplea el comando 5 para establecer o anular la definición de los atributos `preferenced` y `failback` para los grupos de dispositivos de Solaris Volume Manager.

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 para la opción de los grupos de dispositivos y los volúmenes.**
Se muestra el menú Grupos de dispositivos.
- 4 Para modificar las propiedades clave de un grupo de dispositivos, escriba el número para la opción de modificar las propiedades clave de un grupo de dispositivos de Solaris Volume Manager.**
Se muestra el menú Cambiar las propiedades esenciales.
- 5 Para modificar una propiedad de un grupo de dispositivos, escriba el número para la opción de cambiar las preferencias y/o propiedades de recuperación tras los errores.**
Siga las instrucciones para definir las opciones `preferenced` y `failback` para 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 5-5 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* para 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 para 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. Se ejecuta el comando `cldevicegroup` correspondiente, 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:	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 – Entre dichas modificaciones de configuración se incluyen las acciones de agregar o eliminar volúmenes, así como modificar el grupo, el propietario o los permisos de los volúmenes existentes. 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 88](#).

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 5-6 Modificación del número deseado 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
preferred:                  yes
numsecondaries:             1
diskset names:              dg-schost-1
```

Ejemplo 5-7 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
preferred:                  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.**

<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 5–8 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 5–9 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.

`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 5–10 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 de 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 [“Restaure un conjunto de discos dañado” en la página 230](#).

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, inhabilitélo.

```
# 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.

Para 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.

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.

- Para Solaris Volume Manager:

```
# metaset -C release -s diskset
```

5 Ponga en línea el grupo de dispositivos.

```
# cldevicegroup online devicegroup
# cldevicegroup enable devicegroup
```

Ejemplo 5–11 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
```

[Complete all necessary repair procedures.]

[Release ownership.]

For Solaris Volume Manager:

```
# metaset -C release -s dg-schost-1
```

[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 112
- “Visualización del protocolo SCSI de un solo dispositivo de almacenamiento” en la página 113
- “Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento” en la página 114
- “Modificación del protocolo de protección en un solo dispositivo de almacenamiento” en la página 115

▼ 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 del comando `man cluster(1CL)`.

Ejemplo 5–12 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
clusterid:                                    0x4FES2C888
installmode:                                disabled
heartbeat_timeout:                          10000
heartbeat_quantum:                          1000
private_netaddr:                             172.16.0.0
private_netmask:                             255.255.111.0
max_nodes:                                   64
max_privatenets:                             10
udp_session_timeout:                         480
concentrate_load:                            False
global_fencing:                              prefer3
Node List:                                    phys-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 del comando `man cldevice(1CL)`.

Ejemplo 5–13 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/rdisk/d3
Full Device Path:	phappy1:/dev/rdisk/c4t8d0
Full Device Path:	phappy2:/dev/rdisk/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 115](#) 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 proporcione la autorización de RBAC `solaris.cluster.modify`.**
- 2 **Establezca el protocolo de protección para todos los dispositivos de almacenamiento que no sean de quórum.**

<code>cluster set -p global_fencing={pathcount prefer3 nofencing nofencing-noscrub}</code>	
<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 <code>nofencing-noscrub</code> sólo con dispositivos de almacenamiento con problemas graves con las reservas de SCSI.

Ejemplo 5–14 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 proporcione la autorización de 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.

Para obtener más información, consulte la página del comando `man cluster(1CL)`.

Ejemplo 5–15 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 5–3 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	“Adición de un sistema de archivos de clúster” en la página 117
Eliminar un sistema de archivos de clúster	“Eliminación de un sistema de archivos de clúster” en la página 121
Comprobar los puntos de montaje globales de un clúster para verificar la coherencia entre los nodos	“Comprobación de montajes globales en un clúster” en la página 123

▼ 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.
- Un grupo de dispositivos (como un grupo de dispositivos de Solaris Volume Manager) 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.

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 UFS con el comando `newfs`.



Precaución – 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.

`phys - schost# newfs raw-disk-device`

La tabla siguiente muestra ejemplos de nombres para el argumento *dispositivo_discos_básicos*. 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	/dev/md/nfs/rdisk/d1	Dispositivo de disco básico d1 dentro del conjunto de discos nfs
Ninguno	/dev/global/rdisk/d1s3	Dispositivo de disco básico d1s3

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/mount-point/
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 del comando `man vfstab(4)` para obtener información detallada.

- Especifique en cada entrada las opciones de montaje requeridas para el tipo de sistema de archivos que utilice.
- Para montar de forma automática el sistema de archivos de clúster, establezca el campo `mount at boot` en `yes`.
- Para cada sistema de archivos del clúster, asegúrese de que la información de la entrada `/etc/vfstab` sea idéntica en todos los nodos.
- Compruebe que las entradas del archivo `/etc/vfstab` de cada nodo muestren los dispositivos en el mismo orden.
- Compruebe las dependencias de orden de inicio de los sistemas de archivos.

Por ejemplo, considere la siguiente situación: `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 iniciar y montar /global/oracle/logs/ cuando phys-schost-1 inicie 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, no se devuelve nada.

Para obtener más información, consulte la página de comando man [cluster\(1CL\)](#).

6 Monte el sistema de archivos del clúster desde cualquier nodo del clúster.

```
phys-schost# mount /global/device-group/mountpoint/
```

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\(1M\)](#).

Ejemplo 5–16 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 la configuración se haya efectuado correctamente, el sistema de archivos del clúster se monta 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
#
/global/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 5–17 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/dsk/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*](#) para obtener información conceptual sobre el daemon de supervisión de las rutas del disco. Consulte la página del comando `man cldevice(1CL)` para obtener una descripción de las opciones de comandos y los comandos relacionados. Para obtener más información sobre el ajuste del daemon `scdpmd`, consulte la página del comando `man scdpmd.conf(4)`. También consulte la página del comando `man syslogd(1M)` para ver los errores registrados que el daemon informa.

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 5-4 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 124
Anular la supervisión de una ruta de disco	“Anulación de la supervisión de una ruta de disco” en la página 126
Imprimir el estado de rutas de disco erróneas correspondientes a un nodo	“Impresión de rutas de disco erróneas” en la página 126
Supervisar rutas de disco desde un archivo	“Supervisión de rutas de disco desde un archivo” en la página 127
Habilitar o inhabilitar el re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	“Habilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 129 “Inhabilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 130
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 127

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.

`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 5–18 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 5–19 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 5–20 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
```

/dev/did/rdisk/d7	schost-2	Ok
/dev/did/rdisk/d8	schost-2	Ok
	schost-2	Ok

▼ Anulación de la supervisión de una ruta de disco

Siga este procedimiento para anular la supervisión de una ruta 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 5-21 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.

- 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 5–22 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.

`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 5–23 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 rearranque 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 rearranque 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* rearranca 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 rearranque 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. Para 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 131
- “Administración de servidores de quórum de Oracle Solaris Cluster” en la página 153

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 utilizar el comando `clquorum` para realizar todos los procedimientos administrativos de dispositivos de quórum. Además, puede llevar a cabo algunos procedimientos mediante la utilidad interactiva `clsetup`. Siempre que es posible, los procedimientos de quórum de esta sección se describen con la utilidad `clsetup`. Para obtener más información, consulte las páginas del comando `man clquorum(1CL)` y `clsetup(1CL)`.

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	Para obtener instrucciones
Agregar un dispositivo del quórum a un clúster mediante la utilidad <code>clsetup</code>	“Adición de un dispositivo de quórum” en la página 134
Eliminar un dispositivo del quórum de un clúster mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	“Eliminación de un dispositivo de quórum” en la página 142
Eliminar el último dispositivo del quórum de un clúster mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	“Eliminación del último dispositivo de quórum de un clúster” en la página 143
Reemplazar un dispositivo de quórum de un clúster mediante los procedimientos de agregar y quitar	“Sustitución de un dispositivo de quórum” en la página 145
Modificar una lista de dispositivos de quórum mediante los procedimientos de agregar y quitar	“Modificación de una lista de nodos de dispositivo de quórum” en la página 146

TABLA 6-1 Lista de tareas: administrar el quórum (Continuación)

Tarea	Para obtener instrucciones
Poner un dispositivo del quórum en estado de mantenimiento mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 148
Mientras se encuentra en estado de mantenimiento, el dispositivo de quórum no participa en las votaciones para establecer el quórum.	
Restablecer la configuración de quórum a su estado predeterminado mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 149
Enumerar los dispositivos del quórum y los recuentos de votos mediante el comando <code>clquorum</code>	“Enumeración de una lista con la configuración de quórum” en la página 151

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	Para obtener instrucciones
1. Habilitar un nuevo dispositivo de quórum para sustituir el que se va a eliminar	“Adición de un dispositivo de quórum” en la página 134

TABLA 6-2 Mapa de tareas: reconfiguración dinámica con dispositivos de quórum (Continuación)

Tarea	Para obtener instrucciones
2. Inhabilitar el dispositivo de quórum que se va a eliminar	“Eliminación de un dispositivo de quórum” en la página 142
3. Efectuar la operación de eliminación de reconfiguración dinámica en el dispositivo que se va a eliminar	

Adición de un dispositivo de quórum

En esta sección se indican los procedimientos para agregar un dispositivo del quórum. Compruebe que todos los nodos del clúster estén en línea antes de agregar un nuevo dispositivo de quórum. Para obtener información sobre la determinación del número de votos del quórum necesarios para el clúster, las configuraciones del quórum recomendadas y el aislamiento 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:

- LUN compartidos desde:
 - Disco de SCSI compartido
 - Almacenamiento SATA (Serial Attached Technology Attachment)
 - Sun ZFS Storage Appliance de Oracle
- 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 135](#)
- [“Adición de un dispositivo de quórum de servidor de quórum” en la página 138](#)

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 del quórum, puede utilizar un disco que contenga datos del 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`.

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 del 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 1.**

```
# clsetup
```

Aparece el menú principal de `clsetup`.

- 3 Escriba el número para la opción de quórum.**

Aparece el menú Quórum.

- 4 Escriba el número para la opción de agregar un dispositivo del quórum, luego escriba `yes` cuando la utilidad `clsetup` solicite que confirme el dispositivo del quórum que va a agregar.**

La utilidad `clsetup` le pregunta qué tipo de dispositivo del quórum desea agregar.

- 5 Escriba el número para la opción de un dispositivo del 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 del quórum NAS de Sun ZFS Storage Appliance

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 Consulte la documentación de instalación que se envía con Sun ZFS Storage Appliance o la Ayuda en línea de la aplicación para obtener instrucciones sobre cómo configurar un dispositivo iSCSI.
- 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 del 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 página del comando `man clquorum(1CL)` para obtener más información.

▼ Adición de un dispositivo de quórum de servidor de quórum

Antes de empezar

Antes de que pueda agregar un servidor de quórum de Oracle Solaris Cluster como un dispositivo del quórum, el software de servidor de quórum de Oracle Solaris Cluster debe estar instalado en el equipo y el servidor de quórum se debe haber iniciado y puesto en ejecución. Si desea obtener información sobre cómo instalar el servidor de quórum, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software 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 Asegúrese de 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 ralentiza dicha comunicación se ralentiza 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 para la opción de quórum.

Aparece el menú Quórum.

5 Escriba el número para la opción de agregar un dispositivo del 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 para la opción de un dispositivo del 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 del 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 proporcione el nombre del host del servidor 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 6-2 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 142
- “Eliminación del último dispositivo de quórum de un clúster” en la página 143
- “Sustitución de un dispositivo de quórum” en la página 145

▼ Eliminación de un dispositivo de quórum

Cuando un dispositivo del quórum se elimina, ya no participa en la votación para establecer quórum. Todos los clústeres 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 143.

`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`.**
`# clsetup`
Aparece el menú principal.
- 4 **Escriba el número para la opción de quórum.**

- 5 **Escriba el número para la opción de eliminar un dispositivo del 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 6-3 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 clústeres del servidor de quórum” en la página 157](#).

▼ 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 142](#).

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` le permite eliminar el dispositivo de la configuración. Para obtener más información, consulte la página del comando `man clquorum(1CL)`. 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 134](#).

`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 del quórum ha fallado, use la opción de 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 quórum`. Las opciones del menú de administración del clúster `clsetup` no están disponibles cuando el clúster está en modo de instalación. Consulte la página del comando `man “Cómo poner un nodo en estado de mantenimiento” en la página 204` and the `clsetup(1CL)` para obtener más información.

- 3 **Compruebe que se haya quitado el dispositivo de quórum.**

```
# clquorum list -v
```

Ejemplo 6–4 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 134](#) 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 142](#) 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 manual de su hardware para el contenedor de discos. Consulte también el [Oracle Solaris Cluster 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 146](#)
- [“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 148](#)
- [“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 149](#)
- [“Enumeración de una lista con la configuración de quórum” en la página 151](#)

- “Reparación de un dispositivo de quórum” en la página 152
- “Cambio del tiempo de espera predeterminado del quórum” en la página 153

▼ **Modificación de una lista de nodos de dispositivo de quórum**

Puede emplear la utilidad `clsetup` para agregar un nodo o para quitar un nodo de la lista de nodos de un dispositivo del quórum existente. 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. Cuando se agrega un dispositivo del quórum, el comando `clquorum` automáticamente configura las rutas de nodo a disco para todos los nodos conectados al disco. Para obtener más información, consulte la página del comando `man clquorum(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 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 para la opción de quórum.**
Aparece el menú Quórum.
- 5 Escriba el número para la opción de eliminar un dispositivo del 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 para la opción de agregar un dispositivo del 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 6-5 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

Utilice el comando `clquorum` para poner un dispositivo del quórum en estado de mantenimiento. Para obtener más información, consulte la página del comando [man clquorum\(1CL\)](#). La utilidad `clsetup` no tiene actualmente esta capacidad.

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 clústeres 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 “[Cómo poner un nodo en estado de mantenimiento](#)” en la [página 204](#).

`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 6–6 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 149.

Para poner un nodo en estado de mantenimiento, consulte [“Cómo poner un nodo en estado de mantenimiento”](#) en la página 204.

▼ 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 “[Cómo sacar un nodo del estado de mantenimiento](#)” en la página 206.
- Para obtener más información sobre el número de votos del 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 6–7 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/rdisk/d20s2
Type:                    shared_disk
Access Mode:             scsi3
Hosts (enabled):         phys-schost-2, phys-schost-3
```

▼ Enumeración de una lista con la configuración de quórum

Para enumerar la configuración de quórum no hace falta ser superusuario. Puede asumir cualquier rol 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 clústeres, 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.

● Utilice el comando `clquorum` para mostrar la configuración de quórum.

```
% clquorum show +
```

Ejemplo 6-8 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/rdisk/d20s2
Type:                     shared_disk
Access Mode:              scsi3
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 134](#) para agregar un nuevo dispositivo de quórum.

Consulte [“Eliminación de un dispositivo de quórum” en la página 142](#) 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 para el contenedor de discos en la guía del hardware. Consulte también el [Oracle Solaris Cluster Hardware Administration Manual](#).

3 Agregue el disco sustituido como nuevo dispositivo de quórum.

Consulte “[Adición de un dispositivo de quórum](#)” en la página 134 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 142 para quitar el dispositivo de quórum.

Cambio del tiempo de espera predeterminado del quórum

Existe un tiempo de espera predeterminado de 25 segundos para la finalización de operaciones del quórum durante una reconfiguración del clúster. Puede aumentar el tiempo de espera del quórum a un valor superior siguiendo las instrucciones en “[Cómo configurar dispositivos del quórum](#)” de *Guía de instalación del software de Oracle Solaris Cluster*. En lugar de aumentar el valor de tiempo de espera, también se puede cambiar a otro dispositivo del quórum.

Obtendrá información adicional sobre cómo solucionar problemas en “[Cómo configurar dispositivos del quórum](#)” de *Guía de instalación del software de Oracle Solaris Cluster*.

Nota – En el caso de Oracle Real Application Clusters (Oracle RAC), no cambie el tiempo de espera predeterminado del quórum de 25 segundos. En determinados casos en que una parte de la partición del clúster cree que la otra está inactiva ("cerebro dividido"), un tiempo de espera superior puede hacer que falle el proceso de migración tras error de Oracle RAC VIP debido a la finalización del tiempo de espera de recursos VIP. Si el dispositivo del quórum que se utiliza no es adecuado para un tiempo de espera predeterminado de 25 segundos, utilice otro dispositivo.

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:

- “Inicio y detención del software del servidor del quórum” en la página 154
- “Inicio de un servidor de quórum” en la página 154
- “Detención de un servidor de quórum” en la página 155
- “Visualización de información sobre el servidor de quórum” en la página 156
- “Limpieza de la información caducada sobre clústeres del servidor de quórum” en la página 157

Si desea obtener información sobre cómo instalar y configurar servidores de quórum de Oracle Solaris Cluster, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster](#).

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 instalar el software del servidor de quórum, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster](#). Para obtener información sobre cómo cambiar el valor del tiempo de espera del quórum, consulte [“Cambio del tiempo de espera predeterminado del quórum” en la página 153](#).

▼ 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 del puerto de escucha del servidor de quórum. 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 6–9 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 6–10 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 del puerto de escucha del servidor de quórum. 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 6–11 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 6–12 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 clústeres 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`. Para obtener más información sobre los perfiles de derechos de RBAC, consulte la página del 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 6–13 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 clústeres 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 6–14 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 6–15 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 clústeres 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 142](#). 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 clústeres 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 142](#).



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 `cluster show` en un nodo del clúster.

`servidor_quórum` 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, `/etc/scqsd/scqsd.conf`.

-y Fuerce el comando `clquorumserver clear` 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 clústeres 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 6-16 Limpieza de la información obsoleta de clústeres 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 clústeres 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 clústeres 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 clústeres 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 11 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 176](#).

En este capítulo hay información y procedimientos para los temas siguientes.

- [“Administración de interconexiones de clústeres” en la página 161](#)
- [“Administración de redes públicas” en la página 176](#)

Si desea ver una descripción detallada de los procedimientos relacionados con este capítulo, consulte la [Tabla 7-1](#) y la [Tabla 7-3](#).

Consulte la [Oracle Solaris Cluster Concepts Guide](#) para obtener información general sobre las redes públicas y las interconexiones del clúster.

Administración de interconexiones de clústeres

En esta sección se proporcionan procedimientos para reconfigurar interconexiones de clúster, como adaptador de transporte del clúster y cable de transporte de clúster. Estos procedimientos requieren que instale el software de Oracle Solaris Cluster.

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 del 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 de Oracle Solaris Cluster*. Para conocer los procedimientos sobre las tareas de mantenimiento de los componentes de hardware de los clústeres, consulte el *Oracle Solaris Cluster 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 23
Comprobar el estado de la interconexión de los clústeres con <code>clinterconnect status</code>	“Comprobación del estado de la interconexión de clúster” en la página 163
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 164
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 167
Habilitar un cable de transporte de clúster mediante <code>clsetup</code>	“Habilitación de un cable de transporte de clúster” en la página 169
Inhabilitar un cable de transporte de clúster mediante <code>clsetup</code>	“Inhabilitación de un cable de transporte de clúster” en la página 171
Determinar el número de instancia de un adaptador de transporte	“Determinación del número de instancia de un adaptador de transporte” en la página 172
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 173

Reconfiguración dinámica con interconexiones de clústeres

Al completar operaciones de reconfiguración dinámica en las interconexiones de clústeres 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 comunique 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	“Reconfiguración dinámica con interfaces de red pública” en la página 177
2. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	

▼ Comprobación del estado de la interconexión 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.

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:net0 phys-schost-2:net0 Path online
Transport path: phys-schost-1:net4 phys-schost-2:net4 Path online
Transport path: phys-schost-1:net0 phys-schost-3:net0 Path online
Transport path: phys-schost-1:net4 phys-schost-3:net4 Path online
Transport path: phys-schost-2:net0 phys-schost-3:net0 Path online
Transport path: phys-schost-2:net4 phys-schost-3:net4 Path online
```

▼ Adición de dispositivos de cable de transporte de clúster, adaptadores o conmutadores de transporte

Para obtener información sobre los requisitos para el transporte privado del clúster, consulte [“Interconnect Requirements and Restrictions” de Oracle Solaris Cluster Hardware Administration 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 se instale la parte física de cable de transporte de clústers.

Para conocer el procedimiento de instalación de un cable de transporte de clúster, consulte el *Oracle Solaris Cluster 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 para la opción de mostrar el menú de interconexión de clúster.

5 Escriba el número para la opción de agregar un cable de transporte.

Siga las instrucciones y escriba la información que se solicite.

6 Escriba el número para 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
nge	set nge:nge_taskq_disable=1
e1000g	set e1000g:e1000g_taskq_disable=1

7 Escriba el número para 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    net5
    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:net5,hub2
===Transport Cables ===
Transport Cable:          phys-schost-1:net5@0,hub2
Endpoint1:                phys-schost-2:net4@0
Endpoint2:                hub2@2
State:                    Enabled

# clinterconnect show phys-schost-1:net5
=== Transport Adepters for net5
Transport Adapter:          net5
Adapter State:              Enabled
Adapter Transport Type:     dlpi
Adapter Property (device_name): net6
Adapter Property (device_instance): 0
Adapter Property (lazy_free): 1
Adapter Property (dlpi_heartbeat_timeout): 10000
Adpater Property (dlpi_heartbeat_quantum): 1000
Adapter Property (nw_bandwidth): 80
Adapter Property (bandwidth): 70
Adapter Property (ip_address): 172.16.0.129
Adapter Property (netmask): 255.255.255.128
Adapter Port Names:        0
Adapter Port 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
```

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 163.](#)

▼ Eliminación de cable de transporte de clúster, adaptadores de transporte y conmutadores de transporte

Utilice el siguiente procedimiento para quitar cables de transporte de clúster, adaptadores de transporte y conmutadores de transporte de una configuración de 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 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 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 para la opción de acceder al menú de interconexión de clúster.**5 Escriba el número para la opción de deshabilitar 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 para 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 para la opción de quitar el adaptador de transporte de un nodo.

Siga las instrucciones y escriba la información que se solicite. Debe saber cuáles son los nombres válidos de nodos, adaptadores y conmutadores.

Si va a quitar un adaptador físico de un nodo, consulte el [Oracle Solaris Cluster Hardware Administration Manual](#) para conocer los procedimientos de las tareas de mantenimiento de hardware.

8 Escriba el número para 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    net0
    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:net5,hub2@0
===Transport Cables===
Transport Cable:                phys-schost-1:net5,hub2@0
Endpoint1:                     phys-schost-1:net5
Endpoint2:                     hub2@0
State:                         Enabled

# clinterconnect show phys-schost-1:net5
=== Transport Adepters for net5
Transport Adapter:              net5
Adapter State:                 Enabled
Adapter Transport Type:        dlpi
Adapter Property (device_name): net6
Adapter Property (device_instance): 0
Adapter Property (lazy_free): 1
Adapter Property (dlpi_heartbeat_timeout): 10000
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 hub2
=== Transport Switches ===
Transport Switch:              hub2
State:                        Enabled
Type:                         switch
Port Names:                   1 2
Port State(1):                Enabled
Port State(2):                Enabled

```

▼ Habilitación de un cable de transporte de clúster

Esta opción se utiliza para habilitar un cable de trasporte de clúster ya existente.

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 para la opción de acceder al menú de interconexión de clúster y presione la tecla Intro.

4 Escriba el número para la opción de habilitar el cable de transporte y presione 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 net0, 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	net0
switch names	hub1

```
[Verify that the scinterconnect
    command was completed successfully:]
```

```
clinterconnect enable phys-schost-2:net0
```

Command completed successfully.

```
[Quit the clsetup Cluster Interconnect Menu and Main Menu.]
```

```
[   Verify that the cable is enabled:]
```

```
# clinterconnect show phys-schost-1:net5,hub2
Transport cable:  phys-schost-2:net0@0 ethernet-1@2    Enabled
Transport cable:  phys-schost-3:net5@1 ethernet-1@3    Enabled
Transport cable:  phys-schost-1:net5@0 ethernet-1@1    Enabled
```

▼ Inhabilitación de un cable de transporte de clúster

Es posible que necesite deshabilitar un cable de transporte de clúster para cerrar de manera temporal 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 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 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 para la opción de acceder al menú de interconexión de clúster y presione la tecla Intro.**
- 5 **Escriba el número para la opción de deshabilitar el cable de transporte y presione 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 deshabilitar un cable de transporte de clúster en el adaptador `net0`, 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         net0
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:net5,hub2
Transport cable:  phys-schost-2:net0@0 ethernet-1@2    Disabled
Transport cable:  phys-schost-3:net5@1 ethernet-1@3    Enabled
Transport cable:  phys-schost-1:net5@0 ethernet-1@1    Enabled
```

▼ 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
      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/pci11c8,0@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 Asegúrese de que el acceso de shell remoto ([rsh\(1M\)](#)) o de shell seguro ([ssh\(1\)](#)) para superusuario esté habilitado para todos los nodos del clúster.

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 de los nodos de votación en el nodo especificado al siguiente nodo de votación preferido.

```
# 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 Seleccione la opción de menú Cambiar intervalos y asignación de direcciones de red para el transporte del clúster.

La utilidad `clsetup` muestra la configuración de red privada actual y, a continuación, pregunta si se desea modificar esta configuración.

4 Para cambiar la dirección IP de red privada o el rango de direcciones IP, escriba *yes* y presione la tecla de retorno.

La utilidad `clsetup` muestra la dirección IP de red privada predeterminada, 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 presione la tecla de retorno.
- Para cambiar la dirección IP de red privada predeterminada:
 - 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 un máximo de 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**.

- **Para cambiar el rango de direcciones IP:**
 - 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. **Indique 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 `clsetup` sobre si desea continuar con la actualización.**
- 8 Cuando haya finalizado, salga de la utilidad `clsetup`.**
- 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. **Inicie 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.
- 10 Compruebe que el nodo haya arrancado sin errores y esté en línea.**
`# cluster status -t node`

Administración de redes públicas

El software de Oracle Solaris Cluster admite la implementación de software de Oracle Solaris de múltiples rutas de redes IP (IPMP) para 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 11 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.

- No desconfigure (desconecte) o apague un adaptador de un grupo de múltiples rutas de red IP sin primero conmutar mediante direcciones IP del adaptador a eliminar a un adaptador alternativo en el grupo, mediante el comando `if_mpadm`. Para obtener más información, consulte la página del comando `man 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 11	Capítulo 15, “Administering IPMP” de <i>Oracle Solaris Administration: Network Interfaces and Network Virtualization</i>

Si desea información sobre los procedimientos de instalación del software de clúster, consulte *Guía de instalación del software de Oracle Solaris Cluster*. Para conocer los procedimientos sobre las tareas de mantenimiento de los componentes de hardware de las redes públicas, consulte el *Oracle Solaris Cluster 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 quitar una interfaz de red pública, conmute las direcciones IP del adaptador que se va a quitar a otro adaptador del grupo de múltiples rutas mediante el comando `if_mpadm`. Para obtener más información, consulte la página del comando `man 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.

Aplique 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 quitar a otro adaptador del grupo de múltiples rutas mediante el comando <code>if_mpadm</code>	Página del comando <code>man if_mpadm(1M)</code> . “How to Move an Interface From One IPMP Group to Another Group” de Oracle Solaris Administration: Network Interfaces and Network Virtualization
2. Quitar el adaptador del grupo de múltiples rutas mediante el comando <code>ipadm</code>	Página del comando <code>man ipadm(1M)</code> “How to Add or Remove IP Addresses” de Oracle Solaris Administration: Network Interfaces and Network Virtualization
3. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	

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 179
- “Eliminación de nodos de un clúster” en la página 182

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.

Especificar una dirección IP y un NIC para cada nodo de clúster de zona es opcional.

Nota – Si no configura una dirección IP para cada nodo de clúster de zona, ocurrirán dos cosas:

1. Ese clúster de zona específico no podrá configurar dispositivos NAS para utilizar en el clúster de zona. El clúster utiliza la dirección IP del nodo de clúster de zona para comunicarse con el dispositivo NAS, por lo que no tener una dirección IP impide la admisión de clústeres para el aislamiento de dispositivos NAS.
 2. El software del clúster activará cualquier dirección IP de host lógico en cualquier NIC.
-

Si el nodo de clúster de zona original no tiene una dirección IP o un NIC especificados, no tiene que especificar esta información para el nuevo nodo de clúster de zona.

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 Hardware Administration Manual</i>
Agregar almacenamiento compartido	<i>Oracle Solaris Cluster Hardware Administration Manual</i>
Agregar el nodo a la lista de nodos autorizados	<code>/usr/cluster/bin/claccess allow -h node-being-added</code>
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 de Oracle Solaris Cluster</i>
Agregar el nuevo nodo a un clúster existente	“Adición de nodos a un clúster” en la página 179
Si el clúster se configura en asociación con Oracle Solaris Cluster Geographic Edition, configure el nuevo nodo como participante activo en la configuración.	“How to Add a New Node to a Cluster in a Partnership” de <i>Oracle Solaris Cluster Geographic Edition System Administration Guide</i>

▼ Cómo agregar un nodo a un clúster existente

Antes de agregar una máquina virtual o un host de Oracle Solaris a un clúster global o un clúster de zona existentes, asegúrese de que el nodo tenga todo el hardware necesario instalado y configurado correctamente, incluida una conexión física operacional a la interconexión privada del clúster.

Para obtener información sobre la instalación de hardware, consulte el *Oracle Solaris Cluster 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 Instale y configure el software en el nuevo nodo del clúster.

Use la utilidad `scinstall` para completar la instalación y la configuración del nodo nuevo, como se describe en la [Guía de instalación del software de Oracle Solaris Cluster](#).

- 4 Use la utilidad `scinstall` en el nodo nuevo para configurar ese nodo en el clúster.

- 5 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, `sczone` es el nombre de zona y `sc_ipmp0` es el nombre de grupo IPMP.

```
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-3
clzc:sczone:node>set hostname=hostname3
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname3
clzc:sczone:node:net>set physical=sc_ipmp0
clzc:sczone:node:net>end
clzc:sczone:node>end
clzc:sczone>exit
```

Para obtener instrucciones detalladas sobre cómo configurar el nodo, consulte “Configuración de un clúster de zona” de [Guía de instalación del software de Oracle Solaris Cluster](#).

- 6 Después de configurar el nodo, reinicie el nodo en modo de clúster e instale el clúster de zona en el nodo.

```
# clzc install zone-cluster
```

- 7 Para evitar que se agreguen máquinas nuevas al clúster, desde la utilidad `clsetup`, escriba el número correspondiente a la opción que indica al clúster que omita las solicitudes para agregar máquinas nuevas.

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`.

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](#).

Eliminación de nodos de un clúster

En esta sección, se proporcionan instrucciones para 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 la eliminación de un nodo de una configuración de RAC, consulte “[Cómo eliminar Soporte para Oracle RAC de los nodos seleccionados](#)” de *Servicio de datos de Oracle para la Guía de clústeres de aplicación real de Oracle*. Tras completar este proceso, elimine un nodo de una configuración de RAC y siga los pasos adecuados que se muestran a continuación.

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. Si hay un clúster de zona, iniciar sesión en el clúster de zona y evacuar el nodo de clúster de zona que se encuentra en el nodo físico que se está desinstalando. Luego, eliminar el nodo del clúster de zona antes de desactivar el nodo físico.	<code>clnode evacuate nodo</code> “Eliminación de un nodo de un clúster de zona” en la página 184
Verificar que se pueda eliminar el nodo comprobando los hosts permitidos.	<code>claccess show</code> <code>claccess allow -h nodo_que_eliminar</code>
Si el nodo no se puede eliminar, conceder al nodo acceso a la configuración del clúster.	
Eliminar el nodo de todos los grupos de dispositivos.	“Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)” en la página 99
Eliminar todos los dispositivos de quórum conectados al nodo que se va a eliminar.	Este paso es opcional en caso de que se elimine un nodo de un clúster compuesto por dos nodos. “Eliminación de un dispositivo de quórum” en la página 142 Aunque hay que quitar el dispositivo de quórum antes de eliminar el dispositivo de almacenamiento en el paso siguiente, inmediatamente después se puede volver agregar el dispositivo de quórum. “Eliminación del último dispositivo de quórum de un clúster” en la página 143
Poner el nodo que se va a eliminar en un modo que no sea de clúster.	“Cómo poner un nodo en estado de mantenimiento” en la página 204
Eliminar un nodo de un clúster de zona.	“Eliminación de un nodo de un clúster de zona” en la página 184
Eliminar un nodo de la configuración de software del clúster.	“Eliminación de un nodo de la configuración de software del clúster” en la página 184
(Opcional) Desinstalar el software de Oracle Solaris Cluster de un nodo de clúster.	“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208

▼ Eliminación de un nodo de un clúster de zona

Puede eliminar un nodo de un clúster de zona deteniendo el nodo, desinstalando el nodo y eliminando el nodo 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 **Elimine el nodo de todos los grupos de recursos en el clúster de zona.**

```
phys-schost# clrg remove-node -n zonehostname -Z zoneclustername rg-name
```

- 4 **Desinstale el nodo del clúster de zona.**

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

- 5 **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
```

```
clzc:sczone> exit
```

- 6 **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 184](#) 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.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System Interactively” de *Booting and Shutting Down Oracle Solaris on x86 Platforms*](#).

- 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.

- 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/kernel/#ISADIR/unix -B $ZFS-BOOTFS -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.

- e. Escriba `b` para iniciar el nodo en el modo sin clúster.

Este cambio en el comando del parámetro de inicio del núcleo no se conserva tras el inicio del sistema. La siguiente vez que reinicie el nodo, se iniciará en el 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 Elimine el nodo del clúster.

Ejecute el siguiente comando desde un nodo activo:

```
phys-schost# clnode clear -F nodename
```

Si cuenta con grupos de recursos que son `rg_system=true`, debe cambiarlos a `rg_system=false` para que el comando `clnode clear -F` se ejecute con éxito. Después de ejecutar `clnode clear -F`, restablezca los grupos de recursos a `rg_system=true`.

Ejecute el siguiente comando desde el nodo que desea eliminar:

```
phys-schost# clnode remove -F
```

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 [“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208](#). También puede eliminar el nodo del clúster y desinstalar el software de Oracle Solaris Cluster al mismo tiempo. Cambie a un directorio que no contenga ningún archivo de Oracle Solaris Cluster y escriba `scinstall -r`.
- Si no tiene previsto desinstalar el software Oracle Solaris Cluster del nodo eliminado, puede eliminar físicamente el nodo del clúster mediante la eliminación de las conexiones de hardware como se describe en el [Oracle Solaris Cluster 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 [“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208.](#)

Para obtener información sobre los procedimientos de hardware, consulte el [Oracle Solaris Cluster 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 existente, consulte [“Cómo agregar un nodo a un clúster existente” en la página 180.](#)

▼ 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.



Precaución – 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 del nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos de los nodos de votación en el nodo especificado al siguiente nodo de votación preferido.

4 Establezca los grupos de dispositivos en estado de mantenimiento.

Para el procedimiento de establecer un grupo de dispositivos en el estado de mantenimiento, consulte “[Cómo poner un nodo en estado de mantenimiento](#)” en la página 204.

5 Elimine el nodo de los grupos de dispositivos.

Si usa un disco sin procesar, 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 -n node + | resourcegroup  
nodo
```

El nombre del nodo.

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# pkg uninstall /ha-cluster/library/ucmm
```



Precaución – 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.

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.

Si desea obtener información sobre cómo poner en línea un grupo de dispositivos, consulte [“Cómo sacar un nodo del estado de mantenimiento” en la página 206.](#)

▼ 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 98.](#)
 - b. Tras eliminar el nodo de todos los grupos de dispositivos, vuelva a [“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208](#) 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 [“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208](#) 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 191
- “Cómo realizar tareas administrativas del clúster de zona” en la página 222
- “Resolución de problemas” en la página 228

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	Capítulo 8, “Adición y eliminación de un nodo”
Cambiar el nombre del clúster	“Cambio del nombre del clúster” en la página 192
Enumerar los ID de nodos y sus nombres de nodo correspondientes	“Asignación de un ID de nodo a un nombre de nodo” en la página 194
Permitir o denegar que se agreguen nodos nuevos al clúster	“Uso de la autenticación del nodo del clúster nuevo” en la página 194
Cambiar el tiempo para un clúster utilizando el NTP	“Restablecimiento de la hora del día en un clúster” en la página 196
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	“SPARC: Visualización de la PROM OpenBoot en un nodo” en la página 198
Agregar o cambiar el nombre de host privado	“Cómo cambiar un nombre de host privado de nodo” en la página 199
Poner un nodo de clúster en estado de mantenimiento	“Cómo poner un nodo en estado de mantenimiento” en la página 204
Cambiar el nombre de un nodo	“Cómo cambiar el nombre de un nodo” en la página 202
Sacar un nodo del clúster fuera del estado de mantenimiento	“Cómo sacar un nodo del estado de mantenimiento” en la página 206
Desinstalar el software del clúster de un nodo del clúster	“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208
Agregar y administrar una MIB de eventos de SNMP	“Cómo habilitar una MIB de eventos de SNMP” en la página 213 “Cómo agregar un usuario de SNMP a un nodo” en la página 216
Configurar los límites de carga de un nodo	“Cómo configurar límites de carga en un nodo” en la página 219
Mover un clúster de zona, preparar un clúster de zona para aplicaciones, eliminar un clúster de zona	“Cómo realizar tareas administrativas del clúster de zona” en la página 222

▼ 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 8.**
`phys-schost# clsetup`
 Aparece el menú principal.
- 3 Para cambiar el nombre del clúster, escriba el número para 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

En el siguiente ejemplo se muestra el comando `cluster` generado desde la utilidad `clsetup` para cambiar al nuevo nombre de clúster, `dromedary`.

```
phys-schost# cluster rename -c dromedary
```

Para obtener más información, consulte las páginas del comando `man cluster(1CL)` y `clsetup(1CL)`.

▼ Asignación de un ID de nodo a un nombre de nodo

Durante la instalación de Oracle Solaris Cluster, a cada nodo se le asigna automáticamente un número de ID de nodo único. 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 **Utilice el comando `clnode show` para mostrar la información de configuración del clúster para el clúster global.**

```
phys-schost# clnode show | grep Node
```

Para obtener más información, consulte la página del comando `man clnode(1CL)`.

- 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 sczone 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`.**
`phys-schost# clsetup`
Aparece el menú principal.
- 3 Para trabajar con la autenticación del clúster, escriba el número para la opción para nuevos nodos.**
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 de Oracle Solaris Cluster usa 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 la *Oracle Solaris Cluster Concepts Guide* y la *Guía del usuario del protocolo de hora de red* en <http://download.oracle.com/docs/cd/E19065-01/servers.10k/>.



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`, `rdate` o `svcadm` interactivamente o dentro de las secuencias de comandos `cron`. Para obtener más información, consulte las páginas del comando `man date(1)`, `rdate(1M)`, `svcadm(1M)` o `cron(1M)`. La página del comando `man ntpd(1M)` se entrega con el paquete de Oracle Solaris 11 `service/network/ntp`.

`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.

Se muestra el menú de GRUB.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System Interactively” de Booting and Shutting Down Oracle Solaris on x86 Platforms.](#)

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.

Se muestra la pantalla de parámetros de inicio de GRUB.

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/kernel/$ISADIR/unix _B $ZFS-BOOTFS -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.

e. Escriba b para iniciar el nodo en el modo sin clúster.

Nota – Este cambio en el comando del parámetro de inicio del núcleo no se conserva tras el inicio del sistema. La siguiente vez que reinicie el nodo, se iniciará en el 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 `rddate(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(tm).

`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. Este comando también conmuta todos los grupos de recursos del nodo de votación especificado del clúster global al nodo de votación que se sitúe a continuación en el orden de preferencia.

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



Precaución – No use el comando `send brk` en la consola de un clúster para cerrar un nodo de clúster.

3 Ejecute los comandos OBP.

▼ Cómo cambiar un nombre de host privado de nodo

Utilice este procedimiento para cambiar un nombre de host privado de un nodo de 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. El software de clúster los asigna.

`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 Deshabilite, en todos los nodos del clúster, cualquier recurso de servicio de datos u otras aplicaciones que puedan almacenar en antememoria nombres de host privados.

```
phys-schost# clresource disable resource[,...]
```

Incluya lo siguiente en las aplicaciones que deshabilita.

- Servicios HA-DNS y HA-NFS, si están configurados

- Cualquier aplicación que se haya configurado de manera personalizada para utilizar el nombre de host privado
- Cualquier aplicación que utilicen los clientes mediante la interconexión privada

Para obtener más información sobre el uso del comando `clresource`, consulte la página del comando `man clresource(1CL)` y la [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 de NTP en cada nodo del clúster.

Utilice el comando `svcadm` para cerrar el daemon NTP. Consulte la página del comando `man svcadm(1M)` para obtener más información sobre el daemon NTP.

```
phys-schost# svcadm disable ntp
```

3 Ejecute la utilidad `clsetup` para cambiar el nombre de host privado del nodo apropiado.

Ejecute la utilidad desde solamente uno de los nodos del clúster. Para obtener más información, consulte la página del comando `man clsetup(1CL)`.

Nota – Cuando seleccione un nombre nuevo para el sistema privado, compruebe que no se utilice en el nodo del clúster.

También puede ejecutar el comando `clnode` en lugar de la utilidad `clsetup` para cambiar el nombre de host privado. En el siguiente ejemplo, el nombre de nodo del clúster es `pred1`. Después de ejecutar el siguiente comando `clnode`, vaya al [Paso 6](#).

```
phys-schost# /usr/cluster/bin/clnode set -p privatehostname=New-private-nodename pred1
```

4 En la utilidad `clsetup`, escriba el número para la opción para el nombre de host privado.

5 En la utilidad `clsetup`, escriba el número para la opción de cambiar un nombre de host privado.

Responda las preguntas cuando se lo solicite. 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 Purgue la antememoria del servicio de nombres.

Efectúe este paso en todos los nodos del clúster. El vaciado 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 cambió un nombre de host privado en la configuración de NTP o incluye un archivo, actualice el archivo de NTP en cada nodo. Si cambió un nombre de host privado en su archivo de configuración NTP (`/etc/inet/ntp.conf`) y tiene entradas de host de igual o un puntero que

incluye el archivo para los host de iguales en su archivo de configuración NTP (/etc/inet/ntp.conf.include), actualice el archivo en cada nodo. Si cambió un nombre de host privado en su archivo incluido NTP, actualice el archivo /etc/inet/ntp.conf.sc en cada nodo.

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. En general, el archivo ntp.conf.sc 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.

Realice este paso en cada nodo del clúster.

Utilice el comando svcadm para reiniciar el daemon de NTP.

```
# svcadm enable svc:network/ntp:default
```

8 Habilite todos los recursos de servicio de datos y otras aplicaciones deshabilitados en el Paso 1.

```
phys-schost# clresource enable resource[,...]
```

Para obtener más información sobre el uso del comando clresource, consulte la página del comando man clresource(1CL) y la *Oracle Solaris Cluster Data Services Planning and Administration Guide*.

Ejemplo 9–8 Cambio de nombre de host privado

En el siguiente ejemplo, se cambia el nombre de host privado de clusternode2-priv a clusternode4-priv, en el nodo phys-schost-2. Realice esta acción en cada nodo.

```
[Disable all applications and data services as necessary.]
phys-schost-1# svcadm disable ntp
phys-schost-1# clnode show | grep node
...
private hostname:                clusternode1-priv
private hostname:                clusternode2-priv
private hostname:                clusternode3-priv
...
phys-schost-1# clsetup
phys-schost-1# nscd -i hosts
phys-schost-1# vi /etc/inet/ntp.conf.sc
...
peer clusternode1-priv
peer clusternode4-priv
peer clusternode3-priv
phys-schost-1# ping clusternode4-priv
phys-schost-1# svcadm enable ntp
```

[Enable all applications and data services disabled at the beginning of the procedure.]

▼ Cómo cambiar el nombre de un nodo

Puede cambiar el nombre de un nodo que es 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 `clnode 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 cambia el nombre de un nodo en un clúster Geographic Edition de Oracle Solaris Cluster que está en sociedad con una configuración de Oracle Solaris, debe realizar pasos adicionales.**
Para obtener más información sobre los clústeres y los nodos de Geographic Edition, consulte el [Capítulo 5, “Administering Cluster Partnerships” de Oracle Solaris Cluster Geographic Edition System Administration Guide](#).

Si el clúster donde realiza el procedimiento de cambio de nombre es esencial para el grupo de protección y desea que la aplicación esté en el grupo de protección en línea, puede conmutar el grupo de protección al clúster secundario durante el procedimiento de cambio de nombre.
- 3 Cambie los nombres de host de Oracle Solaris siguiendo los pasos indicados en “[How to Change a System’s Identity \(nodename\)](#)” de Oracle Solaris Administration: Common Tasks, 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 cambió el nombre del host de Oracle Solaris, cambie el nombre del nodo y ejecute el comando `cmd` en cada host al que se le cambió el nombre.**
Cambie el nombre de un nodo a la vez.

```
# clnode rename -n newnodename oldnodename
```
- 6 Actualice cualquier referencia existente al nombre de host anterior en las aplicaciones que se ejecutan 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 de clúster.

```
# sync;sync;sync;reboot
```

9 Verifique 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 al que se le modificó el nombre mediante `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 cambiar la propiedad `hostnameList` de los recursos del 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 203](#) 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 202](#). 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 de clúster.

- a. En el modo de clúster, desconecte 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 para entradas anteriores en la propiedad `hostnameList` para hacer referencia a las nuevas entradas.

- e. Habilite los 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 -emM apache-rg
```

- g. Confirme que la aplicación se haya iniciado correctamente; para ello, ejecute el siguiente comando para realizar la comprobación de un cliente.

```
# clrs status apache-rs
```

▼ Cómo poner un nodo en estado de mantenimiento

Ponga un nodo del clúster global en estado de mantenimiento al sacar el nodo fuera de servicio por un período de tiempo extendido. 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 se debe cerrar con los comandos `clnode evacuate` y `cluster shutdown`. Para obtener más información, consulte las páginas del comando [man clnode\(1CL\)](#) y [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.

Quando 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.

Utilice el comando `clquorum disable` para poner un nodo de clúster en estado de mantenimiento. Para obtener más información, consulte la página del comando [man clquorum\(1CL\)](#).

`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 o asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo del clúster global que pone en estado de mantenimiento.**

- 2 **Evacue todos los grupos de recursos y de dispositivos del nodo.**

El comando `clnode evacuate` conmuta todos los grupos de recursos y de dispositivos.

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

- 3 **Cierre el nodo que evacuó.**

```
phys-schost# shutdown -g0 -y -i 0
```

- 4 **Conviértase en superusuario o asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en otro nodo en el clúster y ponga el nodo que cierra en el Paso 3 en estado de mantenimiento.**

```
phys-schost# clquorum disable node
```

node Especifica el nombre de un nodo que desea poner en modo de mantenimiento.

- 5 **Compruebe que 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` (cero) para los votos de quórum `Present` y `Possible`.

Ejemplo 9–9 Cómo poner un nodo del clúster global en estado de mantenimiento

En el siguiente ejemplo se pone un nodo de 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` puede indicar que algunos dispositivos del disco de quórum también se encuentran desconectados.

```
[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 [“Cómo sacar un nodo del estado de mantenimiento” en la página 206](#).

▼ Cómo sacar un nodo del estado de mantenimiento

Utilice el siguiente procedimiento para volver a poner en línea un nodo del clúster global y restablecer el recuento de votos de quórum al valor predeterminado. 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 número 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 un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en cualquier nodo del clúster global distinto del que está en estado de mantenimiento.**
- 2 Según el número de nodos que tenga en la configuración del clúster global, realice uno de los siguientes pasos:**
 - 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 elimina del estado de mantenimiento va a tener dispositivos de quórum, restablezca el recuento de quórum del clúster de un nodo distinto del que está en estado de mantenimiento.

El recuento de quórum debe restablecerse de un nodo que no esté en estado de mantenimiento antes de reiniciar 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 recuento de votos adecuado para los votos de quórum Present y Possible.

Ejemplo 9–10 Eliminación de un nodo de clúster del estado de mantenimiento y restablecimiento del recuento de votos de quórum

En el ejemplo siguiente se restablece el recuento de quórum para un nodo de clúster y sus dispositivos de quórum a los valores predeterminados y se 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. El Quorum Summary debe mostrar también un incremento en los recuentos 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.

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

▼ Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster

Realice este procedimiento para desconfigurar el software de Oracle Solaris Cluster de un nodo del clúster global antes de desconectarlo de una configuración de 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 cambio, vaya a [“Cómo anular la configuración del software Oracle Solaris Cluster para solucionar problemas de instalación” de Guía de instalación del software 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 Asegúrese de haber completado correctamente todas las tareas de requisitos previos en el mapa de tareas para eliminar un nodo de clúster.

Consulte la [Tabla 8-2](#).

Compruebe que haya eliminado el nodo de la configuración del clúster mediante `clnode remove` antes de continuar con este procedimiento. Otros pasos posiblemente incluyan el agregado del nodo que planea desinstalar a la lista de autenticación `node-` del clúster, la desinstalación de un clúster de zona, etc.

Nota – Para desconfigurar el nodo pero dejar el software Oracle Solaris Cluster instalado en el nodo, no realice ninguna acción después de ejecutar el comando `clnode remove`.

2 Conviértase en superusuario en el nodo que vaya a desinstalar.

3 Si su nodo tiene una partición dedicada para el espacio de nombres de dispositivos globales, reinicie el nodo del clúster global en modo que no sea de clúster.

- En un sistema basado en SPARC, ejecute el siguiente comando.

```
# shutdown -g0 -y -i0 ok boot -x
```

- En un sistema basado en x86, ejecute los siguientes comandos.

```
# shutdown -g0 -y -i0
...
<<< Current Boot Parameters >>>
Boot path: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/pci8086,341a@7,1/
sd@0,0:a
Boot args:

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

<<< timeout in 5 seconds >>>
Select (b)oot or (i)nterpreter: b -x
```

4 En el archivo `/etc/vfstab`, elimine todas las entradas del sistema de archivos montadas globalmente *excepto* los montajes globales `/global/.devices`.

5 Reinicie el nodo en un modo que no sea de clúster.

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

```
ok boot -x
```

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

- 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.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Booting an x86 Based System to a Specified State \(Task Map\)”](#) de *Booting and Shutting Down Oracle Solaris on x86 Platforms*.

- b. En la pantalla de parámetros de inicio, use las teclas de flecha para seleccionar la entrada `kernel` y escriba `e` para editarla.

- c. Agregue `-x` al comando para especificar que el sistema se inicia en el modo sin clúster.

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

La pantalla muestra el comando editado.

- e. **Escriba b para iniciar el nodo en el modo sin clúster.**

Nota – Este cambio en el comando del parámetro de inicio del núcleo no se conserva tras el inicio del sistema. La siguiente vez que reinicie el nodo, se iniciará en el modo de clúster. Si, por el contrario, desea iniciar en el modo sin clúster, siga estos pasos para volver a agregar la opción -x al comando del parámetro de inicio del núcleo.

- 6 **Acceda a un directorio como, por ejemplo, el directorio raíz (/), que no contenga ningún archivo proporcionado por los paquetes de Oracle Solaris Cluster.**

```
phys-schost# cd /
```

- 7 **Para anular la configuración del nodo y eliminar el software de Oracle Solaris Cluster, ejecute el siguiente comando.**

```
phys-schost# scinstall -r [-b bename]
```

- | | |
|---------------------------------|---|
| -r | Quita la información de configuración del clúster y desinstala la estructura de Oracle Solaris Cluster y el software de servicios de datos del nodo del clúster. Puede reinstalar el nodo o eliminarlo del clúster. |
| -b <i>nombre_entorno_inicio</i> | Especifica el nombre de un nuevo entorno de inicio, que es donde iniciará una vez completado el proceso de desinstalación. La especificación de un nombre es opcional. Si no especifica un nombre para el entorno de inicio, se genera uno automáticamente. |

Consulte la página del comando `man scinstall(1M)` para obtener más información.

- 8 **Si tiene previsto volver a instalar el software de Oracle Solaris Cluster en este nodo una vez finalizada la desinstalación, reinicie el nodo para iniciar en el nuevo entorno de inicio.**
- 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 abierto 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. Siga la documentación enviada con el servidor y adaptador de host para procedimientos de desconexión.

Consejo – Para obtener más información acerca de la migración de un espacio de nombres de dispositivos globales a Iofi, consulte [“Migración del espacio de nombre de dispositivos globales” en la página 90](#).

Resolución de problemas de desinstalación de nodos

En esta sección se describen los mensajes de error que puede recibir cuando ejecuta el comando `clnode remove` y las medidas correctivas que debe utilizar.

Entradas del sistema de archivos de clúster no eliminadas

Los siguientes mensajes de error indican que el nodo del clúster global que ha eliminado todavía tiene sistemas de archivos del clúster a los que se hace referencia 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 [“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de clúster” en la página 208](#) y repita el procedimiento. Compruebe que haya realizado correctamente el [Paso 4](#) del procedimiento antes de volver a ejecutar el comando `clnode remove`.

Lista no eliminada de grupos de dispositivos

Los siguientes mensajes de error 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.

Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster

En esta sección se describe cómo crear, configurar y gestionar la base de información de administración (MIB) de eventos del Protocolo simple de administración de red (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 `clsnmphost`. La MIB mantiene una tabla de sólo lectura con los 50 eventos más actuales. Debido a que los clústeres 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 reiniciar.

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 9-2 Mapa de tareas: creación, configuración y administración de la MIB de eventos de SNMP de Oracle Solaris Cluster

Tarea	Instrucciones
Habilite una MIB de eventos de SNMP	“Cómo habilitar una MIB de eventos de SNMP” en la página 213
Deshabilite una MIB de eventos de SNMP	“Cómo deshabilitar una MIB de eventos de SNMP” en la página 213
Cambie una MIB de eventos de SNMP	“Cómo cambiar una MIB de eventos de SNMP” en la página 214
Agregue un host de SNMP a la lista de hosts que recibirá notificaciones de captura para las MIB	“Cómo habilitar un host de SNMP para que reciba capturas de SNMP en un nodo” en la página 215

TABLA 9-2 Mapa de tareas: creación, configuración y administración de la MIB de eventos de SNMP de Oracle Solaris Cluster (Continuación)

Tarea	Instrucciones
Elimine un host de SNMP	“Cómo deshabilitar un host de SNMP para que no reciba capturas de SNMP en un nodo” en la página 215
Agregue un usuario de SNMP	“Cómo agregar un usuario de SNMP a un nodo” en la página 216
Elimine un usuario de SNMP	“Cómo eliminar un usuario de SNMP de un nodo” en la página 217

▼ **Cómo habilitar una MIB de eventos de SNMP**

En este procedimiento se 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 proporcione la autorización de RBAC solaris.cluster.modify.**
- 2 **Habilite la MIB de eventos de SNMP.**

phys - schost - l# clsnmpmib enable [-n node] MIB

[-n nodo]

Especifica el *nodo* en el que se ubica la MIB de eventos 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.

▼ **Cómo deshabilitar una MIB de eventos de SNMP**

En este procedimiento se muestra cómo deshabilitar 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 Deshabilite la MIB de eventos de SNMP.

```
phys-schost-1# clsnmpmib disable -n node MIB
```

-n nodo Especifica el *nodo* en el que se ubica la MIB de eventos que desea deshabilitar. 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 tipo de MIB que desea deshabilitar. En este caso, debe especificar el tipo event.

▼ Cómo cambiar una MIB de eventos de SNMP

En este procedimiento se muestra cómo cambiar el protocolo para 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 de eventos 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. Especifica el *valor* como se muestra 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 se especifica este operando, este subcomando utiliza el signo más predeterminado (+), que significa todas las MIB. Si utiliza el operando

MIB, especifique la MIB en una lista delimitada por espacios después de todas las demás opciones de líneas de comandos.

▼ **Cómo habilitar un host de SNMP para que reciba capturas de SNMP en un nodo**

En este procedimiento se muestra cómo agregar un host de SNMP de un nodo a la lista de hosts que recibirá notificaciones de capturas 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# clsnmphost add -c SNMPcommunity [-n node] host
-c comunidad_SNMP
```

Especifica el nombre de la comunidad de SNMP que se usa junto con el nombre del host.

Se debe especificar el nombre de la comunidad de SNMP *comunidad_SNMP* cuando agregue un host a una comunidad que no sea `public`. 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.

▼ **Cómo deshabilitar un host de SNMP para que no reciba capturas de SNMP en un nodo**

En este procedimiento se muestra cómo eliminar un host de SNMP de un nodo de la lista de hosts que recibirá notificaciones de capturas 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 especificado.

-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* del que se ha eliminado el host de SNMP de la configuración. 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 hosts de la comunidad SNMP especificada, use un signo más (+) para el *host* con la opción *-c*. Para eliminar todos los hosts, use el signo más (+) para el *host*.

▼ **Cómo agregar un usuario de SNMP a un nodo**

En este procedimiento se muestra cómo agregar un usuario de SNMP a la configuración de usuario 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 agrega 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 que contiene 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 siguientes caracteres: ▪ ; (punto y coma) ▪ : (dos puntos) ▪ \ (barra diagonal inversa) ▪ \n (línea nueva)
<i>usuario</i>	Especifica el nombre del usuario de SNMP que desea agregar.

▼ **Cómo eliminar un usuario de SNMP de un nodo**

En este procedimiento se muestra cómo eliminar un usuario de SNMP de la configuración de usuario 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 Elimine el usuario de SNMP.**

`phys-schost-1# clsnmpuser delete -n node user`

-n <i>nodo</i>	Especifica el nodo del que se elimina 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.

Configuración de límites de carga

Puede habilitar la distribución automática de la carga del grupo de recursos entre los nodos estableciendo límites de carga. Puede configurar un conjunto de límites de carga para cada nodo del clúster. Asigna factores de carga a grupos de recursos y los factores de carga se corresponden con los límites de carga 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. Debido a 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 respecto a los límites de carga de ese nodo.

Un límite de carga consta de los siguientes elementos:

- Un nombre asignado por el usuario.
- Un valor de límite flexible: un límite de carga flexible se puede exceder temporalmente.
- Un valor de límite fijo: los límites de carga fijos no pueden excederse nunca y se aplican de manera 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 del comando `man clnode(1CL)` para obtener más información.

También puede configurar un grupo de recursos para que tenga una prioridad superior y reducir así la probabilidad de ser desplazado 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 un rol 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 en los que quiera usar equilibrio 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 del 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 del 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 otorgar a algunos grupos de recursos una prioridad más alta que a otros.**

```
# clresourcegroup set -p priority=600 rg1
```

El valor predeterminado de prioridad 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
```

Para obtener más información, consulte la página del comando `man clresourcegroup(1CL)` en las opciones `HAS_COST`, `NO_COST` y `NEVER`.

- 7 **Si lo desea, 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 siguiente ejemplo 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 del clúster global y del clúster de zona del clúster.

```
# clnode status -Z all -v
```

La salida incluye todos los valores de configuración de límite de carga definidos en el nodo.

Cambio de números de puerto para servicios o agentes de gestión

Common Agent Container se inicia automáticamente al iniciar el clúster.

Nota – Si recibe un mensaje de error del sistema al intentar ver información sobre un nodo, compruebe si el parámetro de la dirección de enlace a la red del contenedor de agente común está definido en el valor correcto de 0.0.0.0.

Realice estos pasos en cada nodo del clúster.

1. Muestre el valor del parámetro network-bind-address.

```
# cacaoadm get-param network-bind-address
network-bind-address=0.0.0.0
```

2. Si el valor de parámetro no es 0.0.0.0, cambie el valor del parámetro.

```
# cacaoadm stop
# cacaoadm set-param network-bind-address=0.0.0.0
# cacaoadm start
```

▼ Cómo utilizar Common Agent Container para cambiar los números de puerto para servicios o agentes de gestión

Si los números de puerto predeterminados para los servicios de contenedor de agente común presentan conflictos con otros procesos de ejecución, puede usar el comando `cacaoadm` en cada nodo del clúster para cambiar el número de puerto del agente de gestión o servicio causante del conflicto.

- 1 Detenga el daemon de administración contenedor de agentes común en todos los nodos del clúster.

```
# /opt/bin/cacaoadm stop
```

- 2 Recupere el número de puerto utilizado actualmente por el servicio de contenedor de agente común con el subcomando `get-param`.

```
# /opt/bin/cacaoadm get-param parameterName
```

Puede usar el comando `cacaoadm` para cambiar los números de puerto para los siguientes servicios de contenedor de agente común. La lista siguiente proporciona algunos ejemplos de servicios y agentes que Common Agent Container puede administrar, junto con los nombres de parámetros correspondientes.

Puerto de conector JMX	<code>jmxmp-connector-port</code>
Puerto SNMP	<code>snmp-adapter-port</code>
Puerto de captura SNMP	<code>snmp-adapter-trap-port</code>
Puerto de flujo de comando	<code>commandstream-adapter-port</code>

Nota – Si recibe un mensaje de error del sistema al intentar ver información sobre un nodo, compruebe si el parámetro de la dirección de enlace a la red del contenedor de agente común está definido en el valor correcto de `0.0.0.0`.

Realice estos pasos en cada nodo del clúster.

1. Muestre el valor del parámetro `network-bind-address`.

```
# cacaoadm get-param network-bind-address
network-bind-address=0.0.0.0
```

2. Si el valor de parámetro no es `0.0.0.0`, cambie el valor del parámetro.

```
# cacaoadm stop
# cacaoadm set-param network-bind-address=0.0.0.0
# cacaoadm start
```

- 3 Cambie un número de puerto.
/opt/bin/cacaoadm set-param parameterName=parameterValue
- 4 Repita el Paso 3 en cada nodo del clúster.
- 5 Reinicie el daemon de gestión de contenedor de agentes común en todos los nodos del clúster.
/opt/bin/cacaoadm start

Cómo realizar 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 ejecute aplicaciones y clonar un clúster de zona. Todos estos comandos deben ejecutarse desde el nodo de votación del clúster global.

Nota – Las zonas en un clúster de zona se configuran cuando ejecuta `clzonecluster install -c` para configurar los perfiles. Consulte “Configuración de un clúster de zona” de *Guía de instalación del software de Oracle Solaris Cluster* para obtener instrucciones sobre el uso de la opción `-c config_profile`.

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 del comando `man` de Oracle Solaris Cluster para obtener información sobre la validez del uso de un comando en zonas.

TABLA 9-3 Otras tareas del clúster de zona

Tarea	Instrucciones
Mueva la ruta de zona a una nueva ruta de zona	<code>clzonecluster move -f ruta_zona nombre_clúster_zona</code>
Prepare el clúster de zona para ejecutar aplicaciones	<code>clzonecluster ready -n nombre_nodo nombre_clúster_zona</code>
Clone un clúster de zona	<code>clzonecluster clone -Z nombre_clúster_zona_origen [-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.
Elimine un clúster de zona	“Cómo eliminar un clúster de zona” en la página 223

TABLA 9-3 Otras tareas del clúster de zona (Continuación)

Tarea	Instrucciones
Elimine un sistema de archivos de un clúster de zona	“Cómo eliminar un sistema de archivos de un clúster de zona” en la página 224
Elimine un dispositivo de almacenamiento de un clúster de zona	“Cómo eliminar un dispositivo de almacenamiento de un clúster de zona” en la página 226
Resolución de problemas de desinstalación de nodos	“Resolución de problemas de desinstalación de nodos” en la página 211
Cree, configure y gestione la MIB de eventos de SNMP de Oracle Solaris Cluster	“Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster” en la página 212

▼ Cómo eliminar un clúster de zona

Puede eliminar un clúster de zona específico o usar un comodín para eliminar todos los clústeres de zona configurados en el clúster global. El clúster de zona debe estar configurado antes de eliminarlo.

- 1 **Conviértase en un superusuario o asuma un rol 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 sus recursos del clúster de zona.**

```
phys-schost# clresourcegroup delete -F -Z zoneclustername +
```

Nota – Este paso debe realizarse desde un nodo del clúster global. En cambio, para realizar este paso desde un nodo del clúster de zona, inicie una sesión en el nodo del clúster de zona y omita `-Z clúster_zona` del 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
```

▼ **Cómo eliminar un sistema de archivos de un clúster de zona**

Un sistema de archivos se puede exportar a un clúster de zona mediante un montaje directo o un montaje en bucle de retorno.

Los clústeres de zona admiten montajes directos para lo siguiente:

- Sistema local de archivos UFS
- Oracle Solaris ZFS (exportado como conjunto de datos)
- NFS desde dispositivos NAS admitidos

Los clústeres de zona pueden gestionar montajes en bucle de retorno para lo siguiente:

- Sistema local de archivos UFS
- Sistema de archivos de clúster de UFS

Puede configurar un recurso `HAStoragePlus` o `ScalMountPoint` para gestionar el montaje del sistema de archivos. 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 de 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 de este procedimiento se realizan desde un nodo del 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 del clúster de zona que va a eliminar.

```
phys-schost# clresource delete -F -Z zoneclustername fs_zone_resources
```

- b. Si es necesario, identifique y elimine los recursos de Oracle Solaris Cluster que están configurados en el clúster global para el sistema de archivos que va a eliminar.

```
phys-schost# clresource delete -F fs_global_resources
```

Utilice la opción **-F** con cuidado porque fuerza la eliminación de todos los recursos que especifique, incluso si no los ha deshabilitado 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 del 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 sistema de archivos está especificado por **dir=**.

5 Compruebe que se haya eliminado el sistema de archivos.

```
phys-schost# clzonecluster show -v zoneclustername
```

Ejemplo 9–12 Eliminación de un sistema de archivos local de alta disponibilidad en un clúster de zona

En este ejemplo se 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**. El recurso es **hasp-rs** y es del tipo **HASStoragePlus**.

```
phys-schost# clzonecluster show -v sczone
...
Resource Name:                fs
  dir:                        /local/ufs-1
  special:                    /dev/md/ds1/dsk/d0
  raw:                        /dev/md/ds1/rdisk/d0
  type:                        ufs
  options:                    [logging]
...
phys-schost# clresource delete -F -Z sczone hasp-rs
phys-schost# clzonecluster configure sczone
clzc:sczone> remove fs dir=/local/ufs-1
clzc:sczone> commit
phys-schost# clzonecluster show -v sczone
```

Ejemplo 9-13 Eliminación de un sistema de archivos ZFS de alta disponibilidad en un clúster de zona

En este ejemplo, se muestra cómo eliminar un sistema de archivos ZFS en una agrupación ZFS llamada HAZpool, que está configurada en el clúster de zona sczone del recurso hasp-rs del tipo SUNW.HAStoragePlus.

```
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
```

▼ Cómo eliminar un dispositivo de almacenamiento de un clúster de zona

Puede eliminar dispositivos de almacenamiento, por ejemplo grupos de discos de SVM y 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 de este procedimiento se realizan desde un nodo del 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.HAStoragePlus 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

En este ejemplo se muestra cómo eliminar un conjunto de discos de Solaris Volume Manager 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

En este ejemplo se muestra cómo eliminar 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
```

Resolución de problemas

En esta sección se incluyen procedimientos de resolución de problemas que puede utilizar con fines de prueba.

Ejecución de una aplicación fuera del clúster global

▼ **Cómo tomar un metaconjunto de Solaris Volume Manager de nodos que se han iniciado en un modo que no es de clúster**

Siga este procedimiento para ejecutar una aplicación fuera del clúster global con fines 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 que se tomará más tarde de un modo que no sea de clúster.

```
phys-schost# clquorum add did
```

- b. Quite el dispositivo de quórum anterior.

```
phys-schost# clquorum remove did
```

- c. Si el dispositivo de quórum usa una reserva de SCSI2, quite la reserva de SCSI2 del quórum anterior y compruebe que no queden reservas de SCSI2.

Para saber cómo ejecutar los comandos `pgre`, es necesario instalar y utilizar el paquete Diagnostic Toolkit (`ha-cluster/diagnostic/tool-kit`) proporcionado por el equipo de asistencia de Oracle.

- 2 **Evacue el nodo del clúster global que desee iniciar en un modo que no sea el de clúster.**
`phys-schost# clresourcegroup evacuate -n targetnode`
- 3 **Desconecte todos los grupos de recursos que contengan recursos HASTorage o HASToragePlus 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 **Deshabilite todos los recursos de los grupos de recursos que desconectó.**
`phys-schost# clresource disable resourcename`
- 5 **Anule la gestión de grupos de recursos.**
`phys-schost# clresourcegroup unmanage resourcegroupname`
- 6 **Desconecte el grupo o los grupos de dispositivos correspondientes.**
`phys-schost# cldevicegroup offline devicegroupname`
- 7 **Deshabilite el grupo o los grupos de dispositivos.**
`phys-schost# cldevicegroup disable devicegroupname`
- 8 **Inicie el nodo pasivo en modo que no sea de clúster.**
`phys-schost# reboot -x`
- 9 **Antes de seguir, compruebe que haya finalizado el proceso de inicio en el nodo pasivo.**
`phys-schost# svcs -x`
- 10 **Determine si hay reservas de SCSI3 en los discos de los metaconjuntos.**
 Ejecute el siguiente comando en todos los discos de los metaconjuntos.
`phys-schost# /usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/dids2`
- 11 **Si existe alguna reserva de SCSI3 en los discos, elimínela.**
`phys-schost# /usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/dids2`
- 12 **Tome el metaconjunto del nodo evacuado.**
`phys-schost# metaset -s name -C take -f`
- 13 **Monte el sistema de archivos 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 deseada. Cuando finalice la prueba, detenga la aplicación.**
- 15 **Reinicie el nodo y espere hasta que el proceso de inicio haya finalizado.**
`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 -emM resourcegroupname
```

Restaurar un conjunto de discos dañado

Utilice este procedimiento si un conjunto de discos está dañado o está en un estado en que los nodos del clúster no pueden asumir la propiedad 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 de Oracle (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 `prtvto` 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 231](#).

1 Guarde la tabla de particiones para cada disco del conjunto de discos.

```
# /usr/sbin/prtvto /dev/global/rdisk/diskname > /etc/lvm/diskname.vto
```

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 de 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 de Oracle (SUNWexplo) se ejecuta en un nodo propietario del conjunto, recuperará la información de los comandos `prtvto` 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 de purga 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 (ha-cluster/diagnostic/tool-kit) 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 siguiente comando.

```
# /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. Si existe un problema, utilice el paquete Diagnostic Toolkit (ha-cluster/diagnostic/tool-kit) para eliminar la información de los archivos de configuración dcs.

Si las opciones del comando `purge` no se completan correctamente, compruebe si tiene instaladas las actualizaciones más recientes del núcleo y del metadispositivo, y póngase en contacto 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 nuevo 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 nuevo 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 lo crea nuevamente, el índice de contenido del volumen (VTOC) debe permanecer en los discos para que pueda omitir este paso.

Sin embargo, si vuelve a crear un conjunto para la recuperación, 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/rdisk/d4s2
```

```
# /usr/sbin/fmthard -s /etc/lvm/d8.vtoc /dev/global/rdisk/d8s2
```

Este comando puede ejecutarse en cualquier nodo del clúster.

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/rdisk/metadvice
```

Si el comando `fsck` sólo muestra algunos errores, como el recuento de superbloqueos, probablemente el dispositivo se haya reconstruido de manera 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 del comando `man rg_properties(5)`. Este capítulo proporciona información sobre los temas siguientes:

- “Introducción al control de la CPU” en la página 235
- “Configuración del control de la CPU” en la página 236

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 información sobre proyectos, agrupaciones de recursos, conjuntos de procesadores y clases de programación, consulte *Administración de Oracle Solaris: zonas de Oracle Solaris, zonas de Oracle Solaris 10 y gestión de recursos*.

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 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 grupos de recursos proporcionando valores para <code>project.cpu-shares</code> y <code>zone.cpu-shares</code> .	“Control del uso de la CPU en el nodo de votación de un clúster global” en la página 236

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 236](#)

▼ 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 del 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 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 una clase de programación, 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 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.

Si no asigna un valor a las propiedades `globalzoneshares` y `defaultpsetmin`, estas propiedades asumen sus valores predeterminados.

```
# clnode set [-p globalzoneshares=integer] \
[-p defaultpsetmin=integer] \
node
```

```
-p defaultpsetmin=entero_mínimo_establecido_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.

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.

```
# clresourcegroup create -p RG_SLM_TYPE=automated \  
  [-p RG_SLM_CPU_SHARES=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 CPU asignados al nodo de votación `zone.cpu-shares`.

nombre_grupo_recursos 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 -emM resource_group_name
```

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, puede configurar la propiedad `project.max-lwps`. Para obtener más información, consulte la página del comando [man projmod\(1M\)](#).

Actualización de software

En este capítulo se proporciona información e instrucciones para actualizar software de Oracle Solaris Cluster en las siguientes secciones.

- “Descripción general de actualización de software de Oracle Solaris Cluster” en la página 239
- “Actualización del software de Oracle Solaris Cluster” en la página 240
- “Desinstalación de un paquete” en la página 243

Descripción general de actualización de software de Oracle Solaris Cluster

Todos los nodos miembros del clúster deben tener las mismas actualizaciones aplicadas para la operación de clúster apropiada. Cuando se actualiza un nodo, es posible que ocasionalmente se deba eliminar la pertenencia de un nodo al clúster o detener todo el clúster antes de realizar la actualización.

Hay dos maneras de actualizar el software de Oracle Solaris Cluster.

- Actualización: actualiza el clúster a la última versión de Oracle Solaris Cluster (nueva versión o pequeña actualización) y actualiza el sistema operativo Oracle Solaris mediante la actualización de todos los paquetes. Un ejemplo de una nueva versión sería una actualización de Oracle Solaris Cluster 4.0 a 5.0. Un ejemplo de una pequeña actualización sería una actualización de Oracle Solaris Cluster 4.0 a 4.1. Ejecute la utilidad `scinstall` o el comando `scinstall -u update` para crear un nuevo entorno de inicio (una instancia iniciable de una imagen), monte el entorno de inicio en un punto de montaje que no esté en uso, actualice los bits y active el nuevo entorno de inicio. La creación del entorno de clon no consume inicialmente espacio adicional y se produce instantáneamente. Después de realizar esta actualización debe reiniciar el clúster. La actualización también actualiza el sistema operativo Oracle Solaris a la última versión compatible. Para obtener instrucciones detalladas, consulte la [Oracle Solaris Cluster Upgrade Guide](#).

Si tiene zonas de conmutación por error de tipo de marca `solaris`, siga las instrucciones en [“How to Upgrade Failover Zones” de Oracle Solaris Cluster Upgrade Guide](#).

- **Actualización:** actualiza paquetes de Oracle Solaris Cluster específicos a diferentes niveles de SRU. Puede utilizar uno de los comandos `pkg` para actualizar los paquetes de IPS (Image Packaging System) en una SRU (Service Repository Update). Las SRU se publican generalmente con frecuencia y contienen paquetes actualizados y correcciones para defectos. El depósito contiene todos los paquetes de IPS y los paquetes actualizados. Si ejecuta el comando `pkg update` se actualiza el sistema operativo Oracle Solaris y el software de Oracle Solaris Cluster a versiones compatibles. Después de realizar esta actualización, es posible que necesite reiniciar el clúster. Para obtener instrucciones, consulte [“Cómo actualizar un paquete específico” en la página 241](#).

Actualización del software de Oracle Solaris Cluster

Consulte la siguiente tabla para determinar cómo actualizar una versión o paquete de Oracle Solaris Cluster en el software de Oracle Solaris Cluster.

TABLA 11-1 Actualización del software de Oracle Solaris Cluster

Tarea	Instrucciones
Actualizar todo el clúster a una nueva versión o pequeña actualización	“How to Perform a Standard Upgrade” de Oracle Solaris Cluster Upgrade Guide
Actualizar un paquete específico	“Cómo actualizar un paquete específico” en la página 241
Actualizar un servidor de quórum o servidor de instalación AI	“Cómo actualizar un servidor de quórum o servidor de instalación AI” en la página 242
Eliminar paquetes de Oracle Solaris Cluster	“Cómo desinstalar un paquete” en la página 243 “Cómo desinstalar paquetes de servidor de quórum o servidor de instalación AI” en la página 243

Actualización del clúster a una nueva versión

No es necesario que coloque el clúster en modo que no sea de clúster antes de realizar esta actualización porque la actualización siempre se realiza en el nuevo entorno de inicio y el entorno de inicio existente permanece sin modificaciones. Puede especificar un nombre para el nuevo entorno de inicio o puede utilizar el nombre generado automáticamente. Para obtener instrucciones, consulte [“How to Perform a Standard Upgrade” de Oracle Solaris Cluster Upgrade Guide](#).

Cada vez que se actualice el software de Oracle Solaris Cluster, se deben actualizar también los servicios de datos y el software de Geographic Edition. Sin embargo, si desea actualizar los

servicios de datos por separado, consulte “[Overview of the Installation and Configuration Process](#)” de *Oracle Solaris Cluster Data Services Planning and Administration Guide*. Si desea actualizar Geographic Edition de Oracle Solaris Cluster, consulte la *Oracle Solaris Cluster Geographic Edition Installation Guide*.

El sistema operativo de Oracle Solaris también se actualiza a la última versión cuando actualiza al software de Oracle Solaris Cluster.

Actualización a un paquete específico

Los paquetes IPS se presentaron con el sistema operativo Oracle Solaris 11. El FMRI (Fault Managed Resource Indicator) describe cada paquete IPS y puede utilizar los comandos `pkg(1)` para realizar la actualización de SRU. Como alternativa, puede utilizar el comando `scinstall -u` para realizar una actualización de SRU.

Es posible que desee actualizar un paquete específico para utilizar un agente de servicio de datos de Oracle Solaris Cluster actualizado.

▼ Cómo actualizar un paquete específico

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

- 2 **Actualice el paquete.**

Por ejemplo, para actualizar un paquete de un editor específico, especifique el nombre del editor en *fmri_paquete*.

```
# pkg update pkg-fmri
```



Precaución – Si utiliza el comando `pkg update` sin ningún *fmri_paquete* especificado, se actualizan todos los paquetes instalados que tienen actualizaciones disponibles.

Si una nueva versión de un paquete instalado está disponible y es compatible con el resto de la imagen, el paquete se actualiza a esa versión. Si el paquete contiene binarios que tienen el indicador `reboot-needed` establecido en `true`, entonces si ejecuta `pkg update fmri_paquete` se crea automáticamente un nuevo entorno de inicio y después de la actualización se inicia en el nuevo entorno de inicio. Si el paquete que actualiza no contiene ningún binario que fuerce un reinicio, el comando `pkg update` actualiza la imagen en funcionamiento y no es necesario un reinicio.

- 3 **Si actualiza un agente de servicio de datos (`ha-cluster/data-service/*` o el agente de servicio de datos genérico de `ha-cluster/ha-service/gds`), siga estos pasos.**

a. `# pkg change-facet facet.version-lock.nombre_paquete=false`

b. # pkg update nombre_paquete

Por ejemplo:

```
# pkg change-facet facet.version-lock.ha-cluster/data-service/weblogic=false  
# pkg update ha-cluster/data-service/weblogic
```

Si desea congelar un agente y evitar que se actualice, realice los siguientes pasos.

```
# pkg change-facet facet.version-lock.nombre_paquete=false  
# pkg freeze nombre_paquete
```

Para obtener más información sobre la inmovilización de un determinado agente, consulte [“Control de la instalación de componentes opcionales” de Adición y actualización de paquetes de software de Oracle Solaris 11.](#)

4 Verifique que el paquete se haya actualizado.

```
# pkg verify -v pkg-fmri
```

Actualización de un servidor de quórum o servidor de instalación AI

Utilice el siguiente procedimiento para actualizar paquetes para su servidor de quórum o servidor de instalador AI (Automated Installer). Si desea obtener más información sobre servidores de quórum, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster](#). Si desea obtener más información sobre el uso de AI, consulte [“Instalación y configuración del software de Oracle Solaris Cluster y Oracle Solaris \(Automated Installer\)” de Guía de instalación del software de Oracle Solaris Cluster](#).

▼ Cómo actualizar un servidor de quórum o servidor de instalación AI

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.admin`.**
- 2 Actualice paquetes de servidor de quórum o de servidor de instalación AI.**

```
# pkg update ha-cluster/*
```

Si una nueva versión de los paquetes `ha-cluster` instalados está disponible y es compatible con el resto de la imagen, los paquetes se actualizan a esa versión.



Precaución – La ejecución del comando `pkg update` actualiza todos los paquetes `ha-cluster` instalados en el sistema.

Desinstalación de un paquete

Puede eliminar un solo paquete o varios paquetes.

▼ Cómo desinstalar un paquete

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

- 2 Desinstale un paquete existente.

```
# pkg uninstall pkg-fmri
```

Si desea desinstalar más de un paquete, utilice la siguiente sintaxis.

```
# pkg uninstall pkg-fmri pkg-fmri
```

El comando `pkg uninstall` fallará si hay otros paquetes instalados que dependen del `fmri_paquete` que desinstala. Para desinstalar el `fmri_paquete`, debe proporcionarle al comando `pkg uninstall` todos los dependientes de `fmri_paquete`. Para obtener información adicional sobre la desinstalación de paquetes, consulte [Adición y actualización de paquetes de software de Oracle Solaris 11](#) y la página del comando `man pkg(1)`.

▼ Cómo desinstalar paquetes de servidor de quórum o servidor de instalación AI

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

- 2 Desinstale paquetes de servidor de quórum o de servidor de instalación AI.

```
# pkg uninstall ha-cluster/*
```



Precaución – Este comando desinstala todos los paquetes `ha-cluster` instalados en el sistema.

Consejos de actualización

Utilice los siguientes consejos para administrar actualizaciones de Oracle Solaris Cluster de manera más eficaz:

- Lea el archivo README de SRU antes de realizar la actualización.

- Compruebe los requisitos de actualización de los dispositivos de almacenamiento.
- Aplique todas las actualizaciones 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. Consulte la documentación del hardware para obtener información sobre las actualizaciones del firmware.
- Todos los nodos que actúan como miembros del clúster deben tener las mismas actualizaciones.
- Mantenga actualizados los subsistemas del clúster. Entre otros, estas actualizaciones incluyen la administración de volúmenes, el firmware de dispositivos de almacenamiento y el transporte del clúster.
- Pruebe la conmutación por error tras haber aplicado actualizaciones importantes. Prepárese para retirar la actualización si disminuye el rendimiento del clúster o si no funciona correctamente.
- Si actualiza a una nueva versión de Oracle Solaris Cluster, siga las instrucciones en la [Oracle Solaris Cluster Upgrade Guide](#).

Copias de seguridad y restauraciones de clústeres

Este capítulo contiene las secciones siguientes.

- “Copia de seguridad de un clúster” en la página 245
- “Restauración de los archivos del clúster” en la página 248

Copia de seguridad de un clúster

Antes de realizar una copia de seguridad del clúster, busque los nombres de los sistemas de archivos de los que desea realizar copias de seguridad, calcule cuántas cintas necesita para colocar una copia de seguridad completa y realice una copia de seguridad del sistema de archivos raíz ZFS.

TABLA 12-1 Mapa de tareas: hacer una copia de seguridad de los archivos del clúster

Tarea	Instrucciones
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 245
Hacer una copia de seguridad de la configuración del clúster	“Copias de seguridad de la configuración del clúster” en la página 247
Realizar una copia de seguridad de la configuración de partición del disco de almacenamiento	Consulte la documentación del disco de almacenamiento

▼ 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 - 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 equivalente en el nodo del clúster del que esté haciendo la copia de seguridad.**

- 2 Use el comando `metaset` para determinar qué nodo tiene la propiedad en el volumen con copia de seguridad.**

```
# metaset -s setname
```

`-s setname` Especifica el nombre del conjunto de discos.

Para obtener más información, consulte la página del comando man [metaset\(1M\)](#).

- 3 Use el comando `lockfs` con la opción `-w` para evitar escrituras en el sistema de archivos.**

```
# lockfs -w mountpoint
```

Consulte la página del comando man [lockfs\(1M\)](#) para obtener más información.

- 4 Use el comando `metastat` para determinar los nombres de los subreflejos.**

```
# metastat -s setname -p
```

`-p` Muestra el estado en un formato similar al del archivo `md.tab`.

Consulte la página del comando man [metastat\(1M\)](#) para obtener más información.

- 5 Use el comando `metadetach` para desconectar un subreflejo del reflejo.**

```
# metadetach -s setname mirror submirror
```

Consulte la página del comando man [metadetach\(1M\)](#) para obtener más información.

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.**

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` para volver a conectar el metadispositivo o volumen.**

```
# metattach -s setname mirror submirror
```

Cuando se pone en línea un metadispositivo o un volumen, se vuelve a sincronizar automáticamente con la duplicación. Consulte la página de código man [metattach\(1M\)](#) para obtener más información.

- 10 **Use el comando `metastat` para comprobar que la subduplicación se vuelve a sincronizar.**

```
# metastat -s setname mirror
```

Consulte [Administración de Oracle Solaris: sistemas de archivos ZFS](#) para obtener más información.

▼ 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_configuración 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 la página del comando [man 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

Puede restaurar el sistema de archivos raíz ZFS en un disco nuevo.

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 el sistema de archivos raíz ZFS (/).	“Cómo restaurar el sistema de archivos raíz ZFS (/) (Solaris Volume Manager)” en la página 248

▼ Cómo restaurar el sistema de archivos raíz ZFS (/) (Solaris Volume Manager)

Use este procedimiento para restaurar sistemas de archivos raíz ZFS (/) en un disco nuevo, como, por ejemplo, después de reemplazar un disco raíz defectuoso. 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. Se admite UFS, excepto como un sistema de archivos raíz. UFS se puede usar en metadispositivos de metaconjuntos de Solaris Volume Manager de discos compartidos.

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 setname	Especifica el nombre del conjunto de discos.
-f	Elimina el último host del conjunto de discos.
-d	Elimina del conjunto de discos.
-h nodelist	Especifica el nombre del nodo que se va a eliminar del conjunto de discos.

- 3 **Restaura el sistema de archivos raíz ZFS (/).**

Para recuperar las instantáneas de la agrupación raíz o la agrupación raíz ZFS, siga el procedimiento en [“How to Replace a Disk in a ZFS Root Pool” de Oracle Solaris Administration: ZFS File Systems](#).

Nota – Asegúrese de crear el sistema de archivos `/global/.devices/node@nodeid`.

Si existe el archivo de copia de seguridad `/globaldevices` en el directorio de copia de seguridad, se restaura junto con la restauración raíz ZFS. El archivo no es creado automáticamente por el servicio SMF `globaldevices`.

4 Rearranque el nodo en modo multiusuario.

```
# reboot
```

5 Sustituya el ID del dispositivo.

```
# cldevice repair rootdisk
```

6 Use el comando metadb para recrear las réplicas de la base de datos de estado.

```
# metadb -c copies -af raw-disk-device
```

-c *copies* Especifica el número de réplicas que se van a crear.

-f *dispositivo_disco_sin procesar* Dispositivo de disco básico en el que crear las réplicas.

-a Agrega réplicas.

Consulte la página del comando man [metadb\(1M\)](#) para obtener más información.

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 odelist
```

-a Crea el host y lo agrega al conjunto de discos.

El nodo se rearranca en modo de clúster. El clúster está preparado para utilizarse.

Ejemplo 12–1 Restauración del sistema de archivos raíz ZFS (/) (Solaris Volume Manager)

En el ejemplo siguiente, se muestra el sistema de archivos raíz (/) restaurado en el nodo phys-schost-1. 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/rdisk/c0t0d0s0 y se vuelven a crear tres réplicas de la base de datos en /dev/rdisk/c0t0d0s4. Para obtener más información sobre la restauración de datos, consulte [“Repairing Damaged Data” de Oracle Solaris Administration: ZFS File Systems](#).

```
[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/rdisk/c0t0d0s4
[Add the node back to the metaset:]
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

Ejemplo

Configuración de replicación de datos basada en host con el software 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 entre clústeres. Consulte [“Comprensión de la replicación de datos” en la página 80](#).

En el ejemplo de este apéndice, se muestra cómo configurar la replicación de datos basada en host entre clústeres con el software StorageTek Availability Suite. 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 clústeres.

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 [Administración de Oracle Solaris: servicios de seguridad](#) para obtener más información sobre cómo usar roles 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 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 StorageTek Availability Suite.

La tolerancia a desastres es la capacidad de restaurar una aplicación en un clúster alternativo cuando el clúster principal falla. La tolerancia a desastres se basa en la *replicación de datos* y la *recuperación*. Una recuperación reubica un servicio de aplicaciones en un clúster secundario estableciendo en línea uno o más grupos de recursos y grupos de dispositivos.

Si los datos se replican de manera síncrona entre el clúster principal y el secundario, no se pierde ningún dato comprometido cuando el sitio principal falla. Sin embargo, si los datos se replican de manera asíncrona, es posible que algunos datos no se repliquen en el clúster secundario antes del fallo del sitio principal y, por lo tanto, se pierdan.

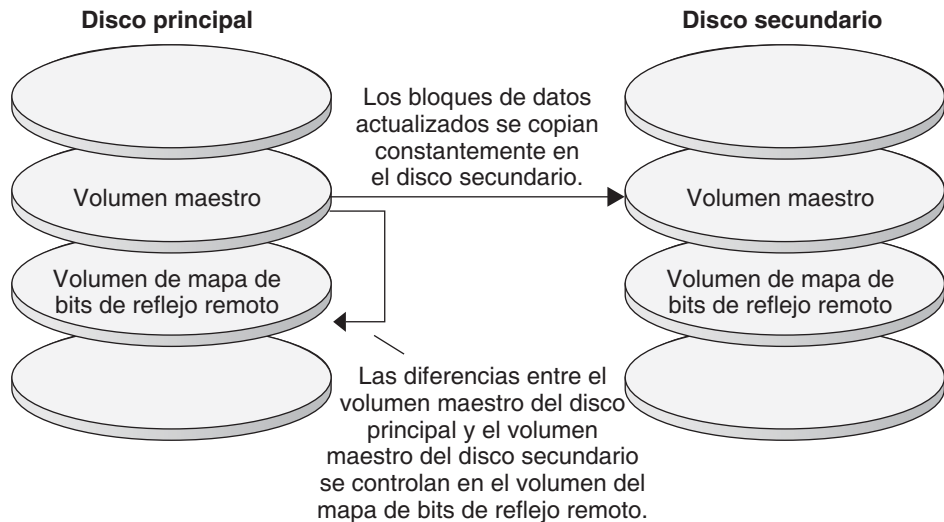
Métodos de replicación de datos utilizados por StorageTek Availability Suite

Esta sección describe los métodos de replicación por duplicación remota y de instantánea de un momento determinado utilizados por StorageTek Availability Suite. Este software utiliza los comandos `sndradm` e `iiadm` para replicar datos. Para obtener más información, consulte las páginas del comando `man sndradm(1M)` and `iiadm(1M)`.

Replicación de reflejo remoto

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 de reflejo remoto



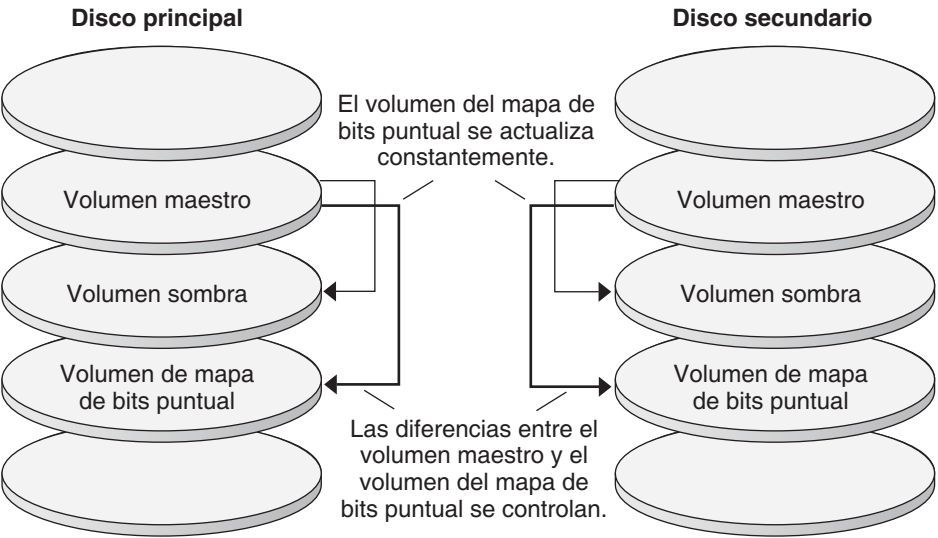
La replicación por duplicación remota se puede efectuar de manera sincrónica en tiempo real o de manera asincrónica. Cada volumen definido en cada clúster se puede configurar de manera individual, para replicación sincrónica o asincrónica.

- 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 puntual

La [Figura A-2](#) muestra una instantánea puntual. 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.

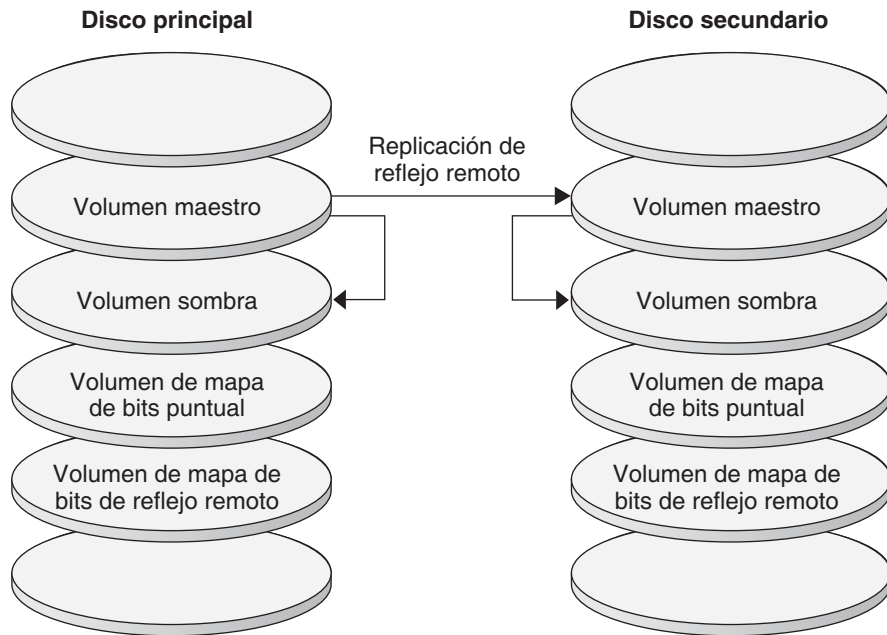
FIGURA A-2 Instantánea puntual



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 clústeres

Esta sección proporciona directrices para la configuración de replicación de datos entre clústeres. 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 255
- “Configuración de grupos de recursos de aplicaciones” en la página 256
 - “Configuración de grupos de recursos en una aplicación de migración tras error” en la página 257
 - “Configuración de grupos de recursos en una aplicación escalable” en la página 258
- “Directrices para la gestión de una recuperación” en la página 259

Configuración de grupos de recursos de replications

Los grupos de recursos de replications sitúan el grupo de dispositivos bajo el control del software StorageTek Availability Suite con un recurso de nombre de host lógico. Debe existir un

nombre de host lógico en cada extremo del flujo de replicación de datos y debe encontrarse en el mismo nodo del clúster que actúa como la ruta de E/S principal para el dispositivo. 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 tener un recurso de nombre de host lógico.

Un nombre de host lógico se aloja en un nodo de cada clúster (principal y secundario) y se utiliza para proporcionar direcciones de origen y destino para el flujo de replicación de datos del software StorageTek Availability Suite.

- Debe tener 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*.

Para obtener más información sobre HASToragePlus, consulte la página del 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.

La propiedad de extensión *ZPoolsSearchDir* se debe definir en el recurso HASToragePlus. Esta propiedad de extensión es necesaria para utilizar el sistema de archivos ZFS.

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 257
- “Configuración de grupos de recursos en una aplicación escalable” en la página 258

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:

- Debe tener un recurso de HASToragePlus para aplicar la conmutación por error del sistema de archivos o zpool cuando el grupo de recursos de aplicaciones se conmuta.

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.

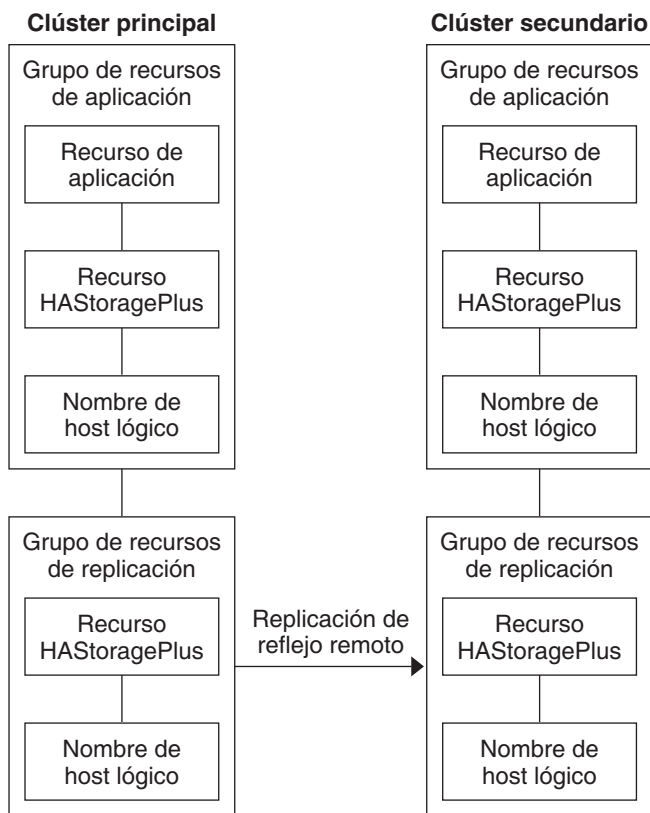
- Si los datos de la aplicación están montados de manera global, la presencia de un recurso de HASToragePlus en el grupo de recursos de aplicaciones no es obligatoria, aunque sí aconsejable.
- Si los datos de la aplicación se montan de manera local, la presencia de un recurso de HASToragePlus en el grupo de recursos de aplicaciones es obligatoria.

Para obtener más información sobre HASToragePlus, consulte la página del comando `man SUNW.HASToragePlus(5)`.

- Debe estar en línea en el clúster principal y sin conexión en el clúster 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.

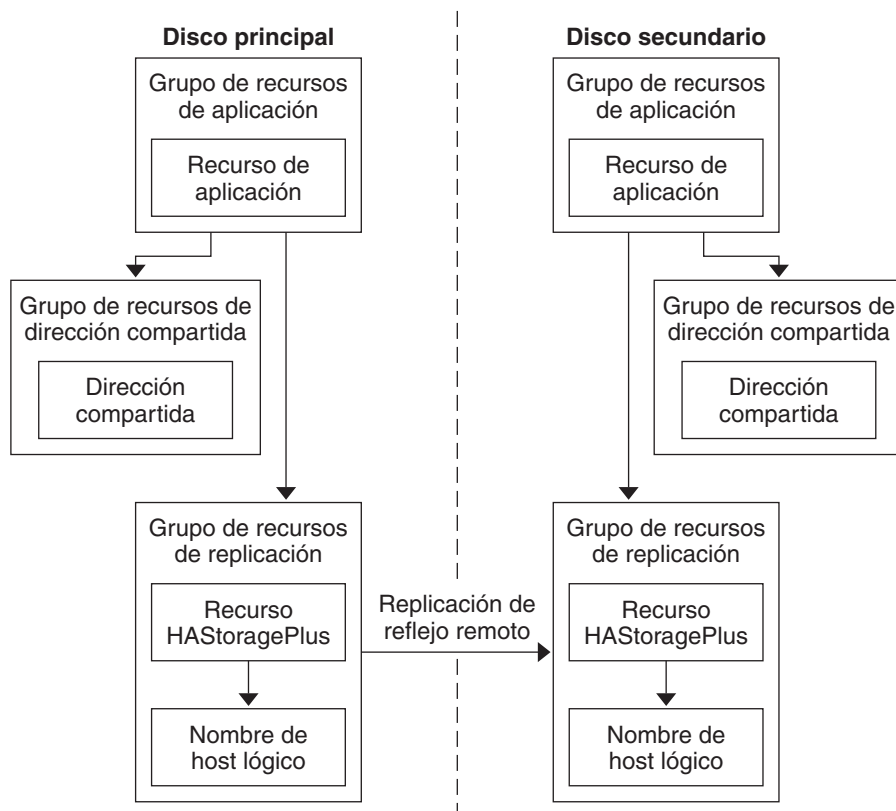
Cuando una aplicación escalable se administra como recurso en un grupo de recursos de aplicaciones, no es necesario acoplar el grupo de recursos de aplicaciones con el grupo de dispositivos. Por este motivo, no es necesario crear un recurso de HASToragePlus para el grupo de recursos de aplicaciones.

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

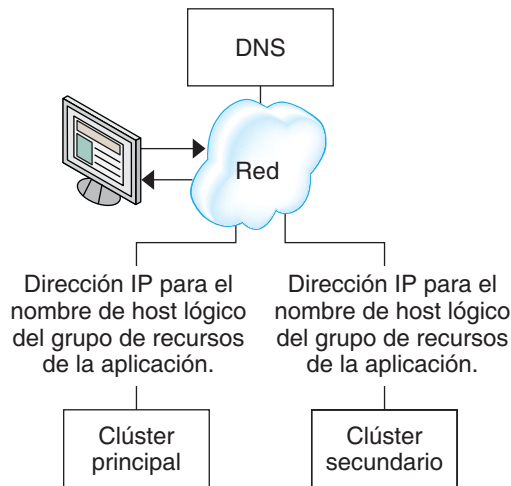


Directrices para la gestión de una recuperación

Si el clúster principal falla, la aplicación se debe conmutar al clúster secundario lo antes posible. Para que el clúster secundario pueda realizar las funciones del principal, se debe actualizar el DNS.

Los clientes usan el DNS para asignar el nombre de host lógico de una aplicación a una dirección IP. Después de una recuperación, donde la aplicación se mueve a un clúster secundario, la información del DNS debe actualizarse para reflejar la asignación entre el nombre de host lógico de la aplicación y la dirección IP nueva.

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 del comando `man nsupdate(1M)`. Por ver un ejemplo de cómo gestionar una recuperación, consulte [“Ejemplo de cómo gestionar una recuperación” en la página 285](#).

Después de la reparación, el clúster principal se puede volver a colocar 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. Puede hacerlo deteniendo el grupo de recursos en el nodo secundario, de modo que el flujo de datos de replicación pueda drenar.
2. Revierta la dirección de la replicación de datos, de modo que el nodo principal original esté ahora, otra vez, replicando los datos en el nodo secundario original.
3. Inicie el grupo de recursos en el clúster principal.
4. 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 StorageTek Availability Suite.

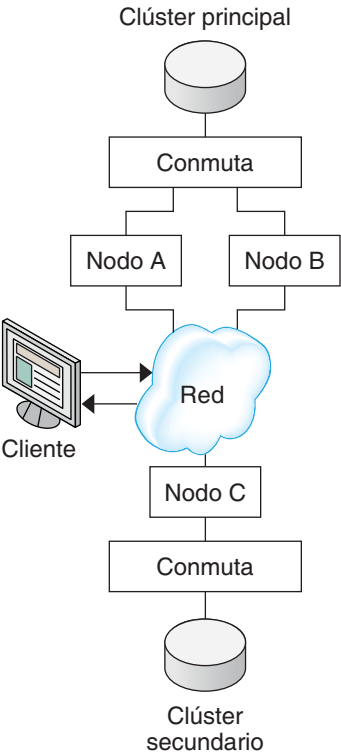
TABLA A-1 Mapa de tareas: ejemplo de configuración de una replicación de datos

Tarea	Instrucciones
1. Conectar e instalar los clústeres	“Conexión e instalación de clústeres” en la página 261
2. Configurar los grupos de dispositivos, los sistemas de archivos para la aplicación de NFS y los grupos de recursos de los clústeres principal y secundario	“Ejemplo de configuración de grupos de dispositivos y de recursos” en la página 263
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 277 “Habilitación de la replicación en el clúster secundario” en la página 278
4. Efectuar la replicación de datos	“Replicación por duplicación remota” en la página 280 “Instantánea de un momento determinado” en la página 281
5. Comprobar la configuración de la replicación de datos	“Procedimiento para comprobar la configuración de la replicación” en la página 282

Conexión e instalación de clústeres

La [Figura A-7](#) ilustra la configuración de clústeres 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 clúster



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 software Volume Manager deben instalarse en los nodos del clúster *antes* de instalar las actualizaciones del software y el software StorageTek Availability Suite.

TABLA A-2 Hardware y software necesarios

Hardware o software	Requisito
Hardware de nodo	El software StorageTek Availability Suite es compatible con todos los servidores que utilicen el sistema operativo Oracle Solaris. Para obtener información sobre el hardware que se debe usar, consulte el Oracle Solaris Cluster Hardware Administration Manual .
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 Guía de instalación del software de Oracle Solaris Cluster
Software Oracle Solaris Cluster	Software Oracle Solaris Cluster 4.0. Para obtener más información sobre la instalación, consulte Guía de instalación del software de Oracle Solaris Cluster.
Software Volume Manager	Software Solaris Volume Manager. Todos los nodos deben usar la misma versión del administrador de volúmenes. Para obtener información sobre la instalación, consulte el Capítulo 4, “Configuración del software de Solaris Volume Manager” de Guía de instalación del software de Oracle Solaris Cluster.
Software StorageTek Availability Suite	Diferentes clústeres pueden usar distintas versiones del sistema operativo Oracle Solaris y del software Oracle Solaris Cluster, pero usted debe usar la misma versión del software StorageTek Availability Suite entre clústeres. Si desea más información sobre cómo instalar el software, consulte los manuales de instalación de su versión del software StorageTek Availability Suite: StorageTek Availability Suite: documentación de StorageTek Availability
Actualizaciones del software StorageTek Availability Suite	Para obtener información sobre las últimas actualizaciones de software, inicie sesión en My Oracle Support.

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 replications](#)” en la página 255 y “[Configuración de grupos de recursos de aplicaciones](#)” en la página 256.

Esta sección incluye los procedimientos siguientes:

- “[Configuración de un grupo de dispositivos en el clúster primario](#)” en la página 265
- “[Configuración de un grupo de dispositivos en el clúster secundario](#)” en la página 266

- “Configuración de sistemas de archivos en el clúster primario para la aplicación NFS” en la página 267
- “Configuración del sistema de archivos en el clúster secundario para la aplicación NFS” en la página 268
- “Creación de un grupo de recursos de replications en el clúster primario” en la página 269
- “Creación de un grupo de recursos de replications en el clúster secundario” en la página 271
- “Creación de un grupo de recursos de aplicaciones NFS en el clúster primario” en la página 272
- “Creación de un grupo de recursos de aplicaciones NFS en el clúster secundario” en la página 274
- “Procedimiento para comprobar la configuración de la replicación” en la página 282

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.

Para obtener información sobre el software Solaris Volume Manager, consulte el [Capítulo 4, “Configuración del software de Solaris Volume Manager”](#) de *Guía de instalación del software de Oracle Solaris Cluster*.

▼ 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 StorageTek Availability Suite en un clúster” en la página 252
 - “Directrices para la configuración de replicación de datos basada en host entre clústeres” en la página 255
- Configure los clústeres primario y secundario como se describe en “Conexión e instalación de clústeres” en la página 261.

1 Obtenga acceso al nodo nodeA como superusuario o asuma un rol que proporcione 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 metaconjunto para colocar los datos NFS y la replicación asociada.

```
nodeA# metaset -s nfsset -h nodeA nodeB
```

3 Agregue discos al metaconjunto.

```
nodeA# metaset -s nfsset -a /dev/did/dsk/d6 /dev/did/dsk/d7
```

4 Agregue mediadores al metaconjunto.

```
nodeA# metaset -s nfsset -a -m nodeA nodeB
```

5 Cree los volúmenes requeridos (o metadispositivos).

Cree dos componentes de un reflejo:

```
nodeA# metainit -s nfsset d101 1 1 /dev/did/dsk/d6s2
```

```
nodeA# metainit -s nfsset d102 1 1 /dev/did/dsk/d7s2
```

Cree el reflejo con uno de los componentes:

```
nodeA# metainit -s nfsset d100 -m d101
```

Conecte el otro componente al reflejo y deje que se sincronice:

```
nodeA# metattach -s nfsset d100 d102
```

Cree particiones de software a partir del reflejo, siguiendo estos ejemplos:

- *d200*: los datos NFS (volumen maestro).


```
nodeA# metainit -s nfsset d200 -p d100 50G
```
- *d201*: el volumen de copia puntual para los datos NFS.


```
nodeA# metainit -s nfsset d201 -p d100 50G
```
- *d202*: el volumen de mapa de bits puntual.

```
nodeA# metainit -s nfsset d202 -p d100 10M
```

- *d203*: el volumen de mapa de bits de sombra remoto.

```
nodeA# metainit -s nfsset d203 -p d100 10M
```

- *d204*: el volumen para la información de configuración de SUNW.NFS de Solaris Cluster.

```
nodeA# metainit -s nfsset d204 -p d100 100M
```

6 Cree sistemas de archivos para los datos NFS y el volumen de configuración.

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d200
```

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d204
```

Pasos siguientes Vaya a “[Configuración de un grupo de dispositivos en el clúster secundario](#)” en la página 266.

▼ 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 265.

1 Obtenga acceso al nodo nodeC como superusuario o asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify`.

2 Cree un metaconjunto para colocar los datos NFS y la replicación asociada.

```
nodeC# metaset -s nfsset a -h nodeC
```

3 Agregue discos al metaconjunto.

En el siguiente ejemplo, se asume que los números DID del disco son diferentes.

```
nodeC# metaset -s nfsset -a /dev/did/dsk/d3 /dev/did/dsk/d4
```

Nota – Los mediadores no son requeridos en un único clúster de nodo.

4 Cree los volúmenes requeridos (o metadispositivos).

Cree dos componentes de un reflejo:

```
nodeC# metainit -s nfsset d101 1 1 /dev/did/dsk/d3s2
```

```
nodeC# metainit -s nfsset d102 1 1 /dev/did/dsk/d4s2
```

Cree el reflejo con uno de los componentes:

```
nodeC# metainit -s nfsset d100 -m d101
```

Conecte el otro componente al reflejo y deje que se sincronice:

```
metattach -s nfsset d100 d102
```

Cree particiones de software a partir del reflejo, siguiendo estos ejemplos:

- *d200*: el volumen maestro de los datos NFS.

```
nodeC# metainit -s nfsset d200 -p d100 50G
```

- *d201*: el volumen de copia puntual para los datos NFS.

```
nodeC# metainit -s nfsset d201 -p d100 50G
```

- *d202*: el volumen de mapa de bits puntual.

```
nodeC# metainit -s nfsset d202 -p d100 10M
```

- *d203*: el volumen de mapa de bits de sombra remoto.

```
nodeC# metainit -s nfsset d203 -p d100 10M
```

- *d204*: el volumen para la información de configuración de SUNW.NFS de Solaris Cluster.

```
nodeC# metainit -s nfsset d204 -p d100 100M
```

5 Cree sistemas de archivos para los datos NFS y el volumen de configuración.

```
nodeC# yes | newfs /dev/md/nfsset/rdisk/d200
```

```
nodeC# yes | newfs /dev/md/nfsset/rdisk/d204
```

Pasos siguientes Vaya a [“Configuración de sistemas de archivos en el clúster primario para la aplicación NFS” en la página 267.](#)

▼ 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 266.](#)

- 1 **Conviértase en superusuario o asuma un rol que proporcione 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 `nodeA` y `nodeB`, configure el volumen maestro para que *no* se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` de `nodeA` y `nodeB`. El texto debe estar en una sola línea.

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \
/global/mountpoint ufs 3 no global,logging
```

- 4 **En `nodeA` y `nodeB`, cree un punto de montaje para el metadispositivo `d204`.**

El ejemplo siguiente crea el punto de montaje `/global/etc`.

```
nodeA# mkdir /global/etc
```

- 5 En **nodeA** y **nodeB**, configure el metadispositivo **d204** 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/md/nfsset/dsk/d204 /dev/md/nfsset/rdisk/d204 \
/global/etc ufs 3 yes global,logging
```

- 6 Monte el metadispositivo **d204** en **nodeA**.

```
nodeA# mount /global/etc
```

- 7 Cree la información y los archivos de configuración para el servicio de datos Oracle Solaris Cluster HA para NFS.

- 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 268.](#)

▼ 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 267.](#)

- 1 Conviértase en superusuario o asuma un rol que proporcione 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/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \
/global/mountpoint ufs 3 yes global,logging
```

- 4 Monte el metadispositivo `d204` en `nodeA`.

```
nodeC# mount /global/etc
```

- 5 Cree la información y los archivos de configuración para el servicio de datos Oracle Solaris Cluster HA para NFS.

- a. Cree un directorio con el nombre `/global/etc/SUNW.nfs` en el nodo `nodeA`.

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

- b. Cree el archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en el nodo `nodeA`.

```
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 `nodeA`.

```
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 269.](#)

▼ 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 268.](#)

- 1 Obtenga acceso al nodo `nodeA` como superusuario o asuma un rol que proporcione las autorizaciones 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 replicaciones. En este nombre, `devgrp` especifica el nombre del grupo de dispositivos.

4 Agregue un recurso SUNW.HASStoragePlus al grupo de recursos de replicación.

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HASStoragePlus \  
-p GlobalDevicePaths=nfsset \  
-p AffinityOn=True \  
devgrp-stor
```

`-g` Especifica el grupo de recursos en el que se agrega el recurso.

`-p GlobalDevicePaths=` Especifica el grupo de dispositivos del que depende el software StorageTek Availability Suite.

`-p AffinityOn=True` Especifica que el recurso SUNW.HASStoragePlus debe realizar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos del clúster definidos por `-p 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.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HASStoragePlus(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 -emM -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 271.](#)

▼ 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 269.

- 1 Obtenga acceso al nodo `nodeC` como superusuario o asuma un rol que proporcione las autorizaciones de `RBAC solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.

- 2 Registre `SUNW.HASStoragePlus` como tipo de recurso.

```
nodeC# clresource_type register SUNW.HASStoragePlus
```

- 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.HASStoragePlus` al grupo de recursos de replicaciones.

```
nodeC# clresource create \
-t SUNW.HASStoragePlus \
-p GlobalDevicePaths=nfsset \
-p AffinityOn=True \
devgrp-stor
```

`create` Crea el recurso.

`-t` Especifica el tipo de recurso.

`-p GlobalDevicePaths=` Especifica el grupo de dispositivos en el que se basa el software StorageTek Availability Suite.

`-p AffinityOn=True` Especifica que el recurso `SUNW.HASStoragePlus` debe realizar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos del clúster definidos por `-p 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 `HASStoragePlus` para el grupo de recursos de replicaciones.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HASStoragePlus(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 secundario es `lhost-reprg-sec`.

6 Habilite los recursos, administre el grupo de recursos y póngalo en línea.

```
nodeC# clresourcegroup online -emM -n nodeC devgrp-stor-rg
```

`online` Lo establece 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 272.](#)

▼ 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 271.](#)

1 Obtenga acceso al nodo `nodeA` como superusuario o asuma un rol que proporcione las autorizaciones 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 servicio NFS.

```
nodeA# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_affinities=+++devgrp-stor-rg \
nfs-rg
```

Pathprefix=/global/etc

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_affinities=+++devgrp-stor-rg

Especifica el grupo de recursos con el que el grupo de recursos de aplicaciones se debe colocar. En este ejemplo, el grupo de recursos de aplicaciones debe colocarse con el grupo de recursos de replications devgrp-stor-rg.

Si el grupo de recursos de replications se conmuta a un nodo principal nuevo, el grupo de recursos de aplicaciones se conmuta automáticamente. Sin embargo, los intentos de conmutar el grupo de recursos de aplicaciones a un nodo principal nuevo se bloquean porque la acción interrumpe el requisito de colocación.

nfs-rg

El nombre del grupo de recursos de aplicaciones.

5 Agregue un recurso de SUNW.HASStoragePlus al grupo de recursos de aplicaciones.

```
nodeA# clresource create -g nfs-rg \
-t SUNW.HASStoragePlus \
-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.HASStoragePlus

Especifica que el recurso es del tipo SUNW.HASStoragePlus.

-p FileSystemMountPoints=/global/punto_montaje

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 para los dispositivos globales y los sistemas de archivos del clúster definidos por -p FileSystemMountPoints. 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 del 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 Establezca en línea el grupo de recursos de aplicaciones.

```
nodeA# clresourcegroup online -emM -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 274.](#)

▼ 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 272.](#)

1 Obtenga acceso al nodo `nodeC` como superusuario o asuma un rol que proporcione las autorizaciones 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.HASStoragePlus no se ha registrado como tipo de recurso, hágalo.

```
nodeC# clresource type register SUNW.HASStoragePlus
```

4 Cree un grupo de recursos de replications para el grupo de dispositivos.

```
nodeC# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_affinities=+++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_affinities=+++devgrp-stor-rg
```

Especifica el grupo de recursos donde el grupo de recursos de aplicaciones debe colocarse. En este ejemplo, el grupo de recursos de aplicaciones debe colocarse con el grupo de recursos de replications devgrp-stor-rg.

Si el grupo de recursos de replications se conmuta a un nodo principal nuevo, el grupo de recursos de aplicaciones se conmuta automáticamente. Sin embargo, los intentos de conmutar el grupo de recursos de aplicaciones a un nodo principal nuevo se bloquean porque la acción interrumpe el requisito de colocación.

```
nfs-rg
```

El nombre del grupo de recursos de aplicaciones.

5 Agregue un recurso de SUNW.HASStoragePlus al grupo de recursos de aplicaciones.

```
nodeC# clresource create -g nfs-rg \
-t SUNW.HASStoragePlus \
-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.HASStoragePlus
```

Especifica que el recurso es del tipo SUNW.HASStoragePlus.

-p

Especifica una propiedad del recurso.

`FileSystemMountPoints=/global/punto_montaje`

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 para los dispositivos globales y los sistemas de archivos del clúster definidos por -p

`FileSystemMountPoints=`. 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.

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 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 276](#).

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 StorageTek Availability Suite. Para obtener más información sobre estos comandos, consulte la documentación de StorageTek Availability.

Esta sección incluye los procedimientos siguientes:

- [“Habilitación de la replicación en el clúster primario” en la página 277](#)
- [“Habilitación de la replicación en el clúster secundario” en la página 278](#)

▼ **Habilitación de la replicación en el clúster primario**

- 1 **Acceda al nodo nodeA como superusuario o asuma un rol que proporcione la autorización de RBAC `solaris.cluster.read`.**

- 2 **Purgue todas las transacciones.**

```
nodeA# lockfs -a -f
```

- 3 **Confirme que los nombres de host 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 permite realizar la replicación del clúster principal al secundario. Este paso permite realizar la replicación del volumen maestro (d200) en el clúster principal al volumen maestro (d200) en el clúster secundario. Además, este paso permite la replicación en el mapa de bits de reflejo remoto de d203.

- Si los clústeres primario y secundario no están sincronizados, ejecute este comando para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

- Si el clúster primario y el secundario están sincronizados, ejecute este comando para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -E lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

- 5 **Habilite la sincronización automática.**

Ejecute este comando para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -a on lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

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 puntual.

Utilice el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/iadm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeA# /usr/sbin/iadm -w \
/dev/md/nfsset/rdisk/d201
```

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 d200, el volumen sombra es d201 y el volumen de mapa de bits puntual es d203.

8 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

Este paso asocia la instantánea de un momento determinado con el conjunto duplicado remoto de volúmenes. El software 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 278.](#)

**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 277.](#)

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 StorageTek Availability Suite:

```
nodeC# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

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 StorageTek Availability Suite.

4 Habilite la instantánea de un momento determinado independiente.

Utilice el comando siguiente para el software StorageTek Availability Suite:

```
nodeC# /usr/sbin/iadm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeC# /usr/sbin/iadm -w \
/dev/md/nfsset/rdisk/d201
```

5 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software StorageTek Availability Suite:

```
nodeC# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

Pasos siguientes Vaya a [“Ejemplo de cómo efectuar una replicación de datos” en la página 279.](#)

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 `iadm` del software StorageTek Availability Suite. Si desea más información sobre estos comandos, consulte la documentación de StorageTek Availability Suite.

Esta sección incluye los procedimientos siguientes:

- [“Replicación por duplicación remota” en la página 280](#)
- [“Instantánea de un momento determinado” en la página 281](#)

- “Procedimiento para comprobar la configuración de la replicación” en la página 282

▼ 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 d200 y el volumen del mapa de bits de reflejo remoto es d203.

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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -m lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

6 Espere hasta que la termine replicación y los volúmenes se sincronicen.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

7 Confirme que el clúster esté en modo de replicación.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

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 StorageTek Availability Suite, actualiza el volumen secundario.

Pasos siguientes Vaya a [“Instantánea de un momento determinado” en la página 281.](#)

▼ 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 d200, el volumen de mapa de bits es d203 y el volumen sombra es d201.

Antes de empezar Realice el procedimiento [“Replicación por duplicación remota” en la página 280.](#)

1 Obtenga acceso al nodo nodeA como superusuario o asuma un rol que proporcione las autorizaciones de RBAC `solaris.cluster.modify` y `solaris.cluster.admin`.

2 Deshabilite el recurso que se esté ejecutando en nodeA.

```
nodeA# clresource disable nfs-rs
```

3 Cambie el clúster primario a modo de registro.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/iidm -u s /dev/md/nfsset/rdisk/d201
nodeA# /usr/sbin/iidm -w /dev/md/nfsset/rdisk/d201
```

5 Sincronice el volumen sombra del clúster secundario con el volumen maestro del clúster secundario.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeC# /usr/sbin/iiadm -u s /dev/md/nfsset/rdisk/d201
nodeC# /usr/sbin/iiadm -w /dev/md/nfsset/rdisk/d201
```

6 Reinicie la aplicación en el nodo nodeA.

```
nodeA# clresource enable nfs-rs
```

7 Vuelva a sincronizar el volumen secundario con el volumen principal.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

Pasos siguientes Vaya a [“Procedimiento para comprobar la configuración de la replicación”](#) en la página 282.

▼ Procedimiento para comprobar la configuración de la replicación

Antes de empezar

Complete el procedimiento [“Instantánea de un momento determinado”](#) en la página 281.

- 1 Obtenga acceso al nodo nodeA y al nodo nodeC como superusuario o asuma un rol que proporcione 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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

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 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 StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
```

```
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

4 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 volumen principal en el directorio de aplicaciones y visualice el directorio montado.

a. Monte el volumen principal en el directorio de aplicaciones.

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

b. Visualice el directorio montado.

```
client-machine# ls /dir
```

6 Desmonte el volumen principal del directorio de aplicaciones.

a. Desmonte el volumen principal del directorio de aplicaciones.

```
client-machine# umount /dir
```

b. Ponga el grupo de recursos de aplicaciones fuera de línea en el clúster primario.

```
nodeA# clresource disable -g nfs-rg +
nodeA# clresourcegroup offline nfs-rg
```

c. Cambie el clúster primario a modo de registro.

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

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. Confirme que el sistema de archivos se adecue para ser montado en el clúster secundario.

```
nodeC# fsck -y /dev/md/nfsset/rdisk/d200
```

- f. Establezca la aplicación en el estado gestionado y establézcala en línea en el clúster secundario.**

```
nodeC# clresourcegroup online -emM nodeC nfs-rg
```

- g. Obtenga acceso al equipo cliente como superusuario.**

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

- h. Monte el directorio de aplicaciones creado en el [Paso 4](#) para el directorio de aplicaciones en el volumen secundario.**

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

- i. 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 volumen principal al directorio de aplicaciones montado.**

- a. Establezca sin conexión el grupo de recursos de aplicaciones en el volumen secundario.**

```
nodeC# clresource disable -g nfs-rg +
nodeC# clresourcegroup offline nfs-rg
```

- b. Asegúrese de que el volumen global se desmonte del volumen secundario.**

```
nodeC# umount /global/mountpoint
```

- c. Establezca el grupo de recursos de aplicaciones en estado gestionado y establézcalo en línea en el clúster principal.**

```
nodeA# clresourcegroup online -emM nodeA nfs-rg
```

- d. Cambie el volumen principal al modo de replicación.**

Ejecute el comando siguiente para el software StorageTek Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

Quando se escribe en el volumen principal, el software StorageTek Availability Suite, actualiza el volumen secundario.

Véase también [“Ejemplo de cómo gestionar una recuperación” en la página 285](#)

Ejemplo de cómo gestionar una recuperación

En esta sección, se describe cómo actualizar las entradas DNS. Para obtener más información, consulte [“Directrices para la gestión de una recuperación” en la página 259](#).

En esta sección, se incluye el siguiente procedimiento:

- [“Actualización de la entrada de DNS” en la página 285](#)

▼ 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](#).

1 Inicie el comando `nsupdate`.

Para obtener más información, consulte la página del comando `man nsupdate(1M)`.

2 Elimine en ambos clústeres 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
```

ipaddress1rev Dirección IP del clúster primario en orden inverso.

ipaddress2rev 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 clústeres.

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

- actualización
 - consejos, 243
 - descripción general, 239
 - zonas de conmutación por error de tipo de marca solaris, 240
- actualizaciones de software, descripción general, 239
- actualizar espacio de nombre global, 88
- actualizar manualmente información de DID, 127
- adaptadores, transporte, 167
- adaptadores de transporte, adición, 27, 31–37, 161–175
- adición
 - cables, adaptadores y conmutadores de transporte, 27, 31–37, 161–175
 - dispositivos del quórum, 135
 - dispositivos del quórum de servidor de quórum, 138
 - dispositivos del quórum NAS de Sun ZFS Storage Appliance, 137
 - hosts de SNMP, 215
 - roles (RBAC), 51
 - usuarios de SNMP, 216
- administración
 - clústeres de zona, 16, 222
 - clústeres globales, 16
 - configuración de clúster global, 191–233
 - interconexiones de clústeres y redes públicas, 161–178
 - IPMP, 161–178
 - sistema de archivos de clúster, 85
- administración de energía, 191
- administrar, quórum, 131–153
- agregación
 - nodos a un clúster de zona, 180
 - nodos a un clúster global, 180
- agregar
 - dispositivos de quórum de disco compartido con conexión directa, 135
 - funciones personalizadas (RBAC), 52
 - grupo de dispositivos, 93, 96
 - grupos de dispositivos de Solaris Volume Manager, 95
 - grupos de dispositivos de ZFS, 96
 - nodos, 179–182
 - sistema de archivos de clúster, 117–120
- almacenamiento SATA, admitido como dispositivo del quórum, 134
- anular registro, grupos de dispositivos de Solaris Volume Manager, 98
- anular supervisión, rutas de disco, 126
- aplicaciones de migración tras error para replicación de datos
 - directrices
 - grupos de recursos, 257
- aplicaciones escalables para replicación de datos, 258–259
- archivo `/etc/vfstab`, 44
 - agregar puntos de montaje, 119
 - verificación de configuración, 120
- archivo `/var/adm/messages`, 77
- archivo `lofi`, desinstalación, 211
- archivo `md.tab`, 20
- archivo `ntp.conf.sc`, 200

archivo vfstab

- agregar puntos de montaje, 119
- verificación de configuración, 120

archivos

- /etc/vfstab, 44
- md.conf, 93
- md.tab, 20
- ntp.conf.sc, 200

arrancar

- clúster de zona, 55–77
- clúster global, 55–77
- modo que no sea de clúster, 74
- nodos, 65–76
- nodos de clúster de zona, 65–76
- nodos de clúster global, 65–76

arrancar en un modo que no sea de clúster, 74

atributos, *Ver* propiedades

Availability Suite, uso para replicación de datos, 251

ayuda, 14

B

búsqueda

- identificadores de nodo para un clúster de zona, 194
- identificadores de nodo para un clúster global, 194

C

cables, transporte, 167

cables de transporte

- adición, 27, 31–37, 161–175
- deshabilitación, 171
- habilitación, 169

cambiar, nombre de clúster, 192–193

cambio

- nombres de host privados, 199
- número de puerto, mediante Common Agent Container, 221
- protocolo de MIB de eventos de SNMP, 214

cambio de nombre de nodos

- en un clúster de zona, 202
- en un clúster global, 202

cerrar

- clúster de zona, 55–77
- clúster global, 55–77
- nodos, 65–76
- nodos de clúster de zona, 65–76
- nodos de clúster global, 65–76

clúster

- autenticación de nodos, 194
- cambiar nombre, 192–193
- configuración de hora del día, 196
- copia de seguridad, 20
- hacer copia de seguridad, 245–248
- restauración de archivos, 248

cluster check

- comando
- cambia a, 38

clúster de zona

- administración, 191–233
- apagar, 55–77
- arrancar, 55–77
- clon, 222
- definición, 16
- desplazamiento de ruta de zona, 222
- eliminación de un sistema de archivos, 222
- estado de componente, 27
- montajes directos admitidos, 224–226
- preparación para aplicaciones, 222
- rearrancar, 61
- validación de configuración, 38
- visualización de información de configuración, 37

clúster global

- administración, 191–233
- arrancar, 55–77
- cerrar, 55–77
- definición, 16
- eliminar nodos, 184
- estado de componente, 27
- rearrancar, 61
- validación de configuración, 38
- visualización de información de configuración, 31–37

clzonecluster

- boot, 59–61
- descripción, 23

clzonecluster (Continuación)

- halt, 55–64

- comando /usr/cluster/bin/cluster check, comprobación de archivo vfstab, 120

- comando boot, 59–61

- comando cconsole, Ver comando pconsole

- comando claccess, 18

- comando cldevice, 18

- comando cldevicegroup, 18

- comando clinterconnect, 18

- comando clnasdevice, 18

- comando clnode, 218, 219–220

- comando clnode check, 18

- comando clquorum, 18

- comando clreslogicalhostname, 18

- comando clresource, 18

- eliminación de recursos y grupos de recursos, 223

- comando clresourcegroup, 18, 219–220

- comando clresourcetype, 18

- comando clressharedaddress, 18

- comando clsnmp host, 18

- comando clsnmpmib, 18

- comando clsnmpuser, 18

- comando cltelemattribute, 18

- comando cluster check, 18

- comprobación de archivo vfstab, 120

- comando cluster shutdown, 55–64

- comando clzonecluster, 18

- comando metaset, 83–85

- comando pconsole, 22

- comando showrev -p, 24

- comandos

- boot, 59–61

- cconsole, 22

- claccess, 18

- cldevice, 18

- cldevicegroup, 18

- clinterconnect, 18

- clnasdevice, 18

- clnode check, 18

- clquorum, 18

- clreslogicalhostname, 18

- clresource, 18

- clresourcegroup, 18

comandos (Continuación)

- clresourcetype, 18

- clressharedaddress, 18

- clsetup, 18

- clsnmp host, 18

- clsnmpmib, 18

- clsnmpuser, 18

- cltelemetryattribute, 18

- cluster check, 18, 21, 38, 44

- cluster scs shutdown, 55–64

- clzonecluster, 18, 55–64

- clzonecluster boot, 59–61

- clzonecluster verify, 38

- metaset, 83–85

- Common Agent Container, cambio de número de puerto, 221

- comprobación, puntos de montaje globales, 44

- comprobar

- configuración de replicación de datos, 282–284

- puntos de montaje globales, 123

- conexiones seguras con consolas del clúster, 23

- configuración, replicación de datos, 251–285

- configuración de hora del clúster, 196

- configuración de límites de carga, en nodos, 219–220

- configurar, funciones (RBAC), 49

- conmutación de afinidad, configuración para replicación de datos, 270

- conmutación para replicación de datos

- conmutación de afinidad, 256

- realización, 285

- conmutación por error de aplicaciones para replicación de datos

- AffinityOn property*, 256

- directrices

- gestión de recuperación, 259

- gestión, 285

- GlobalDevicePaths*, 256

- ZPoolsSearchDir*, 256

- conmutadores, transporte, 167

- conmutadores de transporte, adición, 27, 31–37, 161–175

- conmutar, nodo primario de grupo de dispositivos, 109

- conmutar a nodo primario de grupo de dispositivos, 109
- consola de administración, 20
- consolas, conexión a, 22
- control de acceso basado en funciones, *Ver* RBAC
- convención de asignación de nombre, grupos de recursos de replications, 256
- convención de denominación, dispositivos de discos básicos, 119
- copia de seguridad, clúster, 20

D

- deshabilitación de cables de transporte, 171
- desinstalación
 - archivo de dispositivo Iofi, 211
 - paquetes, 243
 - software de Oracle Solaris Cluster, 208
- detener
 - clúster de zona, 61
 - clúster global, 61
 - nodos, 65–76
 - nodos de clúster de zona, 65–76
 - nodos de clúster global, 65–76
- disco de SCSI compartido, admitido como dispositivo del quórum, 134
- dispositivos
 - globales, 83–130
 - reconfiguración dinámica, 84–85
- dispositivos de discos básicos, convención de denominación, 119
- dispositivos de quórum
 - agregar
 - dispositivos de quórum de disco compartido con conexión directa, 135
 - eliminar, 133
 - estado de mantenimiento, sacar a un dispositivo, 149
 - quitar el último dispositivo de quórum, 143
 - reconfiguración dinámica de dispositivos, 133
 - reemplazar, 145
 - reparar, 152
- dispositivos de quórum de disco compartido con conexión directa, agregar, 135

- dispositivos de quórum de servidor de quórum, resolución de problemas para quitar dispositivos, 143
- dispositivos del quórum
 - adición, 135
 - dispositivos del quórum de servidor de quórum, 138
 - dispositivos del quórum NAS de Sun ZFS Storage Appliance, 137
- cambio del tiempo de espera predeterminado, 153
- eliminación, 142
- enumeración de configuración, 151
- estado de mantenimiento, poner un dispositivo en, 148
- modificación de listas de nodos, 146
- dispositivos del quórum de servidor de quórum
 - adición, 138
 - requisitos para la instalación, 138
- Domain Name System (DNS), actualizar en replicación de datos, 285
- DR, *Ver* reconfiguración dinámica
- duplicación remota, realizar, 280–281
- duplicaciones, copia de seguridad en línea, 245

E

- ejemplos
 - crear un sistema de archivos de clúster, 120
 - ejecutar una comprobación de validación funcional, 42–43
 - listado de comprobaciones de validación interactivas, 42
- eliminación
 - cables, adaptadores y conmutadores de transporte, 167
 - de un clúster de zona, 184
 - dispositivos del quórum, 142
 - hosts de SNMP, 215
 - nodos, 182
 - recursos y grupos de recursos de un clúster de zona, 223
 - usuarios de SNMP, 217
- eliminar
 - dispositivos de quórum, 133

eliminar (*Continuación*)

- grupos de dispositivos de Solaris Volume Manager, 98
- matrices de almacenamiento, 187
- nodos, 184
- nodos de todos los grupos de dispositivos, 98
- enumeración, configuración de quórum, 151
- enumerar, configuración de grupo de dispositivos, 107
- espacio de nombre
 - global, 83–85, 88
 - migrar, 90
- espacio de nombre de dispositivos globales, migrar, 90
- estado
 - componente de clúster de zona, 27
 - componente de clúster global, 27
- estado de mantenimiento
 - nodos, 204
 - poner un dispositivo del quórum en, 148
 - sacar un dispositivo de quórum de, 149

F

- fuera de servicio, dispositivos del quórum, 148
- función
 - agregar funciones personalizadas, 52
 - configurar, 49

G

- global
 - espacio de nombre, 83–85
 - puntos de montaje, comprobar, 123
- GlobalDevicePaths*, propiedad de extensión para replicación de datos, 256
- globales
 - dispositivos, 83–130
 - configurar permisos, 84
 - puntos de montaje, comprobación, 44
- grupo de dispositivos de discos básicos, agregar, 96
- grupos de dispositivos
 - agregar, 95
 - configuración para replicación de datos, 264

grupos de dispositivos (*Continuación*)

- disco básico
 - agregar, 96
 - eliminar y anular registro, 98
 - enumerar configuración, 107
 - estado de mantenimiento, 110
 - información general de administración, 86
 - modificar propiedades, 103
 - propiedad primaria, 103
- SVM
 - agregar, 93
- grupos de recursos
 - replicación de datos
 - configuración, 255
 - directrices para configurar, 255
 - rol en conmutación por error, 255
- grupos de recursos de aplicaciones
 - configurar para replicación de datos, 272–274
 - directrices, 256
- grupos de recursos de direcciones compartidas para replicación de datos, 258

H

- habilitación de cables de transporte, 169
- habilitación y deshabilitación de una MIB de eventos de SNMP, 213
- hacer copia de seguridad, clúster, 245–248
- herramienta de administración de la línea de comandos, 18
- herramienta User Accounts, descripción, 53
- hosts
 - adición y eliminación de SNMP, 215

I

- imprimir, rutas de disco erróneas, 126–127
- información de DID, actualizar manualmente, 127
- información de versión, 24
- información general, quórum, 131–153
- iniciar
 - clúster de zona, 59–61
 - clúster global, 59–61

iniciar (*Continuación*)

- nodos, 65–76
- nodos de clúster global, 65–76
- inicio de sesión, remoto, 22
- inicio de sesión remoto, 22
- instantánea, puntual, 253
- instantánea de un momento determinado, ejecutar, 281–282
- instantánea puntual, definición, 253
- interconexiones de clústeres
 - administración, 161–178
 - reconfiguración dinámica, 163
- IPMP
 - administración, 176
 - comprobación de estado, 30

K

- /kernel/drv/, archivo `md.conf`, 93

L

- límites de carga
 - configuración en nodos, 218, 219–220
 - propiedad `concentrate_load`, 218
 - propiedad `preemption_mode`, 218

M

- mantenimiento, dispositivos del quórum, 148
- mapa de bits
 - instantánea puntual, 253
 - replicación por duplicación remota, 252
- matrices de almacenamiento, eliminar, 187
- mensajes de error
 - archivo `/var/adm/messages`, 77
 - eliminar nodos, 189–190
- MIB
 - cambio de protocolo de eventos de SNMP, 214
 - habilitación y deshabilitación de eventos de SNMP, 213

MIB de eventos

- cambio de protocolo de SNMP, 214
- habilitación y deshabilitación de SNMP, 213
- migrar, espacio de nombre de dispositivos globales, 90
- modificación, listas de nodos de dispositivos del quórum, 146
- modificar
 - nodos primarios, 109
 - propiedad `numsecondaries`, 105
 - propiedades, 103
 - usuarios (RBAC), 53
- montaje directo, exportación de un sistema de archivos a un clúster de zona, 224–226
- montaje en bucle de retorno, exportación de un sistema de archivos a un clúster de zona, 224–226

N

- NAS de Sun, admitido como dispositivo del quórum, 134
- Network File System (NFS), configurar sistemas de archivos de aplicaciones para replicación de datos, 267–268
- nodos
 - agregar, 179–182
 - arrancar, 65–76
 - autenticación, 194
 - búsqueda de ID, 194
 - cambio de nombre en un clúster de zona, 202
 - cambio de nombre en un clúster global, 202
 - cerrar, 65–76
 - cómo poner un nodo en estado de mantenimiento, 204
 - conexión a, 22
 - configuración de límites de carga, 219–220
 - eliminación de un clúster de zona, 184
 - eliminar
 - mensajes de error, 189–190
 - eliminar de grupos de dispositivos, 98
 - eliminar nodos de un clúster global, 184
 - primarios, 84–85, 103
 - secundarios, 103
- nodos de clúster de zona
 - cerrar, 65–76

nodos de clúster de zona (*Continuación*)
 especificando dirección IP y NIC, 179–182
 iniciar, 65–76
 reiniciar, 71–74

nodos de clúster global
 arrancar, 65–76
 cerrar, 65–76
 reiniciar, 71–74

nodos de votación de clúster global
 recursos compartidos de CPU, 236
 sistema de archivos de clúster de administración, 85

nombres de host privados, cambio, 199

número de puerto, cambio mediante Common Agent
 Container, 221

O

opciones de montaje para sistemas de archivos de
 clúster, requisitos, 119

P

paquetes, desinstalación, 243
 pconsole, conexiones seguras, 23
 perfiles, derechos de RBAC, 50–51
 perfiles de derechos, RBAC, 50–51
 permisos, dispositivo global, 84
 planificador por reparto equitativo, configuración de
 recursos compartidos de CPU, 236
 PROM OpenBoot (OBP), 198
 propiedad *AffinityOn*, propiedad de extensión para
 replicación de datos, 256
 propiedad failback, 103
 propiedad numsecondaries, 105
 propiedad primaria de grupos de dispositivos, 103
 propiedades
 failback, 103
 numsecondaries, 105
 preferenced, 103
 propiedades de extensión para replicación de datos
 recurso de aplicación, 273, 275
 recurso de replicación, 271

puntos de montaje
 globales, 44
 modificación del archivo `/etc/vfstab`, 119

Q

quitar
 sistema de archivos de clúster, 121–123
 último dispositivo de quórum, 143
 quórum
 administración, 131–153
 información general, 131–153

R

RBAC, 49–53
 para nodos de votación de clúster global, 50
 perfiles de derechos (descripción), 50–51
 tareas
 adición de roles, 51
 agregar funciones personalizadas, 52
 configurar, 49
 modificar usuarios, 53
 usar, 49
 realizar copia de seguridad, duplicaciones en línea, 245
 rearrancar, clúster de zona, 61
 reconfiguración dinámica, 84–85
 dispositivos de quórum, 133
 interconexiones de clústeres, 163
 interfaces de red pública, 177
 recurso de nombre de host lógico, rol en recuperación
 de replicación de datos, 256
 recursos
 eliminación, 223
 visualización de información de
 configuración, 26–27
 recursos compartidos de CPU
 configurar, 235
 controlar, 235
 nodos de votación de clúster global, 236
 red pública
 administración, 161–178
 reconfiguración dinámica, 177

reemplazar dispositivos de quórum, 145

reiniciar

- clúster global, 61

- nodos de clúster de zona, 71–74

- nodos de clúster global, 71–74

reparación completa de archivo

- /var/adm/messages, 77

reparar, dispositivo de quórum, 152

replicación, *Ver* replicación de datos

replicación de datos, 79–81

- actualizar una entrada DNS, 285

- asincrónica, 253

- basada en host, 80–81, 251–285

- comprobar configuración, 282–284

- configuración

 - conmutación de afinidad, 256, 270

 - grupos de dispositivos, 264

- configuración de ejemplo, 260

- configurar

 - grupos de recursos de aplicaciones

 - NFS, 272–274

 - sistemas de archivos para una aplicación

 - NFS, 267–268

- definición, 80–81

- directrices

 - configurar grupos de recursos, 255

 - gestión de conmutación, 259

 - gestión de recuperación, 259

- duplicación remota, 252, 280–281

- ejemplo, 279–284

- gestión de una recuperación, 285

- grupos de recursos

 - aplicación, 256

 - aplicaciones de migración tras error, 257

 - aplicaciones escalables, 258–259

 - configuración, 255

 - convención de asignación de nombre, 256

 - crear, 269–270

 - dirección compartida, 258

- habilitación, 276–279

- hardware y software necesarios, 262

- instantánea de un momento determinado, 281–282

- instantánea puntual, 253

- presentación, 252

replicación de datos (*Continuación*)

- sincrónica, 253

replicación de datos asincrónica, 253

replicación de datos basada en host

- definición, 80–81

- ejemplo, 251–285

replicación de datos sincrónica, 253

replicación por duplicación remota, definición, 252

restauración

- archivos de clúster, 248

- sistema de archivos raíz, 248

rol, adición de roles, 51

ruta de disco

- anular supervisión, 126

- resolver error de estado, 127

- supervisar, 83–130

 - imprimir rutas de disco erróneas, 126–127

ruta de zona, desplazamiento, 222

rutas de disco compartido

- habilitar rearranque automático, 129–130

- inhabilitar reinicio automático, 130

- supervisar, 123–130

S

SATA, 135

secundarios

- establecer número, 105

- número predeterminado, 103

servicio de asistencia técnica, 14

servidor de quórum de Oracle Solaris Cluster, admitido

- como dispositivo del quórum, 134

servidores de quórum, *Ver* dispositivos del quórum de

- servidor de quórum

sistema de archivos

- aplicación NFS

 - configurar para replicación de datos, 267–268

- eliminación desde un clúster de zona, 222

- restauración de sistema de archivos raíz

 - descripción, 248

sistema de archivos de clúster, 83–130

- administration, 85

- agregar, 117–120

- nodos de votación de clúster global, 85

sistema de archivos de clúster (*Continuación*)

quitar, 121–123

sistema de nombres de dominio (DNS), directrices para actualizar, 259

sistema operativo Oracle Solaris

comando `svcadm`, 199

control CPU, 235

definición de clúster de zona, 15

definición de clúster global, 15

instrucciones especiales para iniciar nodos, 69–71

instrucciones especiales para reiniciar un nodo, 71–74

replicación basada en host, 80–81

tareas administrativas para un clúster global, 16

sistema operativo Solaris, *Ver* sistema operativo Oracle Solaris

sistemas de archivos de clúster

opciones de montaje, 119

verificación de configuración, 120

sistemas de archivos globales, *Ver* sistemas de archivos de clúster

SNMP

adición de usuarios, 216

cambio de protocolo, 214

deshabilitación de hosts, 215

eliminación de usuarios, 217

habilitación de hosts, 215

habilitación y deshabilitación de una MIB de eventos, 213

Solaris Volume Manager, nombres de dispositivos de discos básicos, 119

ssh, 23

StorageTek Availability Suite, uso para replicación de datos, 251

Sun ZFS Storage Appliance

adición como dispositivo del quórum, 137

admitido como dispositivo del quórum, 134

supervisor

rutas de disco, 124–126

rutas de disco compartido, 129–130

switchback, directrices para realizar en replicación de datos, 260

T

tiempo de espera, cambio del valor predeterminado de un dispositivo del quórum, 153

tipos de dispositivo del quórum, lista de tipos admitidos, 134

tipos de dispositivo del quórum admitidos, 134

tolerancia a desastres, definición, 252

U

último dispositivo de quórum, quitar, 143

usar, funciones (RBAC), 49

`/usr/cluster/bin/clresource`, eliminación de grupos de recursos, 223

usuarios

adición de SNMP, 216

eliminación de SNMP, 217

modificar propiedades, 53

utilidad `clsetup`, 18, 23

V

validación

configuración de clúster de zona, 38

configuración de clúster global, 38

verificación, configuración de `vfstab`, 120

Z

ZFS

agregar grupos de dispositivos, 96

eliminación de un sistema de archivos, 224–226
replicación, 96

restricciones para sistemas de archivos root, 85–86

ZFS Storage Appliance, *Ver* dispositivos del quórum

NAS de Sun ZFS Storage Appliance

`ZPoolsSearchDir`, propiedad de extensión para replicación de datos, 256

