

Guía de administración del sistema de Oracle® Solaris Cluster

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 cluster de zona	16
Restricciones de las funciones del sistema operativo Oracle Solaris	17
Herramientas de administración	18
Interfaz gráfica de usuario	18
Interfaz de línea de comandos	18
Preparación para administrar el cluster	20
Documentación de las configuraciones de hardware de Oracle Solaris Cluster	20
Uso de la consola de administración	20
Copia de seguridad del cluster	21
Procedimiento para empezar a administrar el cluster	21
▼ Inicio de sesión remoto en el cluster	23
▼ Conexión segura a las consolas del cluster	24
▼ Obtención de acceso a las utilidades de configuración del cluster	25
▼ Visualización de la información de parches de Oracle Solaris Cluster	26
▼ Visualización de la información de versión de Oracle Solaris Cluster	27
▼ Visualización de tipos de recursos configurados, grupos de recursos y recursos	30
▼ Comprobación del estado de los componentes del cluster	31
▼ Comprobación del estado de la red pública	34
▼ Visualización de la configuración del cluster	35
▼ Validación de una configuración básica de cluster	43
▼ Comprobación de los puntos de montaje globales	49
▼ Visualización del contenido de los registros de comandos de Oracle Solaris Cluster	50

2	Oracle Solaris Cluster y RBAC	53
	Instalación y uso de RBAC con Oracle Solaris Cluster	53
	Perfiles de derechos de RBAC en Oracle Solaris Cluster	54
	Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster	55
	▼ Creación de un rol con la herramienta de roles administrativos	55
	▼ Creación de una función desde la línea de comandos	57
	Modificación de las propiedades de RBAC de un usuario	59
	▼ Modificación de las propiedades de RBAC de un usuario con la herramienta de cuentas de usuario	59
	▼ Modificación de las propiedades de RBAC de un usuario desde la línea de comandos	60
3	Cierre y arranque de un cluster	61
	Descripción general sobre el cierre y el arranque de un cluster	61
	▼ Cierre de un cluster	63
	▼ Arranque de un cluster	65
	▼ Rearranque de un cluster	67
	Cierre y arranque de un solo nodo de un cluster	71
	▼ Cierre de un nodo	72
	▼ Arranque de un nodo	75
	▼ Rearranque de un nodo	78
	▼ Rearranque de un nodo en un modo que no sea de cluster	82
	Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad	85
	▼ Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad	85
4	Métodos de replicación de datos	87
	Comprensión de la replicación de datos	88
	Métodos admitidos de replicación de datos	89
	Uso de la replicación de datos basada en almacenamiento en un cluster	90
	Requisitos y restricciones del uso de la replicación de datos basada en almacenamiento en un cluster	92
	Problemas de recuperación manual en el uso de la replicación de datos basada en almacenamiento en un cluster	93
	Mejores prácticas para utilizar la replicación de datos basada en almacenamiento	94

5 Administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de cluster	95
Información general sobre la administración de dispositivos globales y el espacio de nombre global	95
Permisos de dispositivos globales en Solaris Volume Manager	96
Reconfiguración dinámica con dispositivos globales	96
Administración de dispositivos replicados basados en el almacenamiento	97
Administración de dispositivos replicados de Hitachi TrueCopy	98
Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility	109
Información general sobre la administración de sistemas de archivos de clusters	122
Restricciones del sistema de archivos de cluster	122
Administración de grupos de dispositivos	123
▼ Actualización del espacio de nombre de dispositivos globales	125
▼ Cómo cambiar el tamaño de un dispositivo <code>lofi</code> que se utiliza para el espacio de nombres de dispositivos globales	126
Migración del espacio de nombre de dispositivos globales	127
▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de <code>lofi</code>	127
▼ Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada	129
Adición y registro de grupos de dispositivos	130
▼ Adición y registro de un grupo de dispositivos (Solaris Volume Manager)	130
▼ Adición y registro de un grupo de dispositivos (disco básico)	132
▼ Adición y registro de un grupo de dispositivos replicado (ZFS)	133
Mantenimiento de grupos de dispositivos	135
Eliminación y anulación del registro de un grupo de dispositivos (Solaris Volume Manager)	135
▼ Eliminación de un nodo de todos los grupos de dispositivos	135
▼ Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)	136
▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos	138
▼ Cambio de propiedades de los grupos de dispositivos	140
▼ Establecimiento del número de secundarios para un grupo de dispositivos	142
▼ Enumeración de la configuración de grupos de dispositivos	144
▼ Conmutación al nodo primario de un grupo de dispositivos	146
▼ Colocación de un grupo de dispositivos en estado de mantenimiento	147
Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento	149

▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento	149
▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento	150
▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento	151
▼ Modificación del protocolo de protección en un solo dispositivo de almacenamiento	152
Administración de sistemas de archivos de cluster	154
▼ Adición de un sistema de archivos de cluster	154
▼ Eliminación de un sistema de archivos de cluster	158
▼ Comprobación de montajes globales en un cluster	160
Administración de la supervisión de rutas de disco	160
▼ Supervisión de una ruta de disco	161
▼ Anulación de la supervisión de una ruta de disco	163
▼ Impresión de rutas de disco erróneas	164
▼ Resolución de un error de estado de ruta de disco	164
▼ Supervisión de rutas de disco desde un archivo	165
▼ Habilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	167
▼ Inhabilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	168
6 Administración de quórum	169
Administración de dispositivos de quórum	169
Reconfiguración dinámica con dispositivos de quórum	171
Adición de un dispositivo de quórum	172
Eliminación o sustitución de un dispositivo de quórum	181
Mantenimiento de dispositivos de quórum	184
Cambio del tiempo de espera predeterminado del quórum	192
Administración de servidores de quórum de Oracle Solaris Cluster	193
Inicio y detención del software del servidor del quórum	193
▼ Inicio de un servidor de quórum	193
▼ Detención de un servidor de quórum	194
Visualización de información sobre el servidor de quórum	195
Limpieza de la información caducada sobre clusters del servidor de quórum	197

7 Administración de interconexiones de clusters y redes públicas	199
Administración de interconexiones de clusters	199
Reconfiguración dinámica con interconexiones de clusters	201
▼ Comprobación del estado de la interconexión de cluster	201
▼ Adición de dispositivos de cable de transporte de cluster, adaptadores o conmutadores de transporte	202
▼ Eliminación de cable de transporte de cluster, adaptadores de transporte y conmutadores de transporte	205
▼ Habilitación de un cable de transporte de cluster	208
▼ Inhabilitación de un cable de transporte de cluster	209
▼ Determinación del número de instancia de un adaptador de transporte	211
▼ Modificación de la dirección de red privada o del intervalo de direcciones de un cluster	212
Administración de redes públicas	214
Administración de grupos de varias rutas de red IP en un cluster	215
Reconfiguración dinámica con interfaces de red pública	216
 8 Adición y eliminación de un nodo	 219
Adición de nodos a un cluster	219
▼ Cómo agregar un nodo a un cluster existente	220
Creación de un nodo sin voto (zona) en un cluster global	222
Eliminación de nodos de un cluster	225
▼ Eliminación de un nodo de un cluster de zona	226
▼ Eliminación de un nodo de la configuración de software del cluster	227
▼ Eliminación de un nodo sin voto (zona) de un cluster global	230
▼ Eliminación de la conectividad entre una matriz y un único nodo, en un cluster con conectividad superior a dos nodos	231
▼ Corrección de mensajes de error	233
 9 Administración del cluster	 235
Información general sobre la administración del cluster	235
▼ Cambio del nombre del cluster	236
▼ Asignación de un ID de nodo a un nombre de nodo	238
▼ Uso de la autenticación del nodo del cluster nuevo	238
▼ Restablecimiento de la hora del día en un cluster	240
▼ SPARC: Visualización de la PROM OpenBoot en un nodo	242

▼ Cómo cambiar un nombre de host privado de nodo	243
▼ Agregación de un nombre de host privado para un nodo sin voto en un cluster global	246
▼ Cambio de nombre de host privado en un nodo sin voto de un cluster global	246
▼ Supresión de un nombre de host privado para un nodo sin voto en un cluster global	248
▼ Cómo cambiar el nombre de un nodo	248
▼ Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes	250
▼ Cómo poner un nodo en estado de mantenimiento	251
▼ Cómo sacar un nodo del estado de mantenimiento	252
▼ Cómo instalar el software de Oracle Solaris Cluster desde un nodo de cluster	255
Resolución de problemas de desinstalación de nodos	257
Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster	259
Configuración de límites de carga	265
Cómo realizar tareas administrativas del cluster de zona	267
▼ Cómo agregar una dirección de red a un cluster de zona	268
▼ Cómo eliminar un cluster de zona	270
▼ Cómo eliminar un sistema de archivos de un cluster de zona	270
▼ Cómo eliminar un dispositivo de almacenamiento de un cluster de zona	273
Resolución de problemas	275
Ejecución de una aplicación fuera del cluster global	275
Restauración de un conjunto de discos dañado	277
10 Configuración del control del uso de la CPU	281
Introducción al control de la CPU	281
Elección de una situación hipotética	281
Planificador por reparto equitativo	282
Configuración del control de la CPU	283
▼ Control del uso de la CPU en el nodo de votación de un cluster global	283
▼ Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores predeterminado	285
▼ Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores dedicado	288
11 Aplicación de parches de software y firmware de Oracle Solaris Cluster	293
Descripción general de la aplicación de parches de Oracle Solaris Cluster	293

Sugerencias de parches de Oracle Solaris Cluster	294
Aplicación de parches de software de Oracle Solaris Cluster	295
▼ Aplicación de un parche de reinicio (nodo)	295
▼ Aplicación de un parche de reinicio (cluster)	300
▼ Aplicación de un parche de no reinicio de Oracle Solaris Cluster	303
▼ Aplicación de parches en modo de un solo usuario a nodos con zonas de conmutación por error	305
Cambio de un parche de Oracle Solaris Cluster	309
12 Copias de seguridad y restauraciones de clusters	311
Copia de seguridad de un cluster	311
▼ Búsqueda de nombres de sistemas de archivos para realizar copias de seguridad	312
▼ Determinación del número de cintas necesario para una copia de seguridad completa ...	312
▼ Realización de una copia de seguridad del sistema de archivos raíz (/)	313
▼ Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)	316
▼ Copias de seguridad de la configuración del cluster	319
Restauración de los archivos del cluster	319
▼ Restauración interactiva de archivos individuales (Solaris Volume Manager)	320
▼ Restauración del sistema de archivos raíz (/) (Solaris Volume Manager)	320
▼ Restauración de un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager	323
13 Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario	329
Descripción general sobre Oracle Solaris Cluster Manager	329
Configuración de Oracle Solaris Cluster Manager	330
Configuración de funciones de RBAC	330
▼ Cómo utilizar Common Agent Container para cambiar los números de puerto para servicios o agentes de gestión	331
▼ Cambio de la dirección del servidor para Oracle Solaris Cluster Manager	332
▼ Regeneración de las claves de seguridad de Common Agent Container	332
Inicio del software de Oracle Solaris Cluster Manager	333
▼ Inicio de Oracle Solaris Cluster Manager	334
A Ejemplo	337
Configuración de replicación de datos basada en host con el software de Availability Suite ...	337

Descripción del software de Availability Suite en un cluster	338
Directrices para la configuración de replicación de datos basada en host entre clusters ..	341
Mapa de tareas: ejemplo de configuración de una replicación de datos	346
Conexión e instalación de clusters	347
Ejemplo de configuración de grupos de dispositivos y de recursos	349
Ejemplo de cómo habilitar la replicación de datos	363
Ejemplo de cómo efectuar una replicación de datos	366
Ejemplo de cómo gestionar una recuperación	371
 Índice	 375

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 las familias de arquitectura de procesadores SPARC y x86: UltraSPARC, SPARC64, AMD64 e Intel 64. En este documento, x86 hace referencia a la familia más amplia de productos compatibles con x86 de 64 bits. La información de este documento se aplica a todas las plataformas a menos que se especifique lo contrario.

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

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

Uso de los comandos de UNIX

Este documento contiene información sobre comandos específicos para la administración de una configuración de Oracle Solaris Cluster. Es posible que este documento no contenga información completa sobre los comandos y los procedimientos básicos de UNIX.

Consulte una o varias de las siguientes fuentes para obtener dicha información:

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

Convenciones tipográficas

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

TABLA P-1 Convenciones tipográficas

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

En la tabla siguiente, se muestran los indicadores de sistema UNIX y de superusuario para los shells que se incluyen en el sistema operativo Oracle Solaris. En los ejemplos de comandos, el indicador del shell indica si el comando se debe ejecutar por un usuario común o un usuario con privilegios.

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
Conceptos	<i>Oracle Solaris Cluster Concepts Guide</i>
Administración e instalación de software	<i>Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual</i> y guías de administración de hardware individual
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>

Acceso a My 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

Póngase en contacto con el proveedor de servicios si tiene algún problema al instalar o utilizar Oracle Solaris Cluster. Proporcione la siguiente información al proveedor de servicios.

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

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

Comando	Función
<code>prtconf -v</code>	Muestra el tamaño de la memoria del sistema y ofrece información sobre los dispositivos periféricos.
<code>psrinfo -v</code>	Muestra información acerca de los procesadores.
<code>showrev -p</code>	Indica los parches instalados.
<code>SPARC: prtdiag -v</code>	Muestra información de diagnóstico del sistema.
<code>/usr/cluster/bin/clnode show-rev -v</code>	Muestra la información de la versión de Oracle Solaris Cluster y de los paquetes.

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

Introducción a la administración de Oracle Solaris Cluster

Este capítulo ofrece la información siguiente sobre la administración de clusters 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 cluster” en la página 20
- “Procedimiento para empezar a administrar el cluster” en la página 21

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

Un cluster global se compone de uno o más nodos de cluster global. Un cluster global puede incluir también zonas no globales de la marca `native` y `cluster` que no son nodos, pero que están configuradas con el servicio de datos HA para Zones. Para obtener información general sobre los clusters de zona, consulte la [Oracle Solaris Cluster Concepts Guide](#).

Un cluster de zona consta de una o más zonas no globales de la marca `cluster`. Para crear un cluster de zona, puede utilizar el comando `clzonecluster` o la utilidad `clsetup`. Puede ejecutar servicios compatibles en el cluster de zona del mismo modo que en un cluster global, con el aislamiento proporcionado por las zonas de Oracle Solaris. Un cluster de zona depende de, y por tanto requiere, un cluster global. Un cluster global no contiene un cluster de zona. Un cluster de zona tiene, como máximo, un nodo cluster de zona en un equipo. Un nodo de cluster de zona sigue funcionando únicamente si también lo hace el nodo del cluster global en la misma máquina. Si falla un nodo de cluster global en un equipo, también fallan todos los nodos de cluster de zona de esa máquina. Para obtener información general sobre los clusters de zona, consulte la [Oracle Solaris Cluster Concepts Guide](#).

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 *Guía de instalación del software de Oracle Solaris Cluster* y en *Oracle Solaris Cluster Concepts Guide*. Si desea obtener instrucciones sobre cómo crear un cluster de zona, consulte “Configuración de un cluster 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 cluster 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. Esas tareas se describen en *Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual*.

En general, puede realizar tareas administrativas de Oracle Solaris Cluster mientras el cluster está en funcionamiento. Si necesita retirar un nodo del cluster o incluso cerrar dicho nodo, puede hacerse mientras el resto de los nodos continúan desarrollando las operaciones del cluster. Salvo que se indique lo contrario, las tareas administrativas de Oracle Solaris Cluster deben efectuarse en el nodo de votación del cluster global. En los procedimientos que necesitan el cierre de todo el cluster, 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 cluster o un nodo del cluster, notifíquelo con antelación a los usuarios.

Trabajo con un cluster de zona

En un cluster 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 cluster de zona donde se ejecute dicho comando. Por ejemplo, utilizar el comando `cluster` en el nodo de votación del cluster global recupera toda la información del cluster global de votación y de todos los clusters de zona. Al utilizar el comando `cluster` en un cluster de zona, se recupera la información de ese mismo cluster de zona.

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

Los clusters de zona admiten la administración delegada de los recursos situados jerárquicamente bajo el control del administrador de grupos de recursos. Así, los administradores de clusters de zona pueden ver las dependencias de dichos clusters de zona que superan los límites de los clusters 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 clusters de zona.

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

- Creación de un cluster de zona: emplee la utilidad `clsetup` para iniciar el asistente de configuración del cluster de zona o utilice el comando `clzonecluster install`. Consulte las instrucciones de “Configuración de un cluster de zona” de *Guía de instalación del software de Oracle Solaris Cluster*.
- Inicio y re arranque de un cluster de zona: consulte el [Capítulo 3, “Cierre y arranque de un cluster”](#).
- Adición de un nodo a un cluster de zona: consulte el [Capítulo 8, “Adición y eliminación de un nodo”](#).
- Eliminación de un nodo de un cluster de zona: consulte “[Eliminación de un nodo de un cluster de zona](#)” en la [página 226](#).
- Visualización de la configuración de un cluster de zona: consulte “[Visualización de la configuración del cluster](#)” en la [página 35](#).
- Validación de la configuración de un cluster de zona: consulte “[Validación de una configuración básica de cluster](#)” en la [página 43](#).
- Detención de un cluster de zona: consulte el [Capítulo 3, “Cierre y arranque de un cluster”](#).

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
pnm	svc:/system/cluster/pnm:default
cl_event	svc:/system/cluster/cl_event:default
cl_eventlog	svc:/system/cluster/cl_eventlog:default
rpc_pmf	svc:/system/cluster/rpc_pmf:default

TABLA 1-1 Servicios de Oracle Solaris Cluster (Continuación)

Servicios de Oracle Solaris Cluster	FMRI
rpc_fed	svc:/system/cluster/rpc_fed:default
rgm	svc:/system/cluster/rgm:default
scdpm	svc:/system/cluster/scdpm:default
cl_ccra	svc:/system/cluster/cl_ccra:default
scsymon_srv	svc:/system/cluster/scsymon_srv:default
spm	svc:/system/cluster/spm:default
cl_svc_cluster_milestone	svc:/system/cluster/cl_svc_cluster_milestone:default
cl_svc_enable	svc:/system/cluster/cl_svc_enable:default
network-multipathing	svc:/system/cluster/network-multipathing

Herramientas de administración

Las tareas administrativas de una configuración de Oracle Solaris Cluster se pueden realizar con una interfaz gráfica de usuario (GUI) o mediante la línea de comandos. En la siguiente sección, se ofrece una descripción general de las herramientas de la GUI y la línea de comandos.

Interfaz gráfica de usuario

El software de Oracle Solaris Cluster admite herramientas de GUI que se pueden usar para llevar a cabo diversas tareas administrativas en el cluster. Una de estas herramientas de GUI es Oracle Solaris Cluster Manager. Consulte el [Capítulo 13, “Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario”](#) para obtener más información y conocer los procedimientos relacionados con la configuración de Oracle Solaris Cluster Manager. Para obtener información específica sobre cómo utilizar Oracle Solaris Cluster Manager, consulte la ayuda en pantalla de la GUI.

Interfaz de línea de comandos

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

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

- Quórum

- Grupos de recursos
- Servicios de datos
- Interconexión de cluster
- Grupos de dispositivos y volúmenes
- Nombres de host privados
- Nuevos nodos
- Clusters de zona
- Otras tareas del cluster

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

<code>ccp(1M)</code>	Inicia el acceso remoto de la consola al cluster.
<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 un cluster de Oracle Solaris Cluster de forma interactiva, y también crea y configura un cluster de zona.

<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 cluster 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 cluster utiliza.

Preparación para administrar el cluster

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

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 cluster. Otra forma de facilitar la administración es etiquetar los cables y las conexiones entre los diversos componentes del cluster.

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

Uso de la consola de administración

Para administrar el cluster activo, puede optar por utilizar una estación de trabajo dedicada o una estación de trabajo conectada a través de una red de gestión como *consola de administración*. Por lo general, las herramientas de la interfaz gráfica de usuario (GUI) y Cluster Control Panel (CCP) se instalan y se ejecutan en la consola de administración. Para obtener más información sobre CCP, consulte [“Inicio de sesión remoto en el cluster” en la página 23](#). Para

obtener instrucciones sobre cómo instalar las herramientas de GUI de Oracle Solaris Cluster Manager y el módulo Cluster Control Panel, consulte [Guía de instalación del software de Oracle Solaris Cluster](#).

La consola de administración no es un nodo del cluster. La consola de administración se utiliza para disponer de acceso remoto a los nodos del cluster, 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 cluster 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 cluster

Realice copias de seguridad del cluster 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 cluster

Procedimiento para empezar a administrar el cluster

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

Nota – Los comandos de Oracle Solaris Cluster que sólo ejecuta desde el nodo con voto del cluster global no son válidos para usarlos con los clusters 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 1-2 Herramientas de administración de Oracle Solaris Cluster

Tarea	Herramienta	Instrucciones
Iniciar sesión en el cluster de forma remota	Utilice el comando <code>ccp</code> para iniciar Cluster Control Panel (CCP). A continuación, seleccione uno de los siguientes íconos: <code>cconsole</code> , <code>crlogin</code> , <code>cssh</code> o <code>ctelnet</code> .	“Inicio de sesión remoto en el cluster” en la página 23 “Conexión segura a las consolas del cluster” en la página 24
Configurar el cluster de forma interactiva	Inicie la utilidad <code>clzonecluster(1CL)</code> o la utilidad <code>clsetup(1CL)</code> .	“Obtención de acceso a las utilidades de configuración del cluster” en la página 25
Mostrar información sobre el número y la versión de Oracle Solaris Cluster	Utilice el comando <code>clnode(1CL)</code> con el subcomando y la opción <code>show -rev -v -node</code> .	“Visualización de la información de versión de Oracle Solaris Cluster” en la página 27
Mostrar los recursos, grupos de recursos y tipos de recursos instalados	Utilice los comandos siguientes para mostrar la información sobre recursos: ■ <code>clresource(1CL)</code> ■ <code>clresourcegroup(1CL)</code> ■ <code>clresourcetype(1CL)</code>	“Visualización de tipos de recursos configurados, grupos de recursos y recursos” en la página 30
Supervisar gráficamente los componentes del cluster	Utilice Oracle Solaris Cluster Manager.	Consulte la ayuda en pantalla.
Administrar gráficamente algunos componentes del cluster	Utilice Oracle Solaris Cluster Manager, que sólo está disponible con Oracle Solaris Cluster en los sistemas basados en SPARC.	Para Oracle Solaris Cluster Manager, consulte la ayuda en pantalla.
Comprobar el estado de los componentes del cluster	Utilice el comando <code>cluster(1CL)</code> con el subcomando <code>status</code> .	“Comprobación del estado de los componentes del cluster” en la página 31
Comprobar el estado de los grupos de varias rutas IP en la red pública	En un cluster global, utilice el comando <code>clnode(1CL) status</code> con la opción <code>-m</code> . En un cluster de zona, utilice el comando <code>clzonecluster(1CL) show</code> .	“Comprobación del estado de la red pública” en la página 34
Ver la configuración del cluster	En un cluster global, utilice el comando <code>cluster(1CL)</code> con el subcomando <code>show</code> . En un cluster de zona, utilice el comando <code>clzonecluster(1CL)</code> con el subcomando <code>show</code> .	“Visualización de la configuración del cluster” en la página 35

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

Tarea	Herramienta	Instrucciones
Ver y mostrar los dispositivos NAS configurados	En un cluster de zona o en un cluster global, utilice el comando <code>clzonecluster(1CL)</code> con el subcomando <code>show</code> .	<code>clnasdevice(1CL)</code>
Comprobar los puntos de montaje globales o verificar la configuración del cluster	En un cluster global, utilice el comando <code>cluster(1CL)</code> con el subcomando <code>check</code> . En un cluster de zona, utilice el comando <code>clzonecluster(1CL)</code> con el subcomando <code>verify</code> .	“Validación de una configuración básica de cluster” en la página 43
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 50
Examinar los mensajes del sistema del Oracle Solaris Cluster	Examine el archivo <code>/var/adm/messages</code> .	“Visualización de los mensajes del sistema” de <i>Guía de administración del sistema: Administración avanzada</i>
Supervisar el estado de Solaris Volume Manager	Utilice el comando <code>metastat</code> .	<i>Solaris Volume Manager Administration Guide</i>

▼ Inicio de sesión remoto en el cluster

Cluster Control Panel (CCP) proporciona una pantalla de inicio para las herramientas `cconsole`, `crlogin`, `cssh` y `ctelnet`. Todas las herramientas inician una conexión de ventanas múltiples para un conjunto de nodos especificados. La conexión de ventanas múltiples consta de una ventana de host para cada uno de los nodos especificados y una ventana común. Los datos introducidos en la ventana común se envían a cada una de las ventanas de host, lo que permite ejecutar comandos simultáneamente en todos los nodos del cluster.

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

De forma predeterminada, la utilidad `cconsole` usa una conexión `telnet` con las consolas de los nodos. Para establecer en su lugar conexiones de shell seguro con las consolas, active la casilla de verificación Use SSH (Usar SSH) en el menú Options (Opciones) de la ventana `cconsole`. O bien, especifique la opción `-s` al ejecutar el comando `ccp` o `cconsole`.

Consulte las páginas del comando `man ccp(1M)` and `cconsole(1M)` para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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

Verifique que se cumplen los siguientes requisitos previos antes de iniciar CCP:

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

1 En la consola de administración, abra la pantalla de inicio de CCP.

```
phys-schost# ccp clustername
```

Se muestra la pantalla de inicio de CCP.

2 Para iniciar una sesión remota con el cluster, haga clic en el ícono `cconsole`, `crlogin`, `cssh` o `ctelnet` en la pantalla de inicio de CCP.

▼ Conexión segura a las consolas del cluster

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

Antes de empezar

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

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

Consulte las páginas del comando `man clusters(4)` y `serialports(4)` para obtener más información.

1 Conviértase en superusuario en la consola de administración.

2 Inicie la utilidad `cconsole` en modo seguro.

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

- s Permite activar la conexión de shell seguro.
- l *username* Permite especificar el nombre de usuario para las conexiones remotas. Si no se especifica la opción -l, se utiliza el nombre de usuario que inició la utilidad `cconsole`.
- p *ssh-port* Permite especificar el número de puerto de shell seguro que desea utilizar. Si no se especifica la opción -p, se utiliza el número de puerto predeterminado 22 para las conexiones seguras.

▼ Obtención de acceso a las utilidades de configuración del cluster

La utilidad `clsetup` permite crear un cluster de zona de manera interactiva y configurar opciones de quórum, grupo de recursos, transporte de cluster, nombre de host privado, grupo de dispositivos y nuevo nodo para el cluster global. La utilidad `clzonecluster` realiza tareas de configuración similares para los clusters de zona. Para obtener más información, consulte las páginas del comando `man clsetup(1CL)` y `clzonecluster(1CL)`.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global. Siga todos los pasos de este procedimiento desde un nodo del cluster global.

2 Inicie la utilidad de configuración.

```
phys-schost# clsetup
```

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

```
phys-schost# clsetup
```

Aparece el menú principal.

- En el caso de un cluster de zona, inicie la utilidad con el comando `clzonecluster`. El cluster 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> ?
```

También puede utilizar la utilidad `clsetup` para crear un cluster de zona o agregar un sistema de archivos o un dispositivo de almacenamiento en el ámbito del cluster. Todas las demás tareas de configuración del cluster de zona se realizan con el comando **clzonecluster configure**. Consulte la [Guía de instalación del software de Oracle Solaris Cluster](#) para obtener instrucciones sobre cómo utilizar la utilidad `clsetup`.

- 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 cluster 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 parches de Oracle Solaris Cluster

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

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

Este procedimiento proporciona las formas largas de los comandos de 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 parches de Oracle Solaris Cluster:**

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

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

Ejemplo 1–1 Visualización de información de parches de Oracle Solaris Cluster

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

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

▼ Visualización de la información de versión de Oracle Solaris Cluster

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

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

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

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

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

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

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

En el siguiente ejemplo, se muestra información sobre la versión y el paquete del cluster.

```
phys-schost# clnode show-rev
3.3
```

```
phys-schost# clnode show-rev -v
Oracle Solaris Cluster 3.3 for Solaris 10 sparc
SUNWcccon:      3.3.0,REV=2010.06.14.03.44
SUNWcccon:      3.3.0,REV=2010.06.14.03.44
SUNWcsc:        3.3.0,REV=2010.06.14.03.44
SUNWcscspmu:    3.3.0,REV=2010.06.14.03.44
SUNWcscssv:     3.3.0,REV=2010.06.14.03.44
SUNWecon:       3.3.0,REV=2010.06.14.03.44
SUNWesc:        3.3.0,REV=2010.06.14.03.44
SUNWescspmu:    3.3.0,REV=2010.06.14.03.44
SUNWescssv:     3.3.0,REV=2010.06.14.03.44
SUNWfccon:      3.3.0,REV=2010.06.14.03.44
SUNWfsc:        3.3.0,REV=2010.06.14.03.44
SUNWfscspmu:    3.3.0,REV=2010.06.14.03.44
SUNWfscssv:     3.3.0,REV=2010.06.14.03.44
SUNWjccon:      3.3.0,REV=2010.06.14.03.44
SUNWjcommonS:   3.3.0,REV=2010.06.14.03.44
SUNWjsc:        3.3.0,REV=2010.06.14.03.44
SUNWjscman:     3.3.0,REV=2010.06.14.03.44
SUNWjscspmu:    3.3.0,REV=2010.06.14.03.44
SUNWjscssv:     3.3.0,REV=2010.06.14.03.44
SUNWkcon:       3.3.0,REV=2010.06.14.03.44
SUNWksc:        3.3.0,REV=2010.06.14.03.44
SUNWkscspmu:    3.3.0,REV=2010.06.14.03.44
```

SUNWkscssv: 3.3.0,REV=2010.06.14.03.44
SUNWscu: 3.3.0,REV=2010.06.14.03.44
SUNWsccomu: 3.3.0,REV=2010.06.14.03.44
SUNWsczr: 3.3.0,REV=2010.06.14.03.44
SUNWsccomzu: 3.3.0,REV=2010.06.14.03.44
SUNWsczu: 3.3.0,REV=2010.06.14.03.44
SUNWscsckr: 3.3.0,REV=2010.06.14.03.44
SUNWscscku: 3.3.0,REV=2010.06.14.03.44
SUNWscr: 3.3.0,REV=2010.06.14.03.44
SUNWscrdt: 3.3.0,REV=2010.06.14.03.44
SUNWscrif: 3.3.0,REV=2010.06.14.03.44
SUNWscrtlh: 3.3.0,REV=2010.06.14.03.44
SUNWscnmr: 3.3.0,REV=2010.06.14.03.44
SUNWscnmu: 3.3.0,REV=2010.06.14.03.44
SUNWscdev: 3.3.0,REV=2010.06.14.03.44
SUNWscgds: 3.3.0,REV=2010.06.14.03.44
SUNWscsmf: 3.3.0,REV=2010.06.14.03.44
SUNWscman: 3.3.0,REV=2010.05.21.18.40
SUNWscsal: 3.3.0,REV=2010.06.14.03.44
SUNWscsam: 3.3.0,REV=2010.06.14.03.44
SUNWscvm: 3.3.0,REV=2010.06.14.03.44
SUNWmdmr: 3.3.0,REV=2010.06.14.03.44
SUNWmdmu: 3.3.0,REV=2010.06.14.03.44
SUNWscmasa: 3.3.0,REV=2010.06.14.03.44
SUNWscmasar: 3.3.0,REV=2010.06.14.03.44
SUNWscmasasen: 3.3.0,REV=2010.06.14.03.44
SUNWscmasazu: 3.3.0,REV=2010.06.14.03.44
SUNWscmasau: 3.3.0,REV=2010.06.14.03.44
SUNWscmautil: 3.3.0,REV=2010.06.14.03.44
SUNWscmautilr: 3.3.0,REV=2010.06.14.03.44
SUNWjfreechart: 3.3.0,REV=2010.06.14.03.44
SUNWjfreechartS: 3.3.0,REV=2010.06.14.03.44
ORCLscPeopleSoft: 3.3.0,REV=2010.06.14.03.44
ORCLscobiee: 3.3.0,REV=2010.06.14.03.44
ORCLscoep: 3.3.0,REV=2010.06.14.03.44
ORCLscohs: 3.3.0,REV=2010.06.14.03.44
ORCLscopmn: 3.3.0,REV=2010.06.14.03.44
ORCLscsapnetw: 3.3.0,REV=2010.06.14.03.44
SUNWcvm: 3.3.0,REV=2010.06.14.03.44
SUNWcvmr: 3.3.0,REV=2010.06.14.03.44
SUNWiimsc: 3.3.0,REV=2010.06.14.03.44
SUNWscspmr: 3.3.0,REV=2010.06.14.03.44
SUNWscspmu: 3.3.0,REV=2010.06.14.03.44
SUNWscderby: 3.3.0,REV=2010.06.14.03.44
SUNWscstelemetry: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepavs: 3.2.3,REV=2009.10.23.12.12
SUNWscgrepavsu: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepodg: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepodgu: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepsbpu: 3.3.0,REV=2010.06.14.03.44
SUNWscgrepsrdf: 3.2.3,REV=2009.10.23.12.12
SUNWscgrepsrdfu: 3.3.0,REV=2010.06.14.03.44
SUNWscgreptc: 3.2.3,REV=2009.10.23.12.12
SUNWscgreptcu: 3.3.0,REV=2010.06.14.03.44
SUNWscgspm: 3.3.0,REV=2010.06.14.03.44
SUNWscghb: 3.2.3,REV=2009.10.23.12.12
SUNWscghbr: 3.3.0,REV=2010.06.14.03.44
SUNWscgman: 3.3.0,REV=2010.06.14.03.44

```

ORCLscgprepzfssa: 3.3.0,REV=2010.06.14.03.44
SUNWscgctl: 3.2.3,REV=2009.10.23.12.12
SUNWscgctlr: 3.3.0,REV=2010.06.14.03.44
SUNWscims: 6.0,REV=2003.10.29
SUNWscics: 6.0,REV=2003.11.14
SUNWscids: 3.3.0,REV=2010.06.14.03.44
SUNWscapc: 3.2.0,REV=2006.12.06.18.32
SUNWscdns: 3.2.0,REV=2006.12.06.18.32
SUNWschadb: 3.2.0,REV=2006.12.06.18.32
SUNWschtt: 3.2.0,REV=2006.12.06.18.32
SUNWscslas: 3.2.0,REV=2006.12.06.18.32
SUNWscsrb5: 3.2.0,REV=2006.12.06.18.32
SUNWscnfs: 3.2.0,REV=2006.12.06.18.32
SUNWscor: 3.2.0,REV=2006.12.06.18.32
SUNWscpax: 3.3.0,REV=2010.06.14.03.44
SUNWscslmq: 3.2.0,REV=2006.12.06.18.32
SUNWscsap: 3.2.0,REV=2006.12.06.18.32
SUNWscslc: 3.2.0,REV=2006.12.06.18.32
SUNWscmd: 3.3.0,REV=2010.06.14.03.44
SUNWscsapdb: 3.2.0,REV=2006.12.06.18.32
SUNWscsapenq: 3.2.0,REV=2006.12.06.18.32
SUNWscsaprepl: 3.2.0,REV=2006.12.06.18.32
SUNWscsapscs: 3.2.0,REV=2006.12.06.18.32
SUNWscsapwebas: 3.2.0,REV=2006.12.06.18.32
SUNWscsbl: 3.2.0,REV=2006.12.06.18.32
SUNWscsyb: 3.2.0,REV=2006.12.06.18.32
SUNWscucm: 3.3.0,REV=2010.06.14.03.44
SUNWscwls: 3.3.0,REV=2010.06.14.03.44
SUNWudlm: 3.3.0,REV=2010.06.14.03.44
SUNWudlmr: 3.3.0,REV=2010.06.14.03.44
SUNWscwls: 3.2.0,REV=2006.12.06.18.32
SUNWsc9ias: 3.2.0,REV=2006.12.06.18.32
SUNWscPostgreSQL: 3.2.0,REV=2006.12.06.18.32
SUNWscTimesTen: 3.3.0,REV=2010.06.14.03.44
SUNWsczone: 3.2.0,REV=2006.12.06.18.32
SUNWscdhc: 3.2.0,REV=2006.12.06.18.32
SUNWscsbs: 3.2.0,REV=2006.12.06.18.32
SUNWscmqi: 3.2.0,REV=2006.12.06.18.32
SUNWscmq: 3.2.0,REV=2006.12.06.18.32
SUNWscmys: 3.2.0,REV=2006.12.06.18.32
SUNWscsge: 3.2.0,REV=2006.12.06.18.32
SUNWscsaa: 3.2.0,REV=2006.12.06.18.32
SUNWscsag: 3.2.0,REV=2006.12.06.18.32
SUNWscsmb: 3.2.0,REV=2006.12.06.18.32
SUNWscsps: 3.2.0,REV=2006.12.06.18.32
SUNWscTomcat: 3.2.0,REV=2006.12.06.18.32

```

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

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

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

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

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

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

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

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

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

=== Registered Resource Types ===

Resource Type:	SUNW.qfs
RT_description:	SAM-QFS Agent on Oracle Solaris Cluster
RT_version:	3.1
API_version:	3
RT_basedir:	/opt/SUNWsamfs/sc/bin
Single_instance:	False
Proxy:	False
Init_nodes:	All potential masters

```
Installed_nodes:                <All>
Failover:                       True
Pkglist:                        <NULL>
RT_system:                      False
Global_zone:                    True
=== Resource Groups and Resources ===

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

--- Resources for Group qfs-rg ---

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

▼ Comprobación del estado de los componentes del cluster

Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

phys-schost# **cluster status**

Ejemplo 1-4 Comprobación del estado de los componentes del cluster

En el siguiente ejemplo, se muestra la información sobre el estado de los componentes del cluster que devuelve `cluster(1CL) status`.

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

--- Node Status ---

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

=== Cluster Transport Paths ===

Endpoint1                               Endpoint2                               Status
-----
phys-schost-1:qfe1                      phys-schost-4:qfe1                      Path online
phys-schost-1:hme1                      phys-schost-4:hme1                      Path online

=== Cluster Quorum ===

--- Quorum Votes Summary ---

      Needed   Present   Possible
      -----
      3         3         4

--- Quorum Votes by Node ---

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

--- Quorum Votes by Device ---

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

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

--- Device Group Status ---

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

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

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

=== Cluster Resource Groups ===

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

=== Cluster Resources ===

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

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

=== Zone Clusters ===

--- Zone Cluster Status ---

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

▼ Comprobación del estado de la red pública

Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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 rutas múltiples de red IP, utilice el comando `clnode(1CL)` con el subcomando `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 cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

phys-schost# `clnode status -m`

Ejemplo 1–5 Comprobación del estado de la red pública

En el siguiente ejemplo, se muestra la información sobre el estado de los componentes del cluster que devuelve el comando `clnode status`.

% `clnode status -m`

--- Node IPMP Group Status ---

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

▼ Visualización de la configuración del cluster

Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global o un cluster de zona. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

```
% cluster show
```

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

También puede usar el comando `clzonecluster show` para visualizar la información de configuración sólo de los clusters de zona. Entre las propiedades de un cluster 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 cluster de zona y se aplica sólo a ese cluster de zona. Al ejecutar el comando `clzonecluster show` desde un nodo del cluster de zona, sólo se recupera el estado de los objetos visibles en ese cluster 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 1-6 Visualización de la configuración del cluster global

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

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

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

```
Cluster Name: cluster-1
```

```

clusterid:                0x50C000C4
installmode:              disabled
heartbeat_timeout:        10000
heartbeat_quantum:        1000
private_netaddr:          172.16.0.0
private_netmask:          255.255.248.0
max_nodes:                62
num_zoneclusters:         1
max_privatenets:          10
global_fencing:           pathcount
Node List:                phys-schost-1
Node Zones:               phys_schost-2:za

```

=== Host Access Control ===

```

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

```

=== Cluster Nodes ===

```

Node Name:                 phys-schost-1
Node ID:                   1
Type:                      cluster
Enabled:                   yes
privatehostname:           clusternode1-priv
reboot_on_path_failure:    disabled
globalzoneshares:         3
defaulttpsetmin:           1
quorum_vote:               1
quorum_defaultvote:        1
quorum_resv_key:           0x43CB1E1800000001
Transport Adapter List:    qfe3, hme0

```

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

```

Transport Adapter:         qfe3
  Adapter State:           Enabled
  Adapter Transport Type:  dlpi
  Adapter Property(device_name): qfe
  Adapter Property(device_instance): 3
  Adapter Property(lazy_free): 1
  Adapter Property(dlpi_heartbeat_timeout): 10000
  Adapter Property(dlpi_heartbeat_quantum): 1000
  Adapter Property(nw_bandwidth): 80
  Adapter Property(bandwidth): 10
  Adapter Property(ip_address): 172.16.1.1
  Adapter Property(netmask): 255.255.255.128
  Adapter Port Names:      0
  Adapter Port State(0):   Enabled

```

```

Transport Adapter:         hme0
  Adapter State:           Enabled
  Adapter Transport Type:  dlpi
  Adapter Property(device_name): hme
  Adapter Property(device_instance): 0
  Adapter Property(lazy_free): 0
  Adapter Property(dlpi_heartbeat_timeout): 10000

```

```

Adapter Property(dlpi_heartbeat_quantum):    1000
Adapter Property(nw_bandwidth):              80
Adapter Property(bandwidth):                 10
Adapter Property(ip_address):                 172.16.0.129
Adapter Property(netmask):                    255.255.255.128
Adapter Port Names:                          0
Adapter Port State(0):                       Enabled

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

SNMP MIB Name:                               Event
State:                                       Disabled
Protocol:                                   SNMPv2

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

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

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

Node Name:                                  phys-schost-2:za
Node ID:                                    2
Type:                                        cluster
Enabled:                                    yes
privatehostname:                            clusternode2-priv
reboot_on_path_failure:                     disabled
globalzoneshares:                           1
defaultpsetmin:                              2
quorum_vote:                                1
quorum_defaultvote:                          1
quorum_resv_key:                             0x43CB1E1800000002
Transport Adapter List:                       hme0, qfe3

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

Transport Adapter:                           hme0
Adapter State:                               Enabled
Adapter Transport Type:                       dlpi
Adapter Property(device_name):                 hme
Adapter Property(device_instance):              0
Adapter Property(lazy_free):                   0
Adapter Property(dlpi_heartbeat_timeout):       10000
Adapter Property(dlpi_heartbeat_quantum):       1000
Adapter Property(nw_bandwidth):                 80
Adapter Property(bandwidth):                   10
Adapter Property(ip_address):                   172.16.0.130
Adapter Property(netmask):                     255.255.255.128
Adapter Port Names:                             0
Adapter Port State(0):                         Enabled

Transport Adapter:                           qfe3
Adapter State:                               Enabled
Adapter Transport Type:                       dlpi
Adapter Property(device_name):                 qfe
Adapter Property(device_instance):              3
Adapter Property(lazy_free):                   1

```

```

Adapter Property(dlpi_heartbeat_timeout):    10000
Adapter Property(dlpi_heartbeat_quantum):    1000
Adapter Property(nw_bandwidth):              80
Adapter Property(bandwidth):                 10
Adapter Property(ip_address):                172.16.1.2
Adapter Property(netmask):                   255.255.255.128
Adapter Port Names:                          0
Adapter Port State(0):                      Enabled

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

SNMP MIB Name:                               Event
State:                                       Disabled
Protocol:                                  SNMPv2

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

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

=== Transport Cables ===

Transport Cable:                             phys-schost-1:qfe3,switch2@1
Cable Endpoint1:                           phys-schost-1:qfe3
Cable Endpoint2:                           switch2@1
Cable State:                               Enabled

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

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

Transport Cable:                             phys-schost-2:qfe3,switch2@2
Cable Endpoint1:                           phys-schost-2:qfe3
Cable Endpoint2:                           switch2@2
Cable State:                               Enabled

=== Transport Switches ===

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

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

```

=== Quorum Devices ===

```

Quorum Device Name:          d3
  Enabled:                   yes
  Votes:                     1
  Global Name:               /dev/did/rdisk/d3s2
  Type:                      scsi
  Access Mode:               scsi2
  Hosts (enabled):           phys-schost-1, phys-schost-2

```

```

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

```

=== Device Groups ===

```

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

```

=== Registered Resource Types ===

```

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

```

```

Resource Type:               SUNW.SharedAddress:2
  RT_description:             HA Shared Address Resource Type
  RT_version:                 2
  API_version:                2
  RT_basedir:                 /usr/cluster/lib/rgm/rt/hascip
  Single_instance:            False
  Proxy:                      False
  Init_nodes:                 <Unknown>
  Installed_nodes:            <All>
  Failover:                   True
  Pkglist:                    SUNWscu
  RT_system:                  True

```

```
Resource Type:          SUNW.HAStoragePlus:4
  RT_description:        HA Storage Plus
  RT_version:            4
  API_version:           2
  RT_basedir:            /usr/cluster/lib/rgm/rt/hastorageplus
  Single_instance:       False
  Proxy:                 False
  Init_nodes:            All potential masters
  Installed_nodes:       <All>
  Failover:              False
  Pkglist:               SUNWscu
  RT_system:             False
```

```
Resource Type:          SUNW.haderby
  RT_description:        haderby server for Oracle Solaris Cluster
  RT_version:            1
  API_version:           7
  RT_basedir:            /usr/cluster/lib/rgm/rt/haderby
  Single_instance:       False
  Proxy:                 False
  Init_nodes:            All potential masters
  Installed_nodes:       <All>
  Failover:              False
  Pkglist:               SUNWscderby
  RT_system:             False
```

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

=== Resource Groups and Resources ===

```
Resource Group:          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:          SCSLM_HA_RG
  Resource_project_name:   True
  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:                    default
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:                    default
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/rdsk/c2t2d0
Full Device Path:           phys-schost-1:/dev/rdsk/c2t2d0
Replication:                none
default_fencing:            global

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

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

=== NAS Devices ===

Nas Device:                  nas_filer1
Type:                        sun
User ID:                     root

Nas Device:                  nas2
Type:                        sun
User ID:                     llai

```

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

En el siguiente ejemplo, se enumeran las propiedades de la configuración del cluster de zona.

```

% clzonecluster show
=== Zone Clusters ===

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

--- Solaris Resources for sczone ---

Resource Name:              net
address:                    172.16.0.1
physical:                   auto

Resource Name:              net
address:                    172.16.0.2
physical:                   auto

Resource Name:              fs
dir:                        /gz/db_qfs/CrsHome
special:                    CrsHome

```

```

raw:
type:                                samfs
options:                             []

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

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

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

--- Zone Cluster Nodes for sczone ---

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

Node Name:                            sczone-2
physical-host:                        sczone-2
hostname:                             lzzone-2

```

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

▼ Validación de una configuración básica de cluster

El comando `cluster(1CL)` utiliza el subcomando `check` para validar la configuración básica que necesita un cluster global para funcionar correctamente. Si ninguna comprobación arroja un resultado incorrecto, `cluster check` regresa al indicador de 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` genera 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 cluster.

A partir de Oracle Solaris Cluster 3.3 5/11, el comando `cluster check` se ha mejorado con nuevos tipos de comprobaciones. 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 no se especifica la opción `-k keyword`.

- 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 cluster. 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 functional check-id` para especificar una comprobación funcional. Realice sólo una comprobación funcional cada vez.

Nota – Dado que algunas comprobaciones funcionales implican la interrupción del servicio del cluster, no inicie ninguna comprobación funcional hasta que haya leído la descripción detallada de la comprobación y haya determinado si es necesario retirar antes el cluster 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 con voto del cluster global, se llevan a cabo un conjunto de comprobaciones con el fin de validar la configuración necesaria para que un cluster de zona funcione correctamente. Si todas las comprobaciones son correctas, `clzonecluster verify` devuelve al indicador de shell y el cluster de zona se puede instalar con seguridad. Si falla alguna de las comprobaciones, `clzonecluster verify` informa sobre los nodos del cluster 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 clusters de zona.

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

```
phys-schost# su
```

- 2 Asegúrese de que dispone de las comprobaciones más actuales.**

Vaya a la ficha Patches & Updates (Parches y actualizaciones) de [My Oracle Support](#). Mediante la búsqueda avanzada, seleccione "Solaris Cluster" como producto y especifique "comprobar" en el campo de descripción para localizar los parches de Oracle Solaris Cluster que contengan comprobaciones. Aplique todos los parches que no se encuentren instalados en el cluster.

- 3 Ejecute las comprobaciones de validación básicas.**

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

-v Modo detallado

-o *outputdir* Redirige la salida al subdirectorio *outputdir*.

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

- 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 cluster. No se ve afectada ninguna función del cluster.

- 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 cluster 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 checkID
```

-C *checkID* Especifica una comprobación específica.

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

d. Inicie la comprobación funcional.

```
# cluster check -v -k functional -C checkid -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 cluster de zona para ver si es posible instalar un cluster de zona.

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

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

Consulte “[Cómo registrar los datos de diagnóstico de la configuración del cluster](#)” de *Guía de instalación del software de Oracle Solaris Cluster*.

Ejemplo 1–8 Comprobación de la configuración del cluster 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 1–9 Listado de comprobaciones de validación interactivas

En el siguiente ejemplo se enumeran todas las comprobaciones interactivas que están disponibles para ejecutarse en el cluster. 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 1–10 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 cluster. El cluster 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 1-11 Comprobación de la configuración del cluster global con una comprobación con resultado no satisfactorio

El ejemplo siguiente muestra el nodo `phys-schost-2`, perteneciente al cluster 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.sunccluster.txt
...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE : Global filesystem /etc/vfstab entries are not consistent across
all Oracle Solaris Cluster 3.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
```

▼ Comprobación de los puntos de montaje globales

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

Nota – Ejecute `cluster check` después de efectuar cambios en la configuración del cluster 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 cluster global.

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

```
% su
```

2 Compruebe la configuración del cluster global.

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

Ejemplo 1–12 Comprobación de puntos de montaje globales

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

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

```
cluster check: Requesting explorer data and node report from phys-schost-1.
cluster check: Requesting explorer data and node report from phys-schost-2.
cluster check: phys-schost-1: Explorer finished.
```

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

...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE  : Global filesystem /etc/vfstab entries are not consistent across
all Oracle Solaris Cluster 3.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.phys-schost-1.txt

...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 1398
SEVERITY : HIGH
FAILURE  : An unsupported server is being used as an Oracle Solaris Cluster 3.x node.
ANALYSIS : This server may not been qualified to be used as an Oracle Solaris Cluster 3.x node.
Only servers that have been qualified with Oracle Solaris Cluster 3.x are supported as
Oracle Solaris Cluster 3.x nodes.
RECOMMEND: Because the list of supported servers is always being updated, check with
your Oracle representative to get the latest information on what servers
are currently supported and only use a server that is supported with Oracle Solaris Cluster 3.x.
...
#
```

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

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

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

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

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

- 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 cluster. 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 cluster. 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 1–13 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 53
- “Perfiles de derechos de RBAC en Oracle Solaris Cluster” en la página 54
- “Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster” en la página 55
- “Modificación de las propiedades de RBAC de un usuario” en la página 59

Instalación y uso de RBAC con Oracle Solaris Cluster

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

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

Perfiles de derechos de RBAC en Oracle Solaris Cluster

Oracle Solaris Cluster Manager y 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 con voto y sin voto en un cluster 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 cluster.
- solaris.cluster.modify

Autorización para cambiar las propiedades de un objeto de cluster.

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

Los perfiles de derechos RBAC tienen, al menos, una autorización 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>euid=0</code> .	Ejecutar los comandos seleccionados de Oracle Solaris Cluster que se usen para configurar y administrar un cluster, incluidos los subcomandos siguientes para todos los comandos de Oracle Solaris Cluster: ■ list ■ show ■ status scha_control(1HA) scha_resource_get(1HA) scha_resource_setstatus(1HA) scha_resourcegroup_get(1HA) scha_resourcetype_get(1HA)

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

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 administración de Oracle Solaris Cluster y para asignar usuarios a este rol nuevo.

▼ Creación de un rol con la herramienta de roles administrativos

Antes de empezar

Para crear un rol, debe asumir un rol que tiene el perfil de derechos de administrador principal asignado o ejecutarse como usuario root.

1 Inicie la herramienta Administrative Roles (Roles administrativos).

Para ejecutar la herramienta de roles administrativos, inicie Solaris Management Console, como se describe en [“How to Assume a Role in the Solaris Management Console” de System](#)

Administration Guide: Security Services. Abra User Tool Collection (Recopilación de herramientas de usuario) y haga clic en el ícono Administrative Roles (Roles administrativos).

2 Inicie el Asistente para agregar roles administrativos.

Seleccione Add Administrative Role (Agregar rol administrativo) en el menú Action (Acción) para iniciar el Asistente para agregar roles administrativos con el fin de configurar roles.

3 Configure un rol para asignar el perfil de derechos de gestión de clusters.

Utilice los botones Next (Siguiente) y Back (Atrás) para navegar entre cuadros de diálogo. Tenga en cuenta que el botón Next (Siguiente) no se activa hasta que completa todos los campos necesarios. El último cuadro de diálogo permite revisar los datos introducidos, momento en el cual puede usar el botón Back (Atrás) para cambiar las entradas o hacer clic en Finish (Finalizar) para guardar el rol nuevo. En la lista siguiente, se resumen los botones y los campos del cuadro de diálogo.

Role Name (Nombre de rol)

Nombre abreviado del rol.

Full Name (Nombre completo)

Versión larga del nombre.

Description (Descripción)

Descripción del rol.

Role ID Number (Número de ID de rol)

UID para el rol, incrementado automáticamente.

Role Shell (Shell de rol)

Los shells del perfil que están disponibles para los roles: shell C de administrador, shell Bourne de administrador o shell Korn de administrador.

Create a role mailing list (Crear lista de correo de rol)

Prepara una lista de correo para los usuarios que están asignados a este rol.

Available Rights/Granted Rights (Derechos disponibles/Derechos otorgados)

Asigna o elimina perfiles de derechos de un rol.

Tenga en cuenta que el sistema no impide que se especifiquen varias instancias del mismo comando. Los atributos que se asignan a la primera instancia de un comando en un perfil de derechos tienen prioridad, y las instancias posteriores se ignoran. Utilice las flechas arriba y abajo para cambiar el orden.

Server (Servidor)

Servidor para el directorio principal.

Path (Ruta)

Ruta del directorio principal.

Add (Agregar)

Agrega usuarios que pueden asumir este rol. Deben estar en el mismo ámbito.

Delete (Suprimir)

Suprime usuarios que están asignados a este rol.

Nota – Es necesario colocar este perfil al principio de la lista de perfiles que están asignados al rol.

- 4 Agregue usuarios que necesitan usar las funciones de Oracle Solaris Cluster Manager o los comandos de Oracle Solaris Cluster al rol creado recientemente.**

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

- 5 Haga clic en Finish (Finalizar).**

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

- 7 Inicie y detenga el daemon de antememoria del servicio de nombres.**

El rol nuevo no se activa hasta que se reinicia el daemon de memoria caché del servicio de nombres. Después de convertirse en usuario root, escriba el texto siguiente:

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

▼ Creación de una función desde la línea de comandos

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

- 2 Seleccione un método para crear una función:**

- Para los roles en el ámbito local, use el comando `roleadd(1M)` para especificar un rol local nuevo y sus atributos.
- Asimismo, para los roles en el ámbito local, edite el archivo `user_attr(4)` para agregar un usuario con `type=role`.

Use este método sólo en caso de emergencia.

- Para los roles en el servicio de nombres, use el comando `smrole(1M)` para especificar el rol nuevo y sus atributos.

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

3 Inicie y detenga el daemon de 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 de RBAC de un usuario, elija uno de los procedimientos siguientes.

- “Modificación de las propiedades de RBAC de un usuario con la herramienta de cuentas de usuario” en la página 59
- “Modificación de las propiedades de RBAC de un usuario desde la línea de comandos” en la página 60

▼ Modificación de las propiedades de RBAC de un usuario con la herramienta de cuentas de usuario

Antes de empezar

Para modificar las propiedades de un usuario, debe ejecutar User Tool Collection (Recopilación de herramientas de usuario) como usuario root o asumir un rol que tenga asignado el perfil de derechos de administrador principal.

1 Inicie la herramienta User Accounts (Cuentas de usuario).

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

Una vez que se inicia la herramienta User Accounts (Cuentas de usuario), los íconos de las cuentas de usuario existentes se muestran en el panel de visualización.

2 Haga clic en el ícono User Account (Cuenta de usuario) que se va a modificar y seleccione Properties (Propiedades) en el menú Action (Acción) (o haga doble clic en el ícono de la cuenta de usuario).

3 Haga clic en la ficha adecuada en el cuadro de diálogo de la propiedad que se va a cambiar, como se indica a continuación:

- Para cambiar los roles que están asignados al usuario, haga clic en la ficha Roles y mueva la asignación de rol que se va a cambiar a la columna apropiada: roles disponibles o roles asignados.
- Para cambiar los perfiles de derechos que están asignados al usuario, haga clic en la ficha Rights (Derechos) y mueva la asignación de perfil de derecho que se va a cambiar a la columna adecuada: derechos disponibles o derechos asignados.

Nota – Evite la asignación de perfiles de derechos directamente a los usuarios. El enfoque preferido es requerir que los usuarios asuman roles para ejecutar aplicaciones con privilegios. Esta estrategia desalienta que los usuarios se aprovechen de los privilegios.

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

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

Cierre y arranque de un cluster

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

- “Descripción general sobre el cierre y el arranque de un cluster” en la página 61
- “Cierre y arranque de un solo nodo de un cluster” en la página 71
- “Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad” en la página 85

Para consultar una descripción pormenorizada de los procedimientos tratados en este capítulo, consulte [“Rearranque de un nodo en un modo que no sea de cluster”](#) en la página 82 y la [Tabla 3–2](#).

Descripción general sobre el cierre y el arranque de un cluster

El comando `cluster shutdown` de Oracle Solaris Cluster detiene los servicios del cluster global de manera correcta y cierra sin sobresaltos un cluster global por completo. Puede utilizar el comando `cluster shutdown` al desplazar la ubicación de un cluster global o para cerrar el cluster global si un error de aplicación está dañando los datos. El comando `clzonecluster halt` detiene un cluster de zona que se ejecuta en un determinado nodo o en un cluster de zona completo en todos los nodos configurados. También puede utilizar el comando `cluster shutdown` con los clusters de zona. Para obtener más información, consulte la página del comando [man `cluster\(1CL\)`](#).

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

Nota – Use el comando `cluster shutdown` para asegurarse de que todo el cluster global se cierre correctamente. El comando `shutdown` de Oracle Solaris se utiliza con el comando `clnode(1CL)` evacúe para cerrar nodos individuales. Si desea obtener más información, consulte [“Cierre de un cluster” en la página 63](#) o [“Cierre y arranque de un solo nodo de un cluster” en la página 71](#).

Los comandos `cluster shutdown` y `clzonecluster halt` detienen todos los nodos comprendidos en un cluster global o un cluster 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 cluster correspondientes a un cluster global o uno de zona.
3. El comando `cluster shutdown` cierra los servicios de dispositivos activos en el cluster global o de zona.
4. El comando `cluster shutdown` ejecuta `init 0`; en los sistemas basados en SPARC, lleva a todos los nodos del cluster al indicador de solicitud OpenBoot PROM `ok` o al mensaje `Press any key to continue` (Pulse cualquier tecla para continuar) del menú de GRUB en el caso de los sistemas basados en x86. Los menús de GRUB se describen de forma más detallada en [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica](#). El comando `clzonecluster halt` ejecuta la operación del comando `zoneadm -z zoneclustername halt` para detener (pero no cerrar) las zonas del cluster de zona.

Nota – Si es necesario, tiene la posibilidad de arrancar un nodo en un modo que no sea de cluster para que el nodo no participe como miembro en el cluster. Los modos que no son de cluster son útiles al instalar software de cluster o para efectuar determinados procedimientos administrativos. Si desea obtener más información, consulte [“Rearranque de un nodo en un modo que no sea de cluster” en la página 82](#).

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

Tarea	Instrucciones
Detener el cluster.	“Cierre de un cluster” en la página 63
Iniciar el cluster arrancando todos los nodos. Los nodos deben contar con una conexión operativa con la interconexión de cluster para conseguir convertirse en miembros del cluster.	“Arranque de un cluster” en la página 65
Rearranque el cluster.	“Rearranque de un cluster” en la página 67

▼ Cierre de un cluster

Puede cerrar un cluster global, un cluster de zona o todos los clusters de zona.



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

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global o el cluster 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 cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.
- 3 Cierre el cluster global, el de zona o todos los clusters de zona.
 - Cierre el cluster global. Esta acción cierra también todos los clusters de zona.
`phys-schost# cluster shutdown -g0 -y`
 - Cierre un cluster de zona concreto.
`phys-schost# clzonecluster halt zoneclustername`
 - Cierre todos los clusters de zona.
`phys-schost# clzonecluster halt +`
Para cerrar el cluster de una zona en particular, también puede usar el comando `cluster shutdown` dentro de un cluster de zona.

- 4 **Compruebe que todos los nodos del cluster global o el cluster 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 el estado de uno o más nodos de cluster global de otro nodo de cluster global que aún está activo y en ejecución en el cluster.**

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

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

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

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

Ejemplo 3-1 Cierre de un cluster de zona

En el siguiente ejemplo, se cierra un cluster de zona denominado *sparse-sczone*.

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

Ejemplo 3-2 SPARC: Cierre de un cluster global

En el ejemplo siguiente se muestra la salida de una consola cuando se detiene el funcionamiento normal de un cluster 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 cluster 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 cluster global

En el ejemplo siguiente se muestra la salida de la consola al detenerse el funcionamiento normal del cluster 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 cluster 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 cluster” en la página 65](#) para reiniciar un cluster global o un cluster de zona que se ha cerrado.

▼ Arranque de un cluster

Este procedimiento explica cómo arrancar un cluster global o uno de zona cuyos nodos se han cerrado. En los nodos de un cluster 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster de zona, siga las instrucciones de [“Configuración de un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster](#).

1 Inicie cada nodo en el modo de cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster 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. Aparece el menú de GRUB, que es similar al siguiente:

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

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica](#).

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

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

phys-schost# **clzonecluster boot** zoneclustername

- Si tiene más de un cluster de zona, puede arrancar todos los clusters de zona. Use + en lugar de zoneclustername.

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

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

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

Al ejecutar el comando de estado **clzonecluster(1CL)** desde un nodo del cluster global, este comando informa sobre el estado del nodo del cluster 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 85.](#)

Ejemplo 3–4 SPARC: Arranque de un cluster global

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

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

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

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster ejecuta Oracle RAC, cierre todas las instancias de base de datos del cluster 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 cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

- 3 **Cierre el cluster.**

- **Cierre el cluster global.**

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

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

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

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

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

- 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 cluster global que esté en un sistema basado en SPARC, ejecute el comando siguiente.

```
ok boot
```

- Para un nodo del cluster 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. Aparece el menú de GRUB, que es similar al siguiente:

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

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

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

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

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

A medida que se activan los componentes del cluster, 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 cluster global.

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

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

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

También puede ejecutar el comando `cluster status` en un cluster 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 85.](#)

Ejemplo 3–5 Rearranque de un cluster de zona

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

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

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

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

=== Zone Clusters ===

--- Zone Cluster Status ---

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

phys-schost#

Ejemplo 3-6 SPARC: Rearranque de un cluster global

El ejemplo siguiente muestra la salida de la consola al detenerse el funcionamiento normal del cluster global, todos los nodos se cierran y muestran el indicador ok, y se reinicia el cluster 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 cluster 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 cluster

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

Para cerrar un nodo del cluster 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 cluster global.

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

Para obtener información sobre cómo cerrar e iniciar una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*. Consulte también 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 cluster 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 cluster global, use <code>clnode(1CL)</code> <code>evacuate</code> y <code>shutdown</code> . Para un nodo de un cluster de zona, use <code>clzonecluster(1CL)</code> <code>halt</code> .	“Cierre de un nodo” en la página 72
Iniciar un nodo. El nodo debe disponer de una conexión operativa con la interconexión de cluster para poder convertirse en miembro de este último.	Para un nodo del cluster global, use <code>boot</code> o <code>b</code> . Para un nodo de un cluster de zona, use <code>clzonecluster(1CL)</code> <code>boot</code> .	“Arranque de un nodo” en la página 75
Detener y reiniciar (rearrancar) un nodo de un cluster. El nodo debe disponer de una conexión operativa con la interconexión de cluster para poder convertirse en miembro de este último.	Para un nodo del cluster global, use <code>clnode</code> <code>evacuate</code> y <code>shutdown</code> , seguido de <code>boot</code> o <code>b</code> . Para un nodo de un cluster de zona, use <code>clzonecluster(1CL)</code> <code>reboot</code> .	“Rearranque de un nodo” en la página 78
Arrancar un nodo de manera que no participe como miembro en el cluster.	Para un nodo del cluster global, use los comandos <code>clnode</code> <code>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 cluster global subyacente se arranca en un modo que no sea de cluster, el nodo del cluster de zona pasa automáticamente al modo que no sea de cluster.	“Rearranque de un nodo en un modo que no sea de cluster” en la página 82

▼ **Cierre de un nodo**

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster para cerrar un nodo del cluster global ni de un cluster de zona. El comando no puede utilizarse dentro de un cluster.

- 1 **Si el cluster ejecuta Oracle RAC, cierre todas las instancias de base de datos del cluster 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 cluster que se va a cerrar. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

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

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

Al especificar un determinado nodo de un cluster de zona, sólo se detiene ese nodo. El comando `halt` detiene de forma predeterminada los clusters 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 cluster global.**

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

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

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

- 5 **Cierre el nodo.**

Especifique el nodo del cluster global que quiera cerrar.

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

Compruebe que el nodo del cluster 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 cluster global

En el siguiente ejemplo, se 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 cluster global.

```
phys-schost# clnode evacuate phys-schost-1
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 cluster 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 cluster 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 cluster de zona

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

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

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

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

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

```
phys-schost#
```

```
phys-schost# clzonecluster halt -n schost-4 sparse-sczone
```

```
Waiting for zone halt commands to complete on all the nodes of the zone cluster "sparse-sczone"...
```

```
Sep  5 19:24:00 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone' died.
```

```
phys-host#
```

```
phys-host# clzonecluster status
```

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

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

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

```
phys-schost#
```

Véase también Consulte [“Arranque de un nodo” en la página 75](#) para ver cómo reiniciar un nodo del cluster global que se haya cerrado.

▼ Arranque de un nodo

Si desea cerrar o reiniciar otros nodos activos del cluster global o del cluster 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 cluster que se cierran o rearranquen. Para obtener información sobre cómo iniciar una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

Nota – La configuración del quórum puede afectar el inicio de un nodo. En un cluster de dos nodos, debe tener el dispositivo de quórum configurado de manera que el número total de quórum correspondiente al cluster 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 cluster. Para que el primer nodo retorne al cluster como nodo integrante, el segundo debe estar operativo y en ejecución. Debe estar el correspondiente número de quórum de cluster (dos).

Si ejecuta Oracle Solaris Cluster en un dominio invitado, el reinicio del control o el dominio de E/S puede afectar el dominio invitado en ejecución, incluido el dominio que se va a cerrar. Debe volver a equilibrar la carga de trabajo en otros nodos y detener el dominio invitado que ejecute Oracle Solaris Cluster antes de volver a iniciar el control o el dominio de E/S.

Cuando un control o un dominio de E/S se reinicia, el dominio invitado no envía ni recibe señales. Esto genera una situación de "cerebro dividido" y una reconfiguración del cluster. Como se reinicia el control o el dominio de E/S, el dominio invitado no puede tener acceso a ningún dispositivo compartido. Los otros nodos del cluster ponen una barrera entre el dominio invitado y los dispositivos compartidos. Cuando el control o el dominio de E/S se termina de reiniciar, se reanuda la E/S en el dominio invitado, y cualquier E/S de un almacenamiento compartido hace que el dominio invitado emita un aviso grave a causa de la barrera que le impide acceder a los discos compartidos como parte de la reconfiguración del cluster. Puede mitigar este problema si un invitado utiliza dos dominios de E/S para la redundancia, y usted reinicia los dominios de E/S de uno en uno.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster para poder convertirse en miembros del cluster.

- 1 Para iniciar un nodo del cluster global o un nodo de un cluster de zona que se haya cerrado, arranque el nodo. Siga todos los pasos de este procedimiento desde un nodo del cluster 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. Aparece el menú de GRUB, que es similar al siguiente:

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

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

- Si tiene un cluster 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 cluster global.

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

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

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

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

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 85](#).

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

El ejemplo siguiente muestra la salida de la consola cuando se arranca el nodo `phys-schost-1` en el cluster 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 cluster global o en el cluster 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 cluster que se cierran o rearranquen. Para obtener información sobre cómo reiniciar una zona no global, consulte el [Capítulo 20, “Cómo instalar, iniciar, detener, desinstalar y clonar zonas no globales \(tareas\)”](#) de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global o del cluster 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 cluster global.**
- 3 **Cierre el nodo del cluster global con los comandos `clnode evacuate` y `shutdown`. Cierre el cluster de zona mediante la ejecución del comando `clzonecluster halt` desde un nodo del cluster global. Los comandos `clnode evacuate` y `shutdown` también funcionan en los clusters de zona.**

En el caso de un cluster 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. Este comando también cambia todos los grupos de recursos de las zonas globales o no globales del nodo especificado a las zonas globales o no globales de otros nodos que se sitúen a continuación por 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 cluster.

- 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. Aparece el menú de GRUB, que es similar al siguiente:

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

- Indique el nodo del cluster 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 cluster para poder convertirse en miembros del cluster.

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

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

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

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

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

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

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

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

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

```
The system is ready.
phys-schost-1 console login:
```

Ejemplo 3-12 x86: Rearranque de un nodo del cluster global

En el siguiente ejemplo, se muestra la salida de la consola cuando se reinicia 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 cluster global.

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

GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                     |
|                                                       |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
Hostname: phys-schost-1
Booting as part of a cluster
...
NOTICE: Node phys-schost-1: attempting to join cluster
...
NOTICE: Node phys-schost-1: joined cluster
...
The system is coming up. Please wait.
checking ufs filesystems
...
reservation program successfully exiting
Print services started.
volume management starting.
The system is ready.
phys-schost-1 console login:
```

Ejemplo 3-13 Rearranque de un nodo del cluster global

El ejemplo siguiente muestra cómo rearrancar un nodo de un cluster 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-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#
```

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

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

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster que se va a iniciar en un modo que no es de cluster. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**
- 2 **Cierre el nodo del cluster de zona con el comando `clzonecluster halt` en un nodo del cluster global. Cierre el nodo del cluster global con los comandos `clnode evacuate` y `shutdown`.**

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

- **Cierre un cluster global concreto.**

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

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

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

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

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

- 3 Compruebe que el nodo del cluster 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 cluster global en un modo que no sea de cluster.**

- 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.**

Aparece el menú de GRUB, que es similar al siguiente:

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

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

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

- 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.**

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a)                                         |
| kernel /platform/i86pc/multiboot                     |
| module /platform/i86pc/boot_archive                   |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.

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

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

```
grub edit> kernel /platform/i86pc/multiboot -x
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a)                                     |
| kernel /platform/i86pc/multiboot -x                |
| module /platform/i86pc/boot_archive                |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-

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

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 cluster. Para arrancarlo en el modo que no es de cluster, realice estos pasos de nuevo para agregar la opción -x al comando del parámetro de arranque del núcleo.

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

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 cluster. 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 cluster 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 registran los mensajes de error en el archivo /var/adm/messages, que con el paso del tiempo puede ocupar todo el sistema de archivos /var. Si el sistema de archivos /var de un nodo del cluster alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en ese nodo. Además, es posible que no pueda iniciarse sesión en ese nodo.

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

Si un nodo emite un informe que anuncia que el sistema de archivos /var ha alcanzado su límite de capacidad y continúa ejecutando servicios de Oracle Solaris Cluster, use este procedimiento para borrar el sistema de archivos. Consulte [“Visualización de los mensajes del sistema” de Guía de administración del sistema: Administración avanzada](#) para obtener más información.

- 1 **Conviértase en superusuario en el nodo del cluster 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 cluster.

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

- Entre clusters: use Oracle Solaris Cluster Geographic Edition para recuperar datos en caso de problema grave.
- En un cluster: úselo como sustitución de la duplicación basada en host en un cluster 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 123](#) en el capítulo 5.

Debe comprender la replicación de datos basada en almacenamiento y también la basada en host para poder seleccionar el método de replicación que mejor se adapte a su cluster. Para obtener más información sobre cómo usar Oracle Solaris Cluster Geographic Edition para gestionar la replicación de datos para la recuperación de problemas graves, consulte [Oracle Solaris Cluster Geographic Edition Overview](#).

Este capítulo incluye las secciones siguientes:

- [“Comprensión de la replicación de datos” en la página 88](#)
- [“Uso de la replicación de datos basada en almacenamiento en un cluster” en la página 90](#)

Comprensión de la replicación de datos

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

- *La replicación de datos basada en host* usa el software para replicar en tiempo real volúmenes de discos entre clusters ubicados en puntos geográficos distintos. La replicación por duplicación remota permite replicar los datos del volumen maestro del cluster primario en el volumen maestro del cluster 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 clusters (y entre un cluster y un host que no se encuentra en el cluster) es Sun StorageTek Availability Suite 4.

La replicación de datos basada en host es una solución de replicación más económica, porque usa recursos del host en lugar de matrices de almacenamiento especiales. No se admiten las bases de datos, las aplicaciones ni los sistemas de archivos configurados para permitir que varios hosts ejecuten el sistema operativo Oracle Solaris para escribir datos en un volumen compartido (por ejemplo, Oracle 9iRAC y Oracle Parallel Server). Para obtener más información sobre cómo usar la replicación de datos basada en host entre dos clusters, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Sun StorageTek Availability Suite](#). Para ver un ejemplo de replicación basada en host que no use Oracle Solaris Cluster Geographic Edition, consulte el apéndice A, “Configuración de replicación de datos basada en host con el software de Availability Suite” en la página 337.

- *La replicación de datos basada en almacenamiento* usa el software del controlador de almacenamiento para mover el trabajo de la replicación de datos de los nodos del cluster al dispositivo de almacenamiento. Este software libera algo de la capacidad de procesamiento del nodo para satisfacer las solicitudes del cluster. Hitachi TrueCopy, Hitachi Universal Replicator y EMC SRDF son ejemplos de software basado en almacenamiento que puede replicar datos dentro de un cluster o entre clusters. La replicación de datos basada en almacenamiento puede ser especialmente importante en las configuraciones de cluster de campus y puede simplificar la infraestructura necesaria. Para obtener más información sobre el uso de la replicación de datos basada en almacenamiento en un entorno de cluster de campus, consulte “Uso de la replicación de datos basada en almacenamiento en un cluster” en la página 90.

Para obtener más información sobre cómo usar la replicación basada en almacenamiento entre dos o más clusters y el producto Oracle Solaris Cluster Geographic Edition que automatiza el proceso, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Hitachi TrueCopy and Universal Replicator](#) y [Oracle Solaris Cluster Geographic Edition Data Replication Guide for EMC Symmetrix Remote Data Facility](#). Consulte también el apéndice A, “Configuración de replicación de datos basada en host con el software de Availability Suite” en la página 337 para obtener un ejemplo completo de este tipo de configuración de cluster.

Métodos admitidos de replicación de datos

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

1. Replicación entre clusters: para la recuperación después de un desastre, puede usar una replicación basada en almacenamiento, host o aplicación con el objeto de realizar una replicación de datos entre clusters. Por lo general, se elige un método de replicación en lugar de una combinación de los tres métodos. Puede gestionar todos los tipos de replicación con el software de Oracle Solaris Cluster Geographic Edition.

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

Si desea utilizar la replicación basada en host sin el software de Oracle Solaris Cluster Geographic Edition, consulte las instrucciones en el [Apéndice A, “Ejemplo”, “Configuración de replicación de datos basada en host con el software de Availability Suite” en la página 337](#).

- Replicación basada en almacenamiento
 - Hitachi TrueCopy y Hitachi Universal Replicator, mediante Oracle Solaris Cluster Geographic Edition
 - EMC Symmetrix Remote Data Facility (SRDF), mediante Oracle Solaris Cluster Geographic Edition
 - Sun ZFS Storage Appliance de Oracle. Para obtener más información, consulte [“Data Replication” de Oracle Solaris Cluster Geographic Edition Overview](#).

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

- Replicación basada en aplicación
 - Oracle Data Guard
 - MySQL

Para obtener más información sobre la recuperación después de un desastre, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard](#).

2. Replicación dentro de un cluster: este método se utiliza como sustitución para la duplicación basada en host.
 - Replicación basada en almacenamiento
 - Hitachi TrueCopy y Hitachi Universal Replicator
 - EMC Symmetrix Remote Data Facility (SRDF)

Uso de la replicación de datos basada en almacenamiento en un cluster

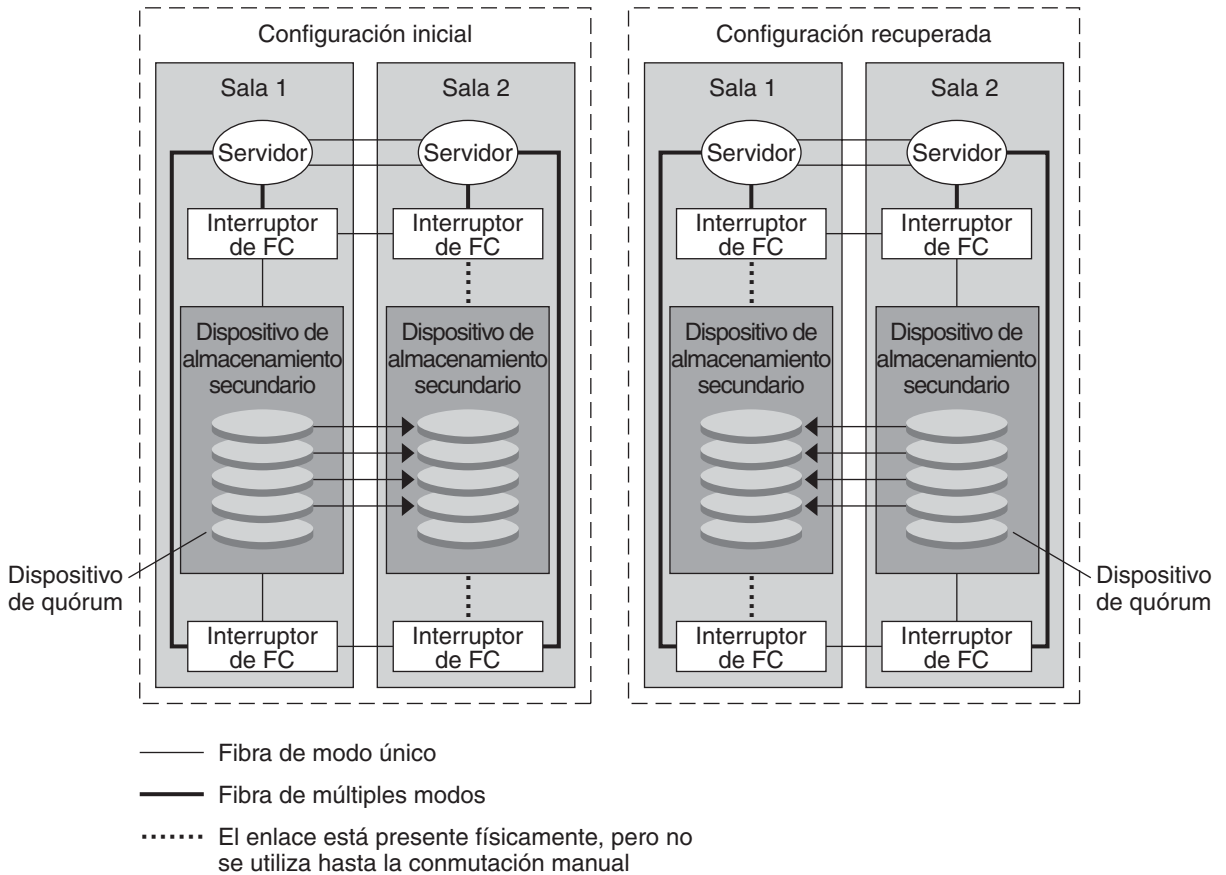
La replicación de datos basada en almacenamiento utiliza software instalado en el dispositivo de almacenamiento para gestionar la replicación dentro de un cluster o un cluster de campus. Dicho software es específico de su dispositivo de almacenamiento particular y si no se utiliza para la recuperación después de un desastre. Consulte la documentación incluida con el dispositivo de almacenamiento cuando vaya a configurar la replicación de datos basada en almacenamiento.

Según el software que utiliza, puede usar la conmutación por error manual o automática con la replicación de datos basada en almacenamiento. Oracle Solaris Cluster admite la conmutación por error manual y automática de los replicadores con el software de EMC SRDF, Hitachi Universal Replicator y Hitachi TrueCopy.

En esta sección, se describe la replicación de datos basada en almacenamiento como se la utiliza en un cluster de campus. En la [Figura 4–1](#), aparece un ejemplo de configuración de dos salas en donde los datos se replican entre dos matrices de almacenamiento. En esta configuración, la matriz de almacenamiento principal se incluye en la sala primaria, donde proporciona datos a los nodos de ambas salas. La matriz de almacenamiento principal también proporciona a la matriz de almacenamiento secundaria datos que se van a replicar.

Nota – En la [Figura 4–1](#), se muestra que el dispositivo del quórum está en un volumen sin replicar. Un volumen replicado no se puede utilizar como dispositivo del quórum.

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



La replicación de datos basada en almacenamiento con Hitachi TrueCopy o Hitachi Universal Replicator se puede realizar de manera síncrona y asíncrona en el entorno de Oracle Solaris Cluster según el tipo de aplicación que se usa. Si desea realizar una conmutación por error automática en un cluster de campus, utilice TrueCopy de forma síncrona. La replicación síncrona basada en almacenamiento con EMC SRDF es compatible con Oracle Solaris Cluster; la replicación asíncrona no es compatible con EMC SRDF.

No use el modo Domino (Dominó) ni el modo Adaptive Copy (Copia adaptable) de EMC SRDF. El modo Domino (Dominó) hace que los volúmenes SRDF locales y de destino no estén disponibles para el host si el destino no está disponible. El modo Adaptive Copy (Copia adaptable) se usa normalmente para migraciones de datos y movimientos del centro de datos, y no se recomienda para la recuperación después de un desastre.

Si se pierde el contacto con el dispositivo de almacenamiento remoto, asegúrese de que una aplicación que se ejecute en el cluster principal no se bloquee mediante la especificación de un

`fence_level` de `never` o `async`. Si especifica un `fence_level` de `data` o `status`, el dispositivo de almacenamiento principal rechaza las actualizaciones si las actualizaciones no se pueden copiar en el dispositivo de almacenamiento remoto.

Requisitos y restricciones del uso de la replicación de datos basada en almacenamiento en un cluster

Para garantizar la integridad de los datos, utilice múltiples rutas y el paquete RAID adecuado. En la lista siguiente, se incluyen consideraciones para la implementación de una configuración de cluster que usa replicación de datos basada en almacenamiento.

- La distancia de nodo a nodo está limitada por la infraestructura de interconexión y el canal de fibra de Oracle Solaris Cluster. Póngase en contacto con el proveedor de servicios de Oracle para obtener más información acerca de las limitaciones actuales y las tecnologías admitidas.
- No configure un volumen replicado como dispositivo del quórum. Localice los dispositivos del quórum en un volumen compartido sin replicar o utilice el servidor del quórum.
- Asegúrese de que sólo la copia principal de los datos permanece visible para los nodos del cluster. De lo contrario, el administrador de volúmenes puede intentar acceder simultáneamente a las copias principales y secundarias de los datos. Consulte la documentación que recibió con la matriz de almacenamiento para obtener más información sobre el control de la visibilidad de las copias de datos.
- EMC SRDF, Hitachi TrueCopy y Hitachi Universal Replicator permiten definir grupos de dispositivos replicados. Cada grupo de dispositivo de replicación requiere grupo de dispositivo de Oracle Solaris Cluster con el mismo nombre.
- Para una configuración de tres sitios o de tres centros de datos que usan EMC SRDF con dispositivos RDF simultáneos o en cascada, debe agregar la siguiente entrada en el archivo de opciones de Solutions Enabler (SYMCLI) en todos los nodos participantes del cluster:

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

Esta entrada permite que el software del cluster automatice el movimiento de la aplicación entre los dos sitios SRDF síncronos. El *rdf-group-number* de la entrada representa el grupo RDF que conecta el symmetrix local del host al symmetrix del segundo sitio.

Para obtener más información sobre las configuraciones de tres centros de datos, consulte [“Three-Data-Center \(3DC\) Topologies” de Oracle Solaris Cluster Geographic Edition Overview](#).

- Es posible que ciertos datos específicos de la aplicación no sean adecuados para la replicación de datos asincrónica. Emplee sus conocimientos acerca del comportamiento de la aplicación para determinar la mejor manera de replicar datos específicos de la aplicación en los dispositivos de almacenamiento.

- Si va a configurar el cluster para la conmutación por error automática, utilice la replicación sincrónica.
Para obtener instrucciones de configuración del cluster para la conmutación por error automática de volúmenes replicados, consulte [“Administración de dispositivos replicados basados en el almacenamiento” en la página 97](#).
- Oracle Real Application Clusters (RAC) no es compatible con SRDF, Hitachi TrueCopy y Hitachi Universal Replicator al replicar dentro de un cluster. Los nodos conectados a las réplicas que no son actualmente la réplica principal no tienen acceso de lectura. Cualquier aplicación escalable que requiere acceso directo de escritura desde todos los nodos del cluster no es compatible con dispositivos replicados.
- No se admite Solaris Volume Manager de múltiples propietarios para el software de Oracle Solaris Cluster.
- No use el modo Domino (Dominó) ni el modo Adaptive Copy (Copia adaptable) de EMC SRDF. Consulte [“Uso de la replicación de datos basada en almacenamiento en un cluster” en la página 90](#) para obtener más información.
- No utilice el modo Data (Datos) o el modo Status (Estado) en Hitachi TrueCopy o Hitachi Universal Replicator. Consulte [“Uso de la replicación de datos basada en almacenamiento en un cluster” en la página 90](#) para obtener más información.

Problemas de recuperación manual en el uso de la replicación de datos basada en almacenamiento en un cluster

Al igual que ocurre con los clusters de campus, los clusters que utilizan la replicación de datos basada en almacenamiento generalmente no necesitan intervención cuando se produce un solo fallo. Sin embargo, si utiliza la conmutación por error manual y se pierde la sala que contiene el dispositivo de almacenamiento principal (como se muestra en la [Figura 4–1](#)), se ocasionan problemas en un cluster de dos nodos. El nodo restante no puede reservar el dispositivo de quórum y no puede iniciar como miembro del cluster. En esta situación, el cluster requiere la siguiente intervención manual:

1. El proveedor de servicios de Oracle debe volver a configurar el nodo restante para iniciar como miembro del cluster.
2. O usted o el proveedor de servicios de Oracle deben configurar un volumen sin replicar del dispositivo de almacenamiento secundario como dispositivo del quórum.
3. O usted o el proveedor de servicios de Oracle deben configurar el nodo restante para utilizar el dispositivo de almacenamiento secundario como almacenamiento principal. Esta reconfiguración puede suponer la reconstrucción de los volúmenes del administrador de volúmenes, la restauración de datos o el cambio de asociaciones de aplicaciones con volúmenes de almacenamiento.

Mejores prácticas para utilizar la replicación de datos basada en almacenamiento

Al configurar grupos de dispositivos que utilizan el software de Hitachi TrueCopy o Hitachi Universal Replicator para la replicación de datos basada en almacenamiento, tenga en cuenta las siguientes prácticas:

- Utilice la replicación síncrona para evitar la posibilidad de perder datos si falla el sitio principal.
- Debe existir una relación de uno a uno entre el grupo de dispositivos global de Oracle Solaris Cluster y el grupo de replicación de TrueCopy definido en el archivo de configuración `horcm`. Esto permite que ambos grupos se muevan de un nodo a otro como una sola unidad.
- Los volúmenes del sistema de archivos global y los volúmenes del sistema de archivos de conmutación por error no pueden mezclarse en el mismo grupo de dispositivos replicado porque se controlan de forma distinta. Los sistemas de archivos globales se controlan mediante el sistema de configuración de dispositivos (DCS), mientras que los volúmenes del sistema de archivos de conmutación por error se controlan mediante HAS+. El nodo principal para cada uno podría ser diferente, lo que provocaría conflictos para determinar qué nodo debería ser el nodo principal de replicación.
- Todas las instancias del gestor de RAID deben estar activas en todo momento.

Al utilizar el software de EMC SRDF para la replicación de datos basada en almacenamiento, use dispositivos dinámicos en lugar de dispositivos estáticos. Los dispositivos estáticos requieren varios minutos para cambiar la replicación principal y pueden afectar el tiempo de conmutación por error.

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

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

- “Información general sobre la administración de dispositivos globales y el espacio de nombre global” en la página 95
- “Administración de dispositivos replicados basados en el almacenamiento” en la página 97
- “Información general sobre la administración de sistemas de archivos de clusters” en la página 122
- “Administración de grupos de dispositivos” en la página 123
- “Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento” en la página 149
- “Administración de sistemas de archivos de cluster” en la página 154
- “Administración de la supervisión de rutas de disco” en la página 160

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

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 cluster, 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 administrador de volúmenes instalado en el cluster. Solaris Volume Manager “reconoce clusters”, de modo puede agregar, registrar y eliminar grupos de dispositivos mediante el comando `metaset(1M)` de Solaris Volume Manager.

Nota – No es posible acceder directamente a los dispositivos globales desde los nodos sin voto del cluster global.

El software Oracle Solaris Cluster crea automáticamente un grupo de dispositivos de discos básicos para cada dispositivo de disco y de cinta del cluster. Sin embargo, los grupos de dispositivos del cluster 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 cluster 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 cluster. Este comando hace que el espacio de nombre global se actualice en todos los demás nodos miembros del cluster y en los nodos que posteriormente tengan la posibilidad de asociarse al cluster.

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 cluster 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 cluster. 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 cluster:

```
# 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 cluster, 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 cluster 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 171](#).



Precaución – Si el nodo primario actual falla durante la operación de reconfiguración dinámica, dicho error repercute en la disponibilidad del cluster. 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 146
2. Efectuar la operación de eliminación de DR en el dispositivo que se va a eliminar.	

Administración de dispositivos replicados basados en el almacenamiento

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

Antes de replicar los datos con el software de Hitachi TrueCopy o EMC Symmetrix Remote Data Facility, debe estar familiarizado con la documentación de replicación basada en almacenamiento y tener instalados en el sistema el producto de replicación basada en

almacenamiento y los parches más recientes. Si desea obtener información sobre cómo instalar el software de replicación basada en almacenamiento, consulte la documentación del producto.

El software de replicación basada en almacenamiento configura un par de dispositivos como réplicas, uno como réplica principal y el otro como réplica secundaria. En un momento dado, el dispositivo conectado a un conjunto de nodos será la réplica principal. El dispositivo conectado al otro conjunto de nodos será la réplica secundaria.

En una configuración de Oracle Solaris Cluster, la réplica principal se mueve automáticamente cada vez que se mueve el grupo de dispositivos de Oracle Solaris Cluster al que pertenece la réplica. Por lo tanto, la réplica principal nunca debe moverse directamente en una configuración de Oracle Solaris Cluster. En su lugar, la recuperación debe realizarse moviendo el grupo de dispositivos de Oracle Solaris Cluster asociado.



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

Esta sección incluye los procedimientos siguientes:

- “Administración de dispositivos replicados de Hitachi TrueCopy” en la página 98
- “Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility” en la página 109

Administración de dispositivos replicados de Hitachi TrueCopy

En la siguiente tabla, se enumeran las tareas que debe realizar para configurar un dispositivo replicado basado en almacenamiento de Hitachi TrueCopy.

TABLA 5-2 Mapa de tareas: administración de un dispositivo replicado basado en almacenamiento de Hitachi TrueCopy

Tarea	Instrucciones
Instalar el software de TrueCopy en el dispositivo de almacenamiento y en los nodos	Consulte la documentación incluida con el dispositivo de almacenamiento de Hitachi.
Configurar el grupo de replicación de Hitachi	“Configuración de un grupo de replicación de Hitachi TrueCopy” en la página 99
Configurar el dispositivo DID	“Configuración de dispositivos DID para la replicación con Hitachi TrueCopy” en la página 101
Registrar el grupo replicado	“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 130

TABLA 5-2 Mapa de tareas: administración de un dispositivo replicado basado en almacenamiento de Hitachi TrueCopy (Continuación)

Tarea	Instrucciones
Comprobar la configuración	“Verificación de la configuración de un grupo de dispositivos globales replicados de Hitachi TrueCopy” en la página 103

▼ Configuración de un grupo de replicación de Hitachi TrueCopy

Antes de empezar

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



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

- 1 **Conviértase en superusuario o asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos conectados a la matriz de almacenamiento.**
- 2 **Agregue la entrada `horcm` al archivo `/etc/services`.**
`horcm 9970/udp`
Especifique un número de puerto y un nombre de protocolo para la nueva entrada.
- 3 **Especifique la información de configuración del grupo de dispositivos de Hitachi TrueCopy en el archivo `/etc/horcm.conf`.**
Para obtener instrucciones, consulte la documentación incluida con el software de TrueCopy.
- 4 **Inicie el daemon CCI de TrueCopy. Para ello, ejecute el comando `horcmstart.sh` en todos los nodos.**
`# /usr/bin/horcmstart.sh`
- 5 **Si aún no ha creado los pares de réplicas, hágalo ahora.**
Utilice el comando `paircreate` para crear los pares de réplicas con el nivel de aislamiento deseado. Para obtener instrucciones sobre la creación de los pares de réplicas, consulte la documentación de TrueCopy.
- 6 **En cada nodo configurado con dispositivos replicados, verifique que la replicación de datos esté configurada correctamente con el comando `pairdisplay`. Un grupo de dispositivos de Hitachi**

TrueCopy o Hitachi Universal Replicator con un `fence_level` de ASYNC no puede compartir su `ctgid` con ningún otro grupo de dispositivos del sistema.

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

7 Verifique que todos los nodos pueden controlar los grupos de replicación.

- a. Determine qué nodo contiene la réplica principal y qué nodo contiene la réplica secundaria mediante el comando `pairdisplay`.**

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

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

- b. Ejecute el comando `horctakeover` en el nodo que contiene la réplica secundaria para convertir el nodo secundario en el nodo maestro.**

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

Espere hasta que la copia de datos inicial se complete antes de continuar con el paso siguiente.

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

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

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

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

- e. Verifique que el nodo principal haya regresado a la configuración original mediante la ejecución del comando `pairdisplay`.**

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

Pasos siguientes Continúe la configuración del dispositivo replicado. Para ello, siga las instrucciones detalladas en [“Configuración de dispositivos DID para la replicación con Hitachi TrueCopy” en la página 101.](#)

▼ Configuración de dispositivos DID para la replicación con Hitachi TrueCopy

Antes de empezar

Después de configurar un grupo de dispositivos para el dispositivo replicado, debe configurar el controlador del identificador de dispositivos (DID) que usa el dispositivo replicado.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**

- 2 **Verifique que el daemon `horcmd` se ejecuta en todos los nodos.**

El siguiente comando iniciará el daemon si no se está ejecutando. El sistema mostrará un mensaje si el daemon ya está en ejecución.

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

- 3 **Determine qué nodo contiene la réplica secundaria mediante el comando `pairdisplay`.**

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

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

- 4 **En el nodo con la réplica secundaria (según se determinó en el paso anterior), configure los dispositivos DID que se utilizarán con la replicación basada en almacenamiento.**

Este comando combina las dos instancias DID independientes para los pares de réplicas del dispositivo en una instancia DID lógica única. La instancia única permite que el software de administración de volúmenes utilice el dispositivo desde ambos lados.



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

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

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

```
-S
```

Especifica un nodo de origen distinto del nodo actual.

secondary replica-nodename

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

Nota – De forma predeterminada, el nodo actual es el nodo de origen. Utilice la opción -S para especificar un nodo de origen diferente.

5 Verifique que las instancias DID se hayan combinado.

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

6 Verifique que la replicación de TrueCopy esté definida.

```
# cldevice show logical_DID_device
```

La salida del comando debe indicar que el tipo de replicación es TrueCopy.

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



Precaución – Tenga mucho cuidado cuando combine instancias DID manualmente. La reasignación inadecuada de dispositivos puede causar daños en los datos.

a. En todos los nodos que contienen la réplica secundaria, ejecute el comando cldevice combine.

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

-d La instancia DID remota, que corresponde a la réplica principal.
destination-instance

source-instance La instancia DID local, que corresponde a la réplica secundaria.

b. Verifique que la reasignación DID se ha producido correctamente.

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

Una de las instancias DID no se debe enumerar.

8 En todos los nodos, compruebe que se pueda acceder a los dispositivos DID para todas las instancias DID combinadas.

```
# cldevice list -v
```

Pasos siguientes

Para completar la configuración del grupo de dispositivos replicado, lleve a cabo los pasos indicados en los siguientes procedimientos.

- [“Adición y registro de un grupo de dispositivos \(Solaris Volume Manager\)” en la página 130](#)

Al registrar el grupo de dispositivos, asegúrese de asignarle el mismo nombre que el grupo de replicación de TrueCopy.

- “Verificación de la configuración de un grupo de dispositivos globales replicados de Hitachi TrueCopy” en la página 103

▼ Verificación de la configuración de un grupo de dispositivos globales replicados de Hitachi TrueCopy

Antes de empezar

Antes de verificar el grupo de dispositivos globales, primero debe crearlo. Puede utilizar grupos de dispositivos de Solaris Volume Manager, ZFS o disco sin procesar. Para obtener más información, consulte lo siguiente:

- “Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 130
- “Adición y registro de un grupo de dispositivos (disco básico)” en la página 132
- “Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 133



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.

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

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

- 1 Compruebe que el grupo de dispositivos principal corresponda al mismo nodo que el nodo que contiene la réplica primaria.

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

- 2 Verifique que la propiedad de replicación esté definida para el grupo de dispositivos.

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

- 3 Verifique que la propiedad replicada esté definida para el dispositivo.

```
# usr/cluster/bin/cldevice status [-s state] [-n node[,?]] [+| [disk-device ]]
```

- 4 Realice una conmutación de prueba para asegurarse de que estos grupos de dispositivos están configurados correctamente y de que las replicas puedan moverse entre los nodos.

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

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

`-n nodename` Nodo en el que se cambia el grupo de dispositivos. Este nodo se convierte en el nuevo nodo principal

5 Compruebe que la operación de conmutación se haya realizado correctamente mediante la comparación de la salida de los siguientes comandos.

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

Ejemplo: configuración de un grupo de replicación de TrueCopy para Oracle Solaris Cluster

En este ejemplo, se completan los pasos específicos de Oracle Solaris Cluster necesarios para configurar la replicación de TrueCopy en el cluster. En el ejemplo, se supone que ya se han realizado las tareas siguientes:

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

Para obtener instrucciones sobre la configuración de los pares de replicación, consulte [“Configuración de un grupo de replicación de Hitachi TrueCopy” en la página 99](#).

En este ejemplo, se utiliza un cluster de tres nodos con TrueCopy. El cluster se distribuye entre dos sitios remotos, con dos nodos en un sitio y un nodo en el otro. Cada sitio tiene su propio dispositivo de almacenamiento de Hitachi.

En los siguientes ejemplos, se muestra el archivo de configuración /etc/horcm.conf de TrueCopy en cada nodo.

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

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

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

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

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

```

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

```

En los ejemplos anteriores, se replican tres LUN entre los dos sitios. Los LUN están todos en un grupo de replicación con el nombre VG01. El comando `pairdisplay` verifica esta información y muestra que el nodo 3 tiene la réplica principal.

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

```

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

```

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

```

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

```

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

```

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

```

Para ver qué discos están en uso, utilice la opción `-fd` del comando `pairdisplay` como se muestra en los siguientes ejemplos.

EJEMPLO 5-7 Salida del comando `pairdisplay` en el nodo 1, que muestra los discos utilizados

```
# pairdisplay -fd -g VG01
Group PairVol(L/R) Device File ,Seq#,LDEV#.P/S,Status,Fence,Seq#,P-LDEV# M
VG01 pair1(L) c6t500060E80000000000000000EEBA0000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair2(L) c6t500060E80000000000000000EEBA0000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R) c5t50060E80000000000000004E600000003Bd0s2 0064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L) c6t500060E80000000000000000EEBA0000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -
```

EJEMPLO 5-8 Salida del comando `pairdisplay` en el nodo 2, que muestra los discos utilizados

```
# pairdisplay -fd -g VG01
Group PairVol(L/R) Device File ,Seq#,LDEV#.P/S,Status,Fence,Seq#,P-LDEV# M
VG01 pair1(L) c5t500060E80000000000000000EEBA0000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair1(R) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair2(L) c5t500060E80000000000000000EEBA0000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair2(R) c5t50060E80000000000000004E600000003Bd0s2 20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair3(L) c5t500060E80000000000000000EEBA0000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -
VG01 pair3(R) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -
```

EJEMPLO 5-9 Salida del comando `pairdisplay` en el nodo 3, que muestra los discos utilizados

```
# pairdisplay -fd -g VG01
Group PairVol(L/R) Device File ,Seq#,LDEV#.P/S,Status,Fence ,Seq#,P-LDEV# M
VG01 pair1(L) c5t50060E80000000000000004E600000003Ad0s2 20064 58..P-VOL PAIR NEVER ,61114 29 -
VG01 pair1(R) c6t500060E80000000000000000EEBA0000001Dd0s2 61114 29..S-VOL PAIR NEVER ,----- 58 -
VG01 pair2(L) c5t50060E80000000000000004E600000003Bd0s2 20064 59..P-VOL PAIR NEVER ,61114 30 -
VG01 pair2(R) c6t500060E80000000000000000EEBA0000001Ed0s2 61114 30..S-VOL PAIR NEVER ,----- 59 -
VG01 pair3(L) c5t50060E80000000000000004E600000003Cd0s2 20064 60..P-VOL PAIR NEVER ,61114 31 -
VG01 pair3(R) c6t500060E80000000000000000EEBA0000001Fd0s2 61114 31..S-VOL PAIR NEVER ,----- 60 -
```

En estos ejemplos, se muestra que se están utilizando los siguientes discos:

- En el nodo 1:
 - c6t500060E80000000000000000EEBA0000001Dd0s2
 - c6t500060E80000000000000000EEBA0000001Ed0s2
 - c6t500060E80000000000000000EEBA0000001Fd0s
- En el nodo 2:
 - c5t500060E80000000000000000EEBA0000001Dd0s2
 - c5t500060E80000000000000000EEBA0000001Ed0s2
 - c5t500060E80000000000000000EEBA0000001Fd0s2
- En el nodo 3:
 - c5t50060E80000000000000004E600000003Ad0s2
 - c5t50060E80000000000000004E600000003Bd0s2
 - c5t50060E80000000000000004E600000003Cd0s2

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

EJEMPLO 5-10 Visualización de DID correspondientes a los discos utilizados**# cldevice list -v**

```

DID Device  Full Device Path
-----
1           node-1:/dev/rdisk/c0t0d0  /dev/did/rdsk/d1
2           node-1:/dev/rdisk/c0t6d0  /dev/did/rdsk/d2
11          node-1:/dev/rdisk/c6t500060E8000000000000EBA00000020d0  /dev/did/rdsk/d11
11          node-2:/dev/rdisk/c5t500060E8000000000000EBA00000020d0  /dev/did/rdsk/d11
12          node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Fd0  /dev/did/rdsk/d12
12          node-2:/dev/rdisk/c5t500060E8000000000000EBA0000001Fd0  /dev/did/rdsk/d12
13          node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Ed0  /dev/did/rdsk/d13
13          node-2:/dev/rdisk/c5t500060E8000000000000EBA0000001Ed0  /dev/did/rdsk/d13
14          node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Dd0  /dev/did/rdsk/d14
14          node-2:/dev/rdisk/c5t500060E8000000000000EBA0000001Dd0  /dev/did/rdsk/d14
18          node-3:/dev/rdisk/c0t0d0  /dev/did/rdsk/d18
19          node-3:/dev/rdisk/c0t6d0  /dev/did/rdsk/d19
20          node-3:/dev/rdisk/c5t50060E800000000000004E6000000013d0  /dev/did/rdsk/d20
21          node-3:/dev/rdisk/c5t50060E800000000000004E600000003Dd0  /dev/did/rdsk/d21
22          node-3:/dev/rdisk/c5t50060E800000000000004E600000003Cd0  /dev/did/rdsk/d2223
23          node-2:/dev/rdisk/c5t50060E800000000000004E600000003Bd0  /dev/did/rdsk/d23
24          node-3:/dev/rdisk/c5t50060E800000000000004E600000003Ad0  /dev/did/rdsk/d24

```

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

En el siguiente ejemplo, se muestra el resultado de combinar instancias DID mediante la ejecución del comando en el nodo 1.

EJEMPLO 5-11 Combinación de instancias DID

```

# cldevice replicate -D node-3
Remapping instances for devices replicated with node-3...
VG01 pair1 L node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Dd0
VG01 pair1 R node-3:/dev/rdisk/c5t50060E800000000000004E600000003Ad0
Combining instance 14 with 24
VG01 pair2 L node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Ed0
VG01 pair2 R node-3:/dev/rdisk/c5t50060E800000000000004E600000003Bd0
Combining instance 13 with 23
VG01 pair3 L node-1:/dev/rdisk/c6t500060E8000000000000EBA0000001Fd0
VG01 pair3 R node-3:/dev/rdisk/c5t50060E800000000000004E600000003Cd0
Combining instance 12 with 22

```

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

EJEMPLO 5-12 Visualización de DID combinados

```
# cldevice list -v
DID Device Full Device Path
-----
1 node-1:/dev/rdisk/c0t0d0 /dev/did/rdsk/d1
2 node-1:/dev/rdisk/c0t6d0 /dev/did/rdsk/d2
11 node-1:/dev/rdisk/c6t500060E8000000000000E000000020d0 /dev/did/rdsk/d11
11 node-2:/dev/rdisk/c5t500060E8000000000000E000000020d0 /dev/did/rdsk/d11
18 node-3:/dev/rdisk/c0t0d0 /dev/did/rdsk/d18
19 node-3:/dev/rdisk/c0t6d0 /dev/did/rdsk/d19
20 node-3:/dev/rdisk/c5t500060E8000000000000004E6000000013d0 /dev/did/rdsk/d20
21 node-3:/dev/rdisk/c5t500060E8000000000000004E600000003Dd0 /dev/did/rdsk/d21
22 node-1:/dev/rdisk/c6t500060E800000000000000E00000001Fd0 /dev/did/rdsk/d1222
22 node-2:/dev/rdisk/c5t500060E800000000000000E00000001Fd0 /dev/did/rdsk/d12
22 node-3:/dev/rdisk/c5t500060E8000000000000004E600000003Cd0 /dev/did/rdsk/d22
23 node-1:/dev/rdisk/c6t500060E800000000000000E00000001Ed0 /dev/did/rdsk/d13
23 node-2:/dev/rdisk/c5t500060E800000000000000E00000001Ed0 /dev/did/rdsk/d13
23 node-3:/dev/rdisk/c5t500060E8000000000000004E600000003Bd0 /dev/did/rdsk/d23
24 node-1:/dev/rdisk/c6t500060E800000000000000E00000001Dd0 /dev/did/rdsk/d24
24 node-2:/dev/rdisk/c5t500060E800000000000000E00000001Dd0 /dev/did/rdsk/d24
24 node-3:/dev/rdisk/c5t500060E8000000000000004E600000003Ad0 /dev/did/rdsk/d24
```

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

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

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

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

Drive Dbase
d22 Yes
d23 Yes
d24 Yes
```

En este punto, el grupo de dispositivos se puede utilizar, es posible crear metadispositivos y el grupo de dispositivos se puede mover a cualquiera de los tres nodos. No obstante, para optimizar los switchovers y las conmutaciones por error, ejecute `cldevicegroup set` para marcar el grupo de dispositivos como replicado en la configuración del cluster.

EJEMPLO 5-14 Optimización de los switchovers y las conmutaciones por error

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

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

La configuración del grupo de replicación ha terminado con este paso. Para verificar que la configuración se realizó correctamente, siga los pasos descritos en [“Verificación de la configuración de un grupo de dispositivos globales replicados de Hitachi TrueCopy” en la página 103](#).

Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility

En la siguiente tabla, se enumeran las tareas que debe realizar para configurar y gestionar un dispositivo replicado basado en almacenamiento de EMC Symmetrix Remote Data Facility (SRDF).

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

Tarea	Instrucciones
Instalar el software de SRDF en el dispositivo de almacenamiento y en los nodos	La documentación incluida con el dispositivo de almacenamiento de EMC.
Configurar el grupo de replications de EMC	“Cómo configurar un grupo de replicación de EMC SRDF” en la página 110
Configurar el dispositivo DID	“Configuración de dispositivos DID para la replicación con EMC SRDF” en la página 112
Registrar el grupo replicado	“Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 130
Comprobar la configuración	“Cómo comprobar la configuración del grupo de dispositivos globales replicados de EMC SRDF” en la página 114

TABLA 5-3 Mapa de tareas: Administración de dispositivo replicado basado en almacenamiento de EMC SRDF (Continuación)

Tarea	Instrucciones
Recuperar manualmente los datos después de que falla por completo la sala primaria de un cluster de campus.	“Cómo recuperar los datos de EMC SRDF después de un fallo completo de la sala primaria” en la página 119

▼

Antes de empezar

Cómo configurar un grupo de replicación de EMC SRDF

El software EMC Solutions Enabler debe estar instalado en todos los nodos del cluster antes de configurar un grupo de replications de EMC Symmetrix Remote Data Facility (SRDF). En primer lugar, configure los grupos de dispositivos EMC SRDF en los discos compartidos del cluster. Para obtener más información acerca de cómo configurar los grupos de dispositivos de EMC SRDF, consulte la documentación del producto EMC SRDF.

Cuando use EMC SRDF, utilice los dispositivos dinámicos en lugar de los estáticos. Los dispositivos estáticos requieren varios minutos para cambiar la replicación principal y pueden afectar el tiempo de conmutación por error.



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 un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos conectados a la matriz de almacenamiento.**
- 2
- Para una implementación de tres sitios o de tres centros de datos con dispositivos en cascada o SRDF simultáneos, defina el parámetro `SYMAPI_2SITE_CLUSTER_DG`.**

Agregue la siguiente entrada al archivo de opciones de Solutions Enabler en todos los nodos participantes del cluster:

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

`device-group`

Especifica el nombre del grupo de dispositivos.

`rdf-group-number`

Especifica el grupo RDF que se conecta el symmetrix local del host al symmetrix del segundo sitio.

Esta entrada permite que el software del cluster automatice el movimiento de la aplicación entre los dos sitios SRDF síncronos.

Para obtener más información sobre las configuraciones de tres centros de datos, consulte [“Three-Data-Center \(3DC\) Topologies” de Oracle Solaris Cluster Geographic Edition Overview](#).
- 3
- En cada nodo configurado con los datos replicados, detecte la configuración de dispositivos de Symmetrix.**

Esta acción puede tardar varios minutos.

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

4 Si aún no ha creado los pares de réplicas, hágalo ahora.

Utilice el comando `symrdf` para crear los pares de réplicas. Para obtener instrucciones sobre la creación de los pares de réplicas, consulte la documentación de SRDF.

Nota – Si utiliza dispositivos RDF simultáneos para una implementación de tres sitios o de tres centros de datos, agregue el siguiente parámetro a todos los comandos `symrdf`:

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

Especificar el número de grupo RDF en el comando `symrdf` garantiza que la operación de `symrdf` se dirija al grupo RDF correcto.

5 En cada nodo configurado con dispositivos replicados, compruebe que la replicación de datos esté configurada de manera correcta.

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

6 Realice un intercambio del grupo de dispositivos.**a. Verifique que la réplica principal y la secundaria estén sincronizadas.**

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

b. Determine qué nodo contiene la réplica principal y qué nodo contiene la réplica secundaria mediante el comando `symdg show`.

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

El nodo con el dispositivo RDF1 contiene la réplica principal y el nodo con el estado del dispositivo RDF2 contiene la réplica secundaria.

c. Active la réplica secundaria.

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

d. Intercambie los dispositivos RDF1 y RDF2.

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

e. Active el par de réplicas.

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

f. Verifique que el nodo principal y las réplicas secundarias estén sincronizados.

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

7 Repita todo el paso 5 en el nodo que tenía originalmente la réplica principal.

Pasos siguientes Después de configurar un grupo de dispositivos para el dispositivo replicado de EMC SRDF, debe configurar el controlador del identificador de dispositivos (DID) que usan los dispositivos replicados.

▼ **Configuración de dispositivos DID para la replicación con EMC SRDF**

Este procedimiento sirve para configurar el controlador del identificador de dispositivos (DID) que usa el dispositivo replicado.

Antes de empezar `phys - schost#` refleja un indicador de cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Determine qué dispositivos DID se corresponden con los dispositivos RDF1 y RDF2 configurados.**

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

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

- 3 **Determine qué dispositivos DID se corresponden con los dispositivos de Oracle Solaris.**
- 4 **Para cada par de dispositivos DID coincidentes, combine las instancias en un único dispositivo DID replicado. Ejecute el siguiente comando desde el lado RDF2/secundario.**

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

Nota – No se admite la opción `-T` para los dispositivos de replicación de datos de SRDF.

-t <i>replication-type</i>	Especifica el tipo de replicación. Para EMC SRDF, el tipo es SRDF .
-g <i>replication-device-group</i>	Especifica el nombre del grupo de dispositivos tal y como se muestra en el comando <code>symdg show</code> .

<code>-d destination-instance</code>	Especifica la instancia DID que corresponde al dispositivo RDF1.
<code>source-instance</code>	Especifica la instancia DID que corresponde al dispositivo RDF2.

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

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

- 5 Si cambia el nombre de un grupo de dispositivos de replicación, se necesitan pasos adicionales para Hitachi TrueCopy y SRDF. Cuando haya completado los pasos del 1 al 4, lleve a cabo los pasos adicionales que correspondan.**

Elemento	Descripción
TrueCopy	Si cambia el nombre del grupo de dispositivos de replicación (y el grupo de dispositivos globales correspondiente), debe volver a ejecutar el comando <code>cldevice replicate</code> para actualizar la información de los dispositivos replicados.
SRDF	Si cambia el nombre del grupo de dispositivos de replicación (con el grupo de dispositivos globales correspondiente), debe actualizar la información de los dispositivos replicados usando primero el comando <code>scdidadm -b</code> para eliminar la información existente. El último paso consiste en utilizar el comando <code>cldevice combine</code> para crear un nuevo dispositivo actualizado.

- 6 Verifique que las instancias DID se hayan combinado.**

`cldevice list -v device`

- 7 Verifique que la replicación de SRDF esté definida.**

`cldevice show device`

- 8 En todos los nodos, compruebe que se pueda acceder a los dispositivos DID para todas las instancias DID combinadas.**

`cldevice list -v`

Pasos siguientes Después de haber configurado el controlador del identificador de dispositivos (DID) que usa el dispositivo replicado, debe comprobar la configuración del grupo de dispositivos globales replicados de EMC SRDF.

▼ Cómo comprobar la configuración del grupo de dispositivos globales replicados de EMC SRDF

Antes de empezar

Antes de verificar el grupo de dispositivos globales, primero debe crearlo. Puede utilizar grupos de dispositivos de Solaris Volume Manager, ZFS o disco sin procesar. Para obtener más información, consulte lo siguiente:

- “Adición y registro de un grupo de dispositivos (Solaris Volume Manager)” en la página 130
- “Adición y registro de un grupo de dispositivos (disco básico)” en la página 132
- “Adición y registro de un grupo de dispositivos replicado (ZFS)” en la página 133



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 nombre del grupo de dispositivos replicados.

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

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

- 1 **Compruebe que el grupo de dispositivos principal corresponda al mismo nodo que el nodo que contiene la réplica primaria.**

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

- 2 **Realice una conmutación de prueba para asegurarse de que estos grupos de dispositivos están configurados correctamente y de que las replicas puedan moverse entre los nodos.**

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

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

`-n nodename` Nodo en el que se cambia el grupo de dispositivos. Este nodo se convierte en el nuevo nodo principal.

- 3 **Compruebe que la operación de conmutación se haya realizado correctamente mediante la comparación de la salida de los siguientes comandos.**

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

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

En este ejemplo, se completan los pasos específicos necesarios de Oracle Solaris Cluster para configurar la replicación de SRDF en el cluster. En el ejemplo, se supone que ya se han realizado las tareas siguientes:

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

En este ejemplo, se muestra un cluster de cuatro nodos en el que hay dos nodos conectados a un Symmetrix y otros dos nodos conectados al segundo Symmetrix. El grupo de dispositivos SRDF se denomina dg1.

EJEMPLO 5-15 Creación de los pares de réplicas

Ejecute el siguiente comando en todos los nodos.

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

Symmetrix ID: 000187990182

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

En todos los nodos del lado RDF1, escriba:

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

En todos los nodos del lado RDF2, escriba:

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

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

Desde un nodo del cluster, escriba:

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

```
# symdg show dg1

Group Name:  dg1

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

Number of STD Devices in Group : 1
Number of Associated GK's      : 0
Number of Locally-associated BCV's : 0
Number of Locally-associated VDEV's : 0
Number of Remotely-associated BCV's (STD RDF): 0
Number of Remotely-associated BCV's (BCV RDF): 0
Number of Remotely-assoc'd RBCV's (RBCV RDF) : 0

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

Device Group RDF Information
...
# symrdf -g dg1 establish

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

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

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

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

#
# symrdf -g dg1 query
```

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

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

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

DEV001	0067	WD	0	0	RW	0067	RW	0	0	S..	Synchronized
Total		-----					-----				
MB(s)		0.0					0.0			0.0	

Legend for MODES:

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

#

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

El mismo procedimiento se aplica a los lados RDF1 y RDF2.

Puede buscar en el campo PdevName de salida del comando dymdg show dg.

En el lado RDF1, escriba:

```
# symdg show dg1
Group Name: dg1
Group Type : RDF1 (RDFA)
...
Standard (STD) Devices (1):
{
-----
LdevName          PdevName          Sym   Cap
Dev  Att. Sts      Dev  (MB)
-----
DEV001            /dev/rdisk/c5t6006048000018790002353594D303637d0s2 0067    RW    4315
}
Device Group RDF Information
...
```

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

Para obtener los DID correspondientes, escriba:

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

Para listar los DID correspondientes, escriba:

```
# cldevice show d217

=== DID Device Instances ===

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

#
```

En el lado RDF2, escriba:

Puede buscar en el campo PdevName de salida del comando dymdg show dg.

```
# symdg show dg1

Group Name:  dg1

Group Type                :  RDF2      (RDFA)

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

Device Group RDF Information
...
```

Para obtener los DID correspondientes, escriba:

```
# scdidadm -L | grep c5t6006048000018799018253594D303637d0
108      pmoney4:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdisk/d108
108      pmoney3:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdisk/d108
#
```

Para listar los DID correspondientes, escriba:

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

```
# cldevice show d108

=== DID Device Instances ===

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

#
```

EJEMPLO 5-18 Combinación de instancias DID

Desde el lado RDF2, escriba:

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

EJEMPLO 5-19 Visualización de DID combinados

Desde cualquier nodo del cluster, escriba:

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

=== DID Device Instances ===

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

#
```

▼ **Cómo recuperar los datos de EMC SRDF después de un fallo completo de la sala primaria**

En este procedimiento, se realiza la recuperación de datos cuando la sala primaria de un cluster de campus ha fallado por completo, la sala primaria conmuta por error a la sala secundaria y, luego, se vuelve a conectar. La sala primaria de un cluster de campus es el nodo principal y el sitio de almacenamiento. Todo el fallo de una sala incluye el fallo del host y del almacenamiento en dicha sala. Si la sala primaria sala falla, Oracle Solaris Cluster pasa automáticamente conmuta por error automáticamente a la sala secundaria, hace que se pueda leer y escribir en el dispositivo de almacenamiento de la sala secundaria, y permite la conmutación por error de los grupos de dispositivos correspondientes y grupos de dispositivos y recursos correspondientes.

Cuando la sala primaria vuelve a conectarse, puede recuperar manualmente los datos del grupo de dispositivos SRDF que se hayan escrito en la sala secundaria y vuelva a sincronizar los datos. En este procedimiento, se recupera el grupo de dispositivos SRDF mediante la sincronización de los datos de la sala de secundaria original (en este procedimiento, se utiliza *phys-campus-2* para la sala secundaria) a la sala primaria original (*phys-campus-1*). También se cambia el tipo de grupo de dispositivos SRDF a RDF1 en *phys-campus-2* y RDF2 en *phys-campus-1*.

Antes de empezar

Debe configurar el grupo de dispositivos DID y replicaciones EMC, además de registrar grupo de replicaciones EMC antes de realizar una conmutación por error manual. Para obtener información sobre cómo crear un grupo de dispositivos de Solaris Volume Manager, consulte [“Adición y registro de un grupo de dispositivos \(Solaris Volume Manager\)” en la página 130.](#)

Nota – En estas instrucciones, se muestra un método que puede utilizar para recuperar datos de SRDF manualmente después de que la sala primaria conmuta por error completamente y, luego, se vuelve a conectar. Compruebe la documentación de EMC para obtener más métodos.

Inicie sesión en la sala primaria de un cluster de campus para realizar estos pasos. En el procedimiento que se indica a continuación, *dg1* es el nombre de grupo de dispositivos SRDF. En el momento del fallo, la sala primaria de este procedimiento es *phys-campus-1* y la segunda sala es *phys-campus-2*.

- 1 Inicie sesión en la sala principal de un cluster de campus, y conviértase en superusuario o asuma un rol que proporcione la autorización `RBAC solaris.cluster.modify`.
- 2 De la sala primaria, utilice el comando `symrdf` para consultar el estado de la replicación de dispositivos RDF y ver la información acerca de los dispositivos.

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

Consejo – Un grupo que está en el estado dividido no está sincronizado.

- 3 Si el par de RDF tiene estado dividido y el tipo de grupo de dispositivos es RDF1, fuerce una conmutación por error del grupo de dispositivos SRDF.

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

- 4 Vea el estado de los dispositivos RDF.

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

- 5 Después de la conmutación por error, puede intercambiar los datos de los dispositivos RDF que tuvieron conmutación por error.

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

6 Verifique el estado y otra información sobre dispositivos RDF.

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

7 Establezca el grupo de dispositivos SRDF en la sala primaria.

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

8 Confirme si el grupo de dispositivos está en un estado sincronizado y si el tipo de grupo de dispositivos es RDF2.

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

Ejemplo 5–20 Recuperación manual de datos de EMC SRDF después de una conmutación por error del sitio principal

En este ejemplo, se proporcionan los pasos específicos de Oracle Solaris Cluster que son necesarios para recuperar manualmente los datos de EMC SRDF después de que una sala primaria de un cluster de campus conmuta por error, una sala secundaria toma su lugar y registra los datos y, luego, la sala primaria vuelve a estar en línea. En el ejemplo, el grupo de dispositivos SRDF se denomina *dg1* y el dispositivo lógico estándar es DEV001. La sala primaria es *phys-campus-1* en el momento del fallo, y la sala secundaria es *phys-campus-2*. Siga los pasos que se indican en la sala primaria de un cluster de campus, *phys-campus-1*.

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

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

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

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

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

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

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

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

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

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

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

Información general sobre la administración de sistemas de archivos de clusters

No se necesita ningún comando especial de Oracle Solaris Cluster para la administración de sistemas de archivos de cluster. Un sistema de archivos de cluster 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 cluster si especifica la opción `-g` en el comando `mount`. Los sistemas de archivos de cluster también se pueden montar automáticamente en el inicio. Los sistemas de archivos de cluster sólo son visibles desde el nodo de votación de un cluster global. Si necesita poder tener acceso a los datos de sistemas de archivos de cluster desde un nodo sin voto, asigne los datos al nodo sin voto con [zoneadm\(1M\)](#) o `HASStoragePlus`.

Nota – Cuando el sistema de archivos de cluster lee archivos, no actualiza la hora de acceso a tales archivos.

Restricciones del sistema de archivos de cluster

La administración del sistema de archivos de cluster presenta las restricciones siguientes:

- El comando [unlink\(1M\)](#) no se admite en los directorios que no están vacíos.
- No se admite el comando `lockfs -d`. Como solución alternativa, utilice el comando `lockfs -n`.
- No puede volver a montar un sistema de archivos de cluster con la opción de montaje `directio` agregada en el momento en que el montaje se realiza de nuevo.
- Se admite ZFS para sistemas de archivos root, con una excepción importante. Si emplea una partición dedicada del disco de arranque para el sistema de archivos de los dispositivos globales, como sistema de archivos se debe usar sólo UFS. El espacio de nombre de dispositivos globales necesita que el sistema de archivos proxy (PxFS) se ejecute en un sistema de archivos UFS. Sin embargo, un sistema de archivos UFS para el espacio de 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.

Administración de grupos de dispositivos

A medida que cambien las necesidades del cluster, tal vez necesite agregar, quitar o modificar los grupos de dispositivos de dicho cluster. 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 cluster que se arranque fuera del cluster si hay otros nodos que son miembros activos del cluster 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 cluster. Sin embargo, los grupos de dispositivos del cluster permanecen fuera de línea hasta que se acceda a dichos grupos como dispositivos globales.

TABLA 5-4 Mapa de tareas: administrar grupos de dispositivos

Tarea	Instrucciones
Actualizar el espacio de nombre los dispositivos globales sin un rearranque de reconfiguración con el comando <code>cldevice populate</code>	“Actualización del espacio de nombre de dispositivos globales” en la página 125
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 126
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 127 “Migración del espacio de nombre de dispositivos globales desde un dispositivo de <code>lofi</code> hasta una partición dedicada” en la página 129
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 130

TABLA 5-4 Mapa de tareas: administrar grupos de dispositivos (Continuación)

Tarea	Instrucciones
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 132
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 133
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 135
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 135
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 136
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 138
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 140
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 144
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 142
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 146
Poner un grupo de dispositivos en estado de mantenimiento con el comando <code>metaset</code> o <code>vxdg</code>	“Colocación de un grupo de dispositivos en estado de mantenimiento” en la página 147

▼ 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 cluster. Tampoco tiene ningún efecto si el sistema de archivos `/global/.devices/node@nodeID` 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 cluster.**
- 2 **En cada nodo del cluster, ejecute el comando `devfsadm(1M)`.**
Este comando puede ejecutarse simultáneamente en todos los nodos del cluster.
- 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 forma 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 cluster.

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

Ejemplo 5–21 Actualización del espacio de nombre de dispositivos globales

El ejemplo siguiente muestra la salida generada al ejecutarse correctamente el comando `cldevice populate`.

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

▼ Cómo cambiar el tamaño de un dispositivo lofi que se utiliza para el espacio de nombres de dispositivos globales

Si utiliza un dispositivo `lofi` para el espacio de nombres de dispositivos globales en uno o más nodos del cluster 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 cluster.**

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 cluster” en la página 82](#).

- 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 cluster.

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 nombres en un dispositivo de interfaz de archivo en bucle de retorno (lofi), en lugar de crear uno para dispositivos globales en una partición dedicada. Esta función es útil si instala el software de Oracle Solaris Cluster en sistemas que vienen preinstalados con el sistema operativo Oracle Solaris 10.

Nota – Se admite ZFS para sistemas de archivos root, con una excepción importante. Si emplea una partición dedicada del disco de arranque para el sistema de archivos de los dispositivos globales, como sistema de archivos se debe usar sólo UFS. El espacio de nombre de dispositivos globales necesita que el sistema de archivos proxy (PxFs) se ejecute en un sistema de archivos UFS. Sin embargo, un sistema de archivos UFS para el espacio de 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 127
- “Migración del espacio de nombre de dispositivos globales desde un dispositivo de lofi hasta una partición dedicada” en la página 129

▼ 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 cluster 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 cluster.

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 cluster” en la página 82.](#)

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 tiene una ruta de montaje que comienza con `/global/.devices/node@nodeID`.**6 Desmonte la partición de dispositivos globales `/global/.devices/node@nodeID`.****7 Inhabilite y vuelva a habilitar los servicios de SMF `globaldevices` y `scmountdev`.**

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

A continuación, se crea un dispositivo de lofi en `/.globaldevices` y se monta en el sistema de archivos de dispositivos globales.

8 Repita estos pasos en otros nodos cuyo espacio de nombre de dispositivos globales quiera migrar desde una partición hasta un dispositivo de lofi.**9 Desde un nodo, complete los espacios de nombre de dispositivos globales.**

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

Antes de llevar a cabo ninguna otra acción en el cluster, 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 cluster 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 cluster.
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 cluster” en la página 82.](#)
- 3 En un disco local del nodo, cree una nueva partición que cumpla con los requisitos siguientes:
 - Un tamaño mínimo de 512 MB
 - Usa el sistema de archivos UFS
- 4 Agregue una entrada al archivo `/etc/vfstab` para que la nueva partición se monte como sistema de archivos de dispositivos globales.
 - Determine el ID de nodo del nodo actual.
`# /usr/sbin/clinfo -nnode ID`
 - Cree la nueva entrada en el archivo `/etc/vfstab` con el formato siguiente:
`blockdevice rawdevice /global/.devices/node@nodeID ufs 2 no global`
Por ejemplo, si la partición que elige es `/dev/did/rdisk/d5s3`, la entrada nueva que se agrega al archivo `/etc/vfstab` debe ser: `/dev/did/dsk/d5s3 /dev/did/rdisk/d5s3 /global/.devices/node@3 ufs 2 no global`
- 5 Desmonte la partición de los dispositivos globales `/global/.devices/node@ID_nodo`.
- 6 Quite el dispositivo de lofi asociado con el archivo `/globaldevices`.
`# lofiadm -d /globaldevices`
- 7 Elimine el archivo `/globaldevices`.
`# rm /globaldevices`
- 8 Inhabilite y vuelva a habilitar los servicios de SMF `globaldevices` y `scmountdev`.
`# svcadm disable globaldevices# svcadm disable scmountdev`
`# svcadm enable scmountdev`
`# svcadm enable globaldevices`

Ahora la partición está montada como sistema de archivos del espacio de nombre de dispositivos globales.

- 9 **Repita estos pasos en otros nodos cuyo espacio de nombre de dispositivos globales desee migrar desde un dispositivo de `lofi` hasta una partición.**
- 10 **Inicie en modo de cluster. Desde un nodo del cluster, ejecute el comando `cldevice populate` para rellenar el espacio de nombres de dispositivos globales.**

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

Antes de pasar a realizar ninguna otra acción en cualquiera de los nodos, compruebe que el proceso concluya en todos los nodos del cluster.

```
# 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 diskset` Especifica el conjunto de discos que se va a crear.

`-a -h nodelist` 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 cluster, 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 142](#) 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 cluster 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:	<code>/dev/did/rdisk/d1</code>
Full Device Path:	<code>phys-schost-1:/dev/rdisk/c0t0d0</code>
DID Device Name:	<code>/dev/did/rdisk/d2</code>
Full Device Path:	<code>phys-schost-1:/dev/rdisk/c0t6d0</code>
DID Device Name:	<code>/dev/did/rdisk/d3</code>
Full Device Path:	<code>phys-schost-1:/dev/rdisk/c1t1d0</code>

```
Full Device Path:          phys-schost-2:/dev/rdisk/clt1d0
...
```

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 cluster, si se utiliza es posible que se prive al metaconjunto de la capacidad de conmutar a otra ubicación.

7 Compruebe el estado del conjunto de discos y de las unidades.

```
# metaset -s setname
```

Ejemplo 5–22 Adición a un grupo de dispositivos de Solaris Volume Manager

En el ejemplo siguiente se muestra la creación del conjunto de discos y el grupo de dispositivos con las unidades de disco /dev/did/rdisk/d1 y /dev/did/rdisk/d2; asimismo, se comprueba que se haya creado el grupo de dispositivos.

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

```
# cldevicegroup list
```

```
dg-schost-1
```

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

▼ Adición y registro de un grupo de dispositivos (disco básico)

El software Oracle Solaris Cluster admite el uso de grupos de dispositivos de discos básicos y otros administradores de volúmenes. Al configurar Oracle Solaris Cluster, los grupos de dispositivos se configuran automáticamente para cada dispositivo básico del cluster. Utilice este procedimiento para reconfigurar automáticamente estos grupos de dispositivos creados a fin de usarlos con el software Sun Oracle Solaris Cluster.

Puede crear un grupo de dispositivos de disco básico por los siguientes motivos:

- Desea agregar más de un DID al grupo de dispositivos.

- Necesita cambiar el nombre del grupo de dispositivos.
- Desea crear una lista de grupos de dispositivos sin utilizar la opción -v del comando cldg.



Precaución – Si crea un grupo de dispositivos en dispositivos replicados, el nombre del grupo de dispositivos que crea (Solaris Volume Manager 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.



Precaución – La compatibilidad total de ZFS con tecnologías de replicación de datos de terceros está pendiente. Consulte las últimas notas de la versión de Oracle Solaris Cluster para conocer las actualizaciones acerca de la compatibilidad de ZFS.

- 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 suprimir los dos grupos de dispositivos predeterminados llamados `d2` y `d13`.

```
# cldevicegroup offline dsk/d2 dsk/d13
# cldevicegroup delete 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; establezca la propiedad `globaldevicepaths` en un grupo de dispositivos del tipo disco sin procesar. Este grupo de dispositivos se creó en el Paso 2.**

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

- 6 **Si el grupo de recursos de aplicaciones se va a ejecutar en zonas locales, cree un nuevo grupo de recursos cuya lista de nodos contenga las zonas locales adecuadas. Las zonas globales que se correspondan con las zonas locales deben figurar en la lista de nodos del grupo de recursos creado en el Paso 4. Establezca el valor `+++` de la propiedad `rg_affinities` de este grupo de recursos al grupo de recursos creado en el Paso 4.**

```
# clrg create -n pnode1:zone-1,pnode2:zone-2 -p \
RG_affinities=+++migrate_truecopydg-rg sybase-rg
```

- 7 **Cree un recurso `HASToragePlus` (`hasp-rs`) para el zpool creado en el Paso 3 en el grupo de recursos creado en el Paso 4 o en el Paso 6. Establezca la propiedad `resource_dependencies` en el recurso `hasp-rs` creado en el Paso 5.**

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

- 8 **Use el nuevo nombre del grupo de recursos cuando se precise un nombre de grupo de dispositivos.**

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 cluster de todos los grupos de dispositivos que incluyen el nodo en sus listas de posibles nodos primarios.

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

Este procedimiento proporciona las formas largas de los comandos de 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 136](#) 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 Disco o Local_Disk, siga los pasos descritos en [“Eliminación de un nodo de un grupo de dispositivos de discos básicos” en la página 138](#) 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 cluster 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 cluster que se arranque fuera del cluster si hay otros nodos activos como miembros del cluster y al menos uno de ellos es propietario del conjunto de discos.

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

Este procedimiento proporciona las formas largas de los comandos de 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 5–23 Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)

En el ejemplo siguiente se muestra cómo se elimina el nombre de host `phys-schost-2` de la configuración de un grupo de dispositivos. En este ejemplo se elimina `phys-schost-2` como nodo primario potencial del grupo de dispositivos designado. Compruebe la eliminación del nodo; para ello, ejecute el comando `cldevicegroup show`. Compruebe que el nodo eliminado ya no aparezca en el texto de la pantalla.

```
[Determine the Solaris Volume Manager
device group for the node:]
# cldevicegroup show dg-schost-1
=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      SVM
failback:                  no
Node List:                 phys-schost-1, phys-schost-2
preferenced:               yes
numsecondaries:            1
diskset name:              dg-schost-1
[Determine which node is the current primary for the device group:]
# cldevicegroup status dg-schost-1
=== Cluster Device Groups ===

--- Device Group Status ---

Device Group Name  Primary          Secondary        Status
-----
dg-schost-1       phys-schost-1   phys-schost-2   Online
[Become superuser on the node that currently owns the device group.]
[Remove the host name from the device group:]
# metaset -s dg-schost-1 -d -h phys-schost-2
[Verify removal of the node:]
phys-schost-1% cldevicegroup list -v dg-schost-1
=== Cluster Device Groups ===

--- Device Group Status ---

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

▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos

Siga este procedimiento para eliminar un nodo de cluster de la lista de nodos primarios potenciales de un grupo de dispositivos de discos básicos.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster *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 5–24 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 cluster (`phys-schost-1`).

```
[Identify the device groups connected to the node being removed, and determine which are raw-disk
 device groups:]
phys-schost-1# cldevicegroup show -n phys-schost-2 -t rawdisk -v +
Device Group Name:          dsk/d4
Type:                       Disk
failback:                   false
Node List:                  phys-schost-2
preferenced:                false
localonly:                  false
autogen                     true
numsecondaries:             1
device names:               phys-schost-2

Device Group Name:          dsk/d2
Type:                       Disk
failback:                   true
Node List:                  pbrave2
preferenced:                false
localonly:                  false
autogen                     true
numsecondaries:             1
diskgroup name:             vxdg1
```

```
Device Group Name:          dsk/d1
Type:                      SVM
failback:                  false
Node List:                  pbrave1, pbrave2
preferenced:                true
localonly:                 false
autogen:                   true
numsecondaries:             1
diskset name:              ms1
(dsk/d4) Device group node list: phys-schost-2
(dsk/d2) Device group node list: phys-schost-1, phys-schost-2
(dsk/d1) Device group node list: phys-schost-1, phys-schost-2
[Disable the localonly flag for each local disk on the node:]
phys-schost-1# cldevicegroup set -p localonly=false dsk/d4
[Verify that the localonly flag is disabled:]
phys-schost-1# cldevicegroup show -n phys-schost-2 -t rawdisk +
(dsk/d4) Device group type:      Disk
(dsk/d8) Device group type:      Local_Disk
[Remove the node from all raw-disk device groups:]

phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d4
phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d2
phys-schost-1# cldevicegroup remove-node -n phys-schost-2 dsk/d1
```

▼ Cambio de propiedades de los grupos de dispositivos

El método para establecer la propiedad primaria de un grupo de dispositivos se basa en el establecimiento de un atributo de preferencia de propiedad denominado `preferenced`. Si el atributo no está definido, el propietario primario de un grupo de dispositivos que no está sujeto a ninguna otra relación de propiedad es el primer nodo que intente acceder a un disco de dicho grupo. Sin embargo, si este atributo está definido, debe especificar el orden de preferencia en el cual los nodos intentan establecer la propiedad.

Si inhabilita el atributo `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.

Este procedimiento utiliza el comando `clsetup` 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Inicie la utilidad `clsetup`.**
`# clsetup`
Aparece el menú principal.
- 3 **Para trabajar con grupos de dispositivos, escriba el número correspondiente a la opción de grupos de dispositivos y volúmenes.**
Se muestra el menú Grupos de dispositivos.
- 4 **Para modificar las propiedades clave de un grupo de dispositivos, escriba el número correspondiente a la opción para modificar 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 correspondiente a la opción para cambiar las preferencias o las propiedades de conmutación por recuperación.**
Siga las instrucciones para definir las opciones `preferenced` y `failback` de un grupo de dispositivos.
- 6 **Compruebe que se hayan modificado los atributos del grupo de dispositivos.**
Busque la información del grupo de dispositivos mostrada por el comando siguiente.
`# cldevicegroup show -v devicegroup`

Ejemplo 5–25 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 cluster. 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 cluster. 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**

2 Inicie la utilidad `clsetup`.

```
# clsetup
```

Aparece el menú principal.

3 Para trabajar con grupos de dispositivos, seleccione la opción del menú Grupos de dispositivos y volúmenes.

Se muestra el menú Grupos de dispositivos.

4 Para modificar las propiedades esenciales de un grupo de dispositivos, seleccione la opción Cambiar las propiedades esenciales de un grupo de dispositivos.

Se muestra el menú Cambiar las propiedades esenciales.

5 Para modificar el número de nodos secundarios, escriba el número correspondiente a la opción para cambiar la propiedad `numsecondaries`.

Siga las instrucciones y escriba el número de nodos secundarios que se van a configurar para el grupo de dispositivos. El comando `cldevicegroup` correspondiente se ejecuta a continuación, se imprime un registro y la utilidad vuelve al menú anterior.

6 Valide la configuración del grupo de dispositivos.

```
# cldevicegroup show dg-schost-1
=== Device Groups ===
```

Device Group Name:	dg-schost-1	
Type:	Local_Disk	<i>This might also be SDS.</i>
failback:	yes	
Node List:	phys-schost-1, phys-schost-2 phys-schost-3	
preferenced:	yes	
numsecondaries:	1	
diskgroup names:	dg-schost-1	

Nota – Si cambia algún tipo de información de configuración para un volumen o grupo de discos registrado con el cluster, debe volver a registrar el grupo de dispositivos mediante `clsetup`. Entre los cambios de configuración, se incluyen 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 125](#).

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 5–26 Modificación del número de nodos secundarios (Solaris Volume Manager)

En el ejemplo siguiente se muestra el comando `cldevicegroup` que genera `clsetup` al configurar el número de nodos secundarios para un grupo de dispositivos (`dg-schost-1`). En este ejemplo se supone que el grupo de discos y el volumen ya se habían creado.

```
# cldevicegroup set -p numsecondaries=1 dg-schost-1
# cldevicegroup show -v dg-schost-1

=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      SVM
failback:                   yes
Node List:                  phys-schost-1, phys-schost-2
preferenced:                yes
numsecondaries:             1
diskset names:              dg-schost-1
```

Ejemplo 5–27 Definición del número de nodos secundarios con el valor predeterminado

En el ejemplo siguiente se muestra el uso de un valor nulo de secuencia de comandos para configurar el número predeterminado de nodos secundarios. El grupo de dispositivos se configura para usar el valor predeterminado, aunque dicho valor sufra modificaciones.

```
# cldevicegroup set -p numsecondaries= dg-schost-1
# cldevicegroup show -v dg-schost-1

=== Device Groups ===

Device Group Name:          dg-schost-1
Type:                      SVM
failback:                   yes
Node List:                  phys-schost-1, phys-schost-2 phys-schost-3
preferenced:                yes
numsecondaries:             1
diskset names:              dg-schost-1
```

▼ Enumeración de la configuración de grupos de dispositivos

No es necesario ser un superusuario para enumerar un listado con la configuración. Sin embargo, se debe disponer de autorización `solaris.cluster.read`.

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

Este procedimiento proporciona las formas largas de los comandos de 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.**

GUI de Oracle Solaris Cluster Manager	Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.
<code>cldevicegroup show</code>	Use <code>cldevicegroup show</code> para enumerar la configuración de todos los grupos de dispositivos del cluster.
<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 cluster.

Use la opción `-v` con cualquiera de estos comandos para obtener información más detallada.

Ejemplo 5-28 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 5-29 Enumeración de la configuración de un determinado grupo de dispositivos

```
# cldevicegroup show dg-schost-1
=== Device Groups ===

Device Group Name:      dg-schost-1
Type:                   SVM
failback:               yes
```

Node List:	phys-schost-2, phys-schost-3
preferenced:	yes
numsecondaries:	1
diskset names:	dg-schost-1

▼ Conmutación al nodo primario de un grupo de dispositivos

Este procedimiento también es válido para iniciar (poner en línea) un grupo de dispositivos inactivo.

También puede poner en línea un grupo de dispositivos inactivo o cambiar el nodo principal de un grupo de dispositivos mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**

- 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.

`devicegroup` 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 5-30 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. Debe poner los grupos de dispositivos en estado de mantenimiento al finalizar los procedimientos de reparación que requieran consentimiento para todas las actividades de E/S hasta que termine la reparación. Asimismo, poner un grupo de dispositivos en estado de mantenimiento ayuda a impedir la pérdida de datos, al asegurar que el grupo de dispositivos no se pone en línea en un nodo mientras el conjunto o el grupo de discos se esté reparando en otro nodo.

Para obtener instrucciones sobre cómo restaurar un conjunto de discos dañado, consulte [“Restauración de un conjunto de discos dañado” en la página 277](#).

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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 5–31 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 149
- “Visualización del protocolo SCSI de un solo dispositivo de almacenamiento” en la página 150
- “Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento” en la página 151
- “Modificación del protocolo de protección en un solo dispositivo de almacenamiento” en la página 152

▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento

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

Este procedimiento proporciona las formas largas de los comandos de 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 5–32 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 cluster.

```
# cluster show -t global

=== Cluster ===

Cluster Name:                racerxx
installmode:                 disabled
heartbeat_timeout:          10000
heartbeat_quantum:           1000
private_netaddr:             172.16.0.0
private_netmask:             255.255.248.0
max_nodes:                   64
max_privatenets:             10
global_fencing:              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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 5-33 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:
default_fencing:

none
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 cluster. La configuración de aislamiento predeterminada de un solo dispositivo de almacenamiento anula la configuración global al establecer el aislamiento predeterminado del dispositivo en `pathcount`, `prefer3` o `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 152](#) 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 cluster desde hosts situados fuera de dicho cluster.

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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster 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 5–34 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 el protocolo de protección para todos los dispositivos de almacenamiento del cluster.

```
# 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster desde hosts situados fuera de dicho cluster.

- 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 cluster tengan acceso al almacenamiento. Use la opción `nofencing-noscrub` sólo con dispositivos de almacenamiento con problemas graves con las reservas de SCSI.

device 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 5–35 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 cluster

El sistema de archivos de cluster está disponible globalmente. Se puede leer y tener acceso a él desde cualquier nodo del cluster.

TABLA 5–5 Mapa de tareas: administrar sistemas de archivos de cluster

Tarea	Instrucciones
Agregar sistemas de archivos de cluster tras la instalación inicial de Oracle Solaris Cluster	“Adición de un sistema de archivos de cluster” en la página 154
Eliminar un sistema de archivos de cluster	“Eliminación de un sistema de archivos de cluster” en la página 158
Comprobar los puntos de montaje globales de un cluster para verificar la coherencia entre los nodos	“Comprobación de montajes globales en un cluster” en la página 160

▼ Adición de un sistema de archivos de cluster

Efectúe esta tarea en cada uno de los sistemas de archivos de cluster 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 cluster 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 cluster adicional, compruebe que se cumplan los requisitos siguientes:

- Se cuenta con privilegios de superusuario en un nodo del cluster.
- Software de administración de volúmenes instalado y configurado en el cluster.
- Un grupo de dispositivos (grupo de dispositivos de Solaris Volume Manager) o segmento de disco de bloques donde poder crear el sistema de archivos de cluster.

Si ha usado Oracle Solaris Cluster Manager para instalar servicios de datos, ya hay uno o más sistemas de archivos de cluster si los discos compartidos en los que crear los sistemas de archivos de cluster eran suficientes.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.

Lleve a cabo este procedimiento desde la zona global si hay zonas no globales configuradas en el cluster.

Consejo – Para crear sistemas de archivos con mayor rapidez, conviértase en superusuario en el nodo principal actual del dispositivo global para el que desea crear un sistema de archivos.

2 Cree un sistema de archivos.



Caution – Todos los datos de los discos se destruyen al crear un sistema de archivos. Compruebe que haya especificado el nombre del dispositivo de disco correcto. Si se especifica un nombre equivocado, podría borrar datos que no tuviera previsto eliminar.

- **Para un sistema de archivos UFS, use el comando `newfs(1M)`.**

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

En la tabla siguiente, se muestran ejemplos de nombres para el argumento *raw-disk-device*. 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 cluster para el sistema de archivos de dicho cluster.

Todos los nodos deben tener un punto de montaje, aunque no se acceda al sistema de archivos del cluster en un nodo concreto.

Consejo – Para facilitar la administración, cree el punto de montaje en el directorio `/global/device-group/`. Esta ubicación permite distinguir fácilmente los sistemas de archivos de cluster disponibles de forma global de los sistemas de archivos locales.

`phys-schost# mkdir -p /global/device-group/mountpoint/`
device-group Nombre del directorio correspondiente al nombre del grupo de dispositivos que contiene el dispositivo.
mountpoint Nombre del directorio en el que se monta el sistema de archivos de cluster.

4 En cada uno de los nodos del cluster, agregue una entrada en el archivo `/etc/vfstab` para el punto de montaje.

Consulte la página del comando `man vfstab(4)` para obtener más información.

Nota – Si se configuran zonas no globales en el cluster, asegúrese de montar sistemas de archivos de cluster en la zona global en una ruta del directorio raíz de la zona global.

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

e. Compruebe las dependencias de orden de inicio de los sistemas de archivos.

Por ejemplo, considere la siguiente situación hipotética: `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 En cualquier nodo del cluster, ejecute la utilidad de comprobación de la configuración.

```
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 cluster. Si no hay ningún error, el comando no devuelve nada.

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

6 Monte el sistema de archivos de cluster.

Para sistemas de archivos UFS y QFS, monte el sistema de archivos de cluster desde cualquier nodo del cluster.

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

7 Compruebe que el sistema de archivos de cluster esté montado en todos los nodos de dicho cluster.

Puede utilizar los comandos `df` o `mount` para enumerar los sistemas de archivos montados. Para obtener más información, consulte la página del comando `man df(1M)` o `mount(1M)`.

Se puede obtener acceso a los sistemas de archivos de cluster desde la zona global y desde la zona no global.

Ejemplo 5-36 Creación de un sistema de archivos de cluster UFS

En el ejemplo siguiente, se crea un sistema de archivos de cluster UFS en el volumen de Solaris Volume Manager `/dev/md/oracle/rdisk/d1`. Se agrega una entrada para el sistema de archivos de cluster en el archivo `vfstab` de cada nodo. A continuación, se ejecuta el comando `cluster check` desde un nodo. Una vez que el procesamiento de comprobación de la configuración se completa correctamente, el sistema de archivos de cluster 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
#
/dev/md/oracle/dsk/d1 /dev/md/oracle/rdisk/d1 /global/oracle/d1 ufs 2 yes global,logging
...
```

```
phys-schost# cluster check -k vfstab
phys-schost# mount /global/oracle/d1
phys-schost# mount
...
/global/oracle/d1 on /dev/md/oracle/dsk/d1 read/write/setuid/global/logging/largefiles
on Sun Oct 3 08:56:16 2005
```

▼ Eliminación de un sistema de archivos de cluster

Para *quitar* un sistema de archivos de cluster, 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 cluster se desmontan automáticamente como parte del cierre del sistema sucedido al ejecutar `cluster shutdown` para detener todo el cluster. Un sistema de archivos de cluster 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 cluster en ese disco da como resultado un error.

Compruebe que se cumplan los requisitos siguientes antes de desmontar los sistemas de archivos de cluster:

- Se cuenta con privilegios de superusuario en un nodo del cluster.
- 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 cluster.

1 Conviértase en superusuario en un nodo de cluster.

2 Determine los sistemas de archivos de cluster que están montados.

```
# mount -v
```

3 En cada nodo, enumere todos los procesos que utilicen el sistema de archivos de cluster 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.

mountpoint Especifica el nombre del sistema de archivos del cluster para el que desea detener procesos.

4 Detenga todos los procesos del sistema de archivos de cluster 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 cluster.

```
# fuser -c -k mountpoint
```

Se envía un SIGKILL a cada proceso que usa el sistema de archivos de cluster.

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
```

mountpoint Especifica el nombre del sistema de archivos del cluster que desea desmontar. Puede ser el nombre del directorio en el que está montado el sistema de archivos de cluster 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 cluster que se va a eliminar.

Realice este paso en cada nodo del cluster que tenga una entrada de este sistema de archivos de cluster 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 5-37 Eliminación de un sistema de archivos de cluster

En el ejemplo siguiente se quita un sistema de archivos de cluster UFS montado en el metadispositivo o el volumen de Solaris Volume Manager `/dev/md/oracle/rdisk/d1`.

```
# mount -v
...
/global/oracle/d1 on /dev/md/oracle/dsk/d1 read/write/setuid/global/logging/largefiles
# fuser -c /global/oracle/d1
/global/oracle/d1: 4006c
# fuser -c -k /global/oracle/d1
/global/oracle/d1: 4006c
# fuser -c /global/oracle/d1
/global/oracle/d1:
# umount /global/oracle/d1
```

```
(On each node, remove the highlighted entry:)
# vi /etc/vfstab
#device          device          mount   FS      fsck    mount   mount
#to mount        to fsck      point   type    pass    at boot options
#
/dev/md/oracle/dsk/d1 /dev/md/oracle/rdisk/d1 /global/oracle/d1 ufs 2 yes global,logging

[Save and exit.]
```

Para borrar los datos del sistema de archivos de cluster, elimine el dispositivo subyacente. Para obtener más información, consulte la documentación del administrador de volúmenes.

▼ Comprobación de montajes globales en un cluster

La utilidad `cluster(1CL)` comprueba la sintaxis de las entradas de sistemas de archivos de cluster 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 cluster, que puedan haber afectado a los dispositivos o a los componentes de administración de volúmenes).

- 1 **Conviértase en superusuario en un nodo de cluster.**
- 2 **Compruebe los montajes globales del cluster.**

```
# 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)`. Consulte también la página del comando `man syslogd(1M)` para ver los errores registrados que informa el daemon.

Nota – Las rutas de disco se agregan automáticamente a la lista de supervisión al incorporar dispositivos de E/S a un nodo mediante el comando `cldevice`. Asimismo, la supervisión de rutas de disco se detiene automáticamente al quitar los dispositivos de un nodo con los comandos de Oracle Solaris Cluster.

TABLA 5–6 Mapa de tareas: administrar la supervisión de rutas de disco

Tarea	Instrucciones
Supervisar una ruta de disco	“Supervisión de una ruta de disco” en la página 161
Anular la supervisión de una ruta de disco	“Anulación de la supervisión de una ruta de disco” en la página 163
Imprimir el estado de rutas de disco erróneas correspondientes a un nodo	“Impresión de rutas de disco erróneas” en la página 164
Supervisar rutas de disco desde un archivo	“Supervisión de rutas de disco desde un archivo” en la página 165
Habilitar o inhabilitar el rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	“Habilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 167 “Inhabilitación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas” en la página 168
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 164

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 cluster.



Precaución – DPM no se admite en los nodos que ejecutan versiones que se publicaron antes que el software Sun Cluster 3.1 10/03. No utilice comandos de DPM cuando hay una actualización gradual en curso. Después de actualizar todos los nodos, los nodos deben estar en línea para utilizar los comandos de DPM.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Supervise una ruta de disco.**
`# cldevice monitor -n node disk`
- 3 **Compruebe que se supervise la ruta de disco.**
`# cldevice status device`

Ejemplo 5–38 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 5–39 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 5-40 Relectura de la configuración de disco desde CCR

En el ejemplo siguiente se fuerza al daemon a que relea la configuración de disco desde CCR e imprima las rutas de disco supervisadas con su estado.

```
# cldevice monitor +
# cldevice status
Device Instance      Node      Status
-----
/dev/did/rdisk/d1    schost-1    Ok
/dev/did/rdisk/d2    schost-1    Ok
/dev/did/rdisk/d3    schost-1    Ok
                   schost-2    Ok
/dev/did/rdisk/d4    schost-1    Ok
                   schost-2    Ok
/dev/did/rdisk/d5    schost-1    Ok
                   schost-2    Ok
/dev/did/rdisk/d6    schost-1    Ok
                   schost-2    Ok
/dev/did/rdisk/d7    schost-2    Ok
/dev/did/rdisk/d8    schost-2    Ok
```

▼ Anulación de la supervisión de una ruta de disco

Siga este procedimiento para anular la supervisión de una ruta de disco.



Precaución – DPM no se admite en los nodos que ejecutan versiones que se publicaron antes que el software Sun Cluster 3.1 10/03. No utilice comandos de DPM cuando hay una actualización gradual en curso. Después de actualizar todos los nodos, los nodos deben estar en línea para utilizar los comandos de DPM.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 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 5-41 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/rdsk/d1` y se imprimen las rutas de disco de todo el cluster, con sus estados.

```
# cldevice unmonitor -n schost2 /dev/did/rdsk/d1
# cldevice status -n schost2 /dev/did/rdsk/d1
```

Device Instance	Node	Status
-----	----	-----
/dev/did/rdsk/d1	schost-2	Unmonitored

▼ Impresión de rutas de disco erróneas

Siga este procedimiento para imprimir las rutas de disco erróneas correspondientes a un cluster.



Precaución – DPM no se admite en los nodos que ejecutan versiones que se publicaron antes que el software Sun Cluster 3.1 10/03. No utilice comandos de DPM cuando hay una actualización gradual en curso. Después de actualizar todos los nodos, los nodos deben estar en línea para utilizar los comandos de DPM.

- 1 **Conviértase en superusuario en un nodo de cluster.**
- 2 **Imprima las rutas de disco erróneas de todo el cluster.**
`# cldevice status -s fail`

Ejemplo 5-42 Impresión de rutas de disco erróneas

En el ejemplo siguiente se imprimen las rutas de disco erróneas de todo el cluster.

```
# 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 cluster.
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 cluster mediante un archivo, primero se debe exportar la configuración actual. Esta operación de exportación crea un archivo XML que luego puede modificarse para establecer los elementos de la configuración que vaya a cambiar. Las instrucciones de este procedimiento describen el proceso completo.



Precaución – DPM no se admite en los nodos que ejecutan versiones que se publicaron antes que el software Sun Cluster 3.1 10/03. No utilice comandos de DPM cuando hay una actualización gradual en curso. Después de actualizar todos los nodos, los nodos deben estar en línea para utilizar los comandos de DPM.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 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 *configurationfile* Especifique el nombre de archivo del archivo XML modificado.
- 5 **Compruebe que la ruta de dispositivo ya se supervise.**

```
# cldevice status
```

Ejemplo 5-43 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 cluster 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 cluster y usar el archivo XML resultante para establecer la configuración del cluster, consulte las páginas del 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 supervisados como mínimo desde un nodo diferente del cluster. El daemon `scdpm` utiliza la interconexión privada para comprobar si se puede acceder a los discos desde un nodo diferente del cluster. Si la interconexión privada está desactivada, el daemon `scdpm` no puede obtener el estado de los discos desde otro nodo.

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 cluster o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.
- 2 Habilite el re arranque automático de un nodo para *todos* los nodos del cluster si fallan todas las rutas de disco compartido supervisadas de dicho nodo.

```
# clnode set -p reboot_on_path_failure=enabled +
```

▼ Inhabilitación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas

Si se inhabilita esta función y fallan todas las rutas de disco compartido supervisadas de un nodo, dicho nodo *no* re arranca automáticamente.

- 1 Conviértase en superusuario en uno de los nodos del cluster o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.
- 2 Inhabilite el re arranque automático de un nodo para *todos* los nodos del cluster 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 169
- “Administración de servidores de quórum de Oracle Solaris Cluster” en la página 193

Administración de dispositivos de quórum

Un dispositivo de quórum es un dispositivo de almacenamiento o servidor de quórum compartido por dos o más nodos y que aporta votos usados para establecer un quórum. Esta sección explica los procedimientos para administrar dispositivos de quórum.

Puede usar el comando `clquorum(1CL)` para realizar todos los procedimientos administrativos de los dispositivos de quórum. Además, es posible llevar a cabo algunos procedimientos mediante la utilidad interactiva `clsetup(1CL)` o la GUI de Oracle Solaris Cluster Manager. Siempre que es posible, los procedimientos de quórum de esta sección se describen con la utilidad `clsetup`. La ayuda en pantalla de Oracle Solaris Cluster Manager describe cómo realizar procedimientos de quórum mediante la GUI. Al trabajar con dispositivos de quórum, tenga en cuenta las directrices siguientes:

- Todos los comandos relacionados con el quórum deben ejecutarse en el nodo de votación del cluster 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 cluster. De darse esta incoherencia, vuelva a ejecutar el comando o ejecute el comando `clquorum reset` para restablecer la configuración de quórum.

- Para que el cluster 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 cluster 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 cluster. 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 cluster. 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 cluster de tres nodos.

TABLA 6-1 Lista de tareas: administrar el quórum

Tarea	Para obtener instrucciones
Agregar un dispositivo del quórum a un cluster mediante la utilidad <code>clsetup</code>	“Adición de un dispositivo de quórum” en la página 172
Eliminar un dispositivo de quórum de un cluster mediante <code>clsetup</code> (para generar <code>clquorum</code>)	“Eliminación de un dispositivo de quórum” en la página 181
Eliminar el último dispositivo de quórum de un cluster mediante <code>clsetup</code> (para generar <code>clquorum</code>)	“Eliminación del último dispositivo de quórum de un cluster” en la página 182
Reemplazar un dispositivo de quórum de un cluster mediante los procedimientos de agregar y quitar	“Sustitución de un dispositivo de quórum” en la página 184
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 185

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

Tarea	Para obtener instrucciones
Poner un dispositivo de quórum en estado de mantenimiento mediante <code>clsetup</code> (para generar <code>clquorum</code>)	“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 187
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 <code>clsetup</code> (para generar <code>clquorum</code>)	“Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 189
Enumerar los dispositivos de quórum y el número de votos mediante <code>clquorum</code>	“Enumeración de una lista con la configuración de quórum” en la página 190

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 cluster.

- 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 172
2. Inhabilitar el dispositivo de quórum que se va a eliminar	“Eliminación de un dispositivo de quórum” en la página 181

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

Tarea	Para obtener instrucciones
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 de quórum. Compruebe que todos los nodos del cluster 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 cluster, 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 cluster. 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)
 - NAS de Sun
 - 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 173](#)

- [“Adición de un dispositivo de quórum de servidor de quórum” en la página 177](#)

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 cluster. 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 los dispositivos de quórum, puede utilizar un disco que contenga los datos de usuario o sea miembro de un grupo de dispositivos. El protocolo que utiliza el subsistema de quórum con un disco compartido puede verse si mira el valor de `access-mode` para el disco compartido en la salida del comando `cluster show`.

Estos procedimientos también se pueden llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Consulte las páginas del 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 cluster 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 Inicie la utilidad `clsetup`.**
`# clsetup`
Aparece el menú principal de `clsetup`.
- 3 Escriba el número correspondiente a la opción de quórum.**
Aparece el menú Quórum.
- 4 Escriba el número correspondiente a la opción para agregar un dispositivo de quórum; a continuación, escriba *yes* cuando la utilidad `clsetup` le solicite que confirme el dispositivo de quórum que desea agregar.**
La utilidad `clsetup` pregunta qué tipo de dispositivo de quórum se desea agregar.
- 5 Escriba el número correspondiente a la opción de un dispositivo de quórum de disco compartido.**
La utilidad `clsetup` pregunta qué dispositivo global quiere utilizar.
- 6 Escriba el dispositivo global que va a usar.**
La utilidad `clsetup` solicita que confirme que el nuevo dispositivo de quórum debe agregarse al dispositivo global especificado.
- 7 Escriba *yes* para seguir agregando el nuevo dispositivo de quórum.**
Si se agrega correctamente el nuevo dispositivo de quórum, la utilidad `clsetup` muestra el correspondiente mensaje.
- 8 Compruebe que se haya agregado el dispositivo de quórum.**
`# clquorum list -v`

Ejemplo 6–1 Adición de un dispositivo de quórum de disco compartido

En el ejemplo siguiente se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum de disco compartido y el paso de comprobación.

Become superuser or assume a role that provides `solaris.cluster.modify` RBAC authorization on any cluster node.

[Start the `clsetup` utility:]

clsetup

[Select Quorum>Add a quorum device]

[Answer the questions when prompted.]

[You will need the following information.]

[Information:	Example:]
[Directly attached shared disk	shared_disk]
[Global device	d20]

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

clquorum add d20

Command completed successfully.

[Quit the `clsetup` Quorum Menu and Main Menu.]

[Verify that the quorum device is added:]

clquorum list -v

Quorum	Type
-----	----
d20	shared_disk
scphyshost-1	node
scphyshost-2	node

▼ Adición de un dispositivo de quórum NAS de Sun o Sun ZFS Storage Appliance

Compruebe que todos los nodos del cluster estén en línea antes de agregar un nuevo dispositivo de quórum.

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

Este procedimiento proporciona las formas largas de los comandos de 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 **Utilice la GUI de NAS de Sun para configurar un dispositivo iSCSI en el archivador NAS de Sun. 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.**

Si tiene un dispositivo NAS de Sun, utilice los comandos siguientes:

- a. Cree un volumen de archivos cuyo tamaño sea aproximadamente de 50 MB.
- b. Para cada nodo, cree una lista de acceso iSCSI.
 - i. Utilice el nombre del cluster como nombre de la lista de acceso iSCSI.

- ii. Agregue el nombre de nodo de iniciador de cada nodo del cluster a la lista de acceso. CHAP e IQN no son necesarios.

c. Configure el LUN iSCSI.

Puede utilizar el nombre del volumen de archivos de copia de seguridad como nombre del LUN. Agregue la lista de acceso de cada nodo al LUN.

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 cluster, 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 cluster 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 cluster.

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 cluster tiene reglas predeterminadas para decidir si se deben usar `scsi-2`, `scsi-3` o los protocolos de quórum del software. Consulte `clquorum(1CL)` para obtener más información.

Ejemplo 6–2 Adición de un dispositivo de quórum NAS de Sun o Sun ZFS Storage Appliance

En el siguiente ejemplo, se muestra el comando `clquorum` generado por `clsetup` al agregar un dispositivo de quórum NAS de Sun y un paso de verificación. 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.

```

Add an iSCSI device on the Sun NAS filer.
Use the Sun NAS GUI to create a file volume that is approximately 50mb in size.
File Volume Operations -> Create File Volume
For each node, create an iSCSI access list.
iSCSI Configuration -> Configure Access List
Add the initiator node name of each cluster node to the access list.
*** Need GUI or command syntax for this step. ***
Configure the iSCSI LUN
iSCSI Configuration -> Configure iSCSI LUN
On each of the cluster nodes, discover the iSCSI LUN and set the iSCSI access list to static configuration.
iscsiadm modify discovery -s enable
iscsiadm list discovery
Discovery:
    Static: enabled
    Send Targets: enabled
    iSNS: disabled
iscsiadm add static-config
iqn.1986-03.com.sun0-1:000e0c66efe8.4604DE16.thinqorum,10.11.160.20
devsadm -i iscsi
From one cluster node, configure the DID devices for the iSCSI LUN.
/usr/cluster/bin/sclddevice populate
/usr/cluster/bin/sclddevice populate
Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on any cluster node.

[Add the NAS device as a quorum device
using the DID device:]
clquorum add d20

Command completed successfully.

```

▼ Adición de un dispositivo de quórum de servidor de quórum

Antes de empezar

Antes de poder agregar un servidor de quórum de Oracle Solaris Cluster como dispositivo de quórum, el software del servidor de quórum de Oracle Solaris Cluster debe estar instalado en el equipo host, y el servidor de quórum debe haberse iniciado y estar en ejecución. Si desea obtener información sobre cómo instalar el servidor de quórum, consulte [“Instalación y configuración del software Servidor de quórum” de Guía de instalación del software de Oracle Solaris Cluster](#).

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**

2 Compruebe que todos los nodos de Oracle Solaris Cluster estén en línea y puedan comunicarse con el servidor de quórum de Oracle Solaris Cluster.

a. Compruebe que los conmutadores de red conectados directamente con los nodos del cluster 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 cluster y el servidor de quórum sea inmediata. Si el conmutador ralentizada dicha comunicación se ralentizase de forma significativa, el cluster 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 cluster.

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 cluster 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
```

ipaddress Dirección IP del equipo donde se ejecuta el servidor de quórum.

qshost1 Nombre de host del equipo donde se ejecuta el servidor de quórum.

d. Si usa un servicio de nombres, agregue la asignación entre nombre y dirección de host del servidor de quórum a la base de datos del servicio de nombres.

3 Inicie la utilidad `clsetup`.

```
# clsetup
```

Aparece el menú principal de `clsetup`.

4 Escriba el número correspondiente a la opción de quórum.

Aparece el menú Quórum.

5 Escriba el número correspondiente a la opción para agregar un dispositivo de quórum. A continuación, escriba yes para confirmar que va a agregar un dispositivo de quórum.

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

6 Escriba el número correspondiente a la opción de un dispositivo de quórum de servidor de quórum. A continuación, escriba yes para confirmar que va a agregar un dispositivo de quórum de servidor de quórum.

La utilidad `clsetup` le solicita que indique el nombre del nuevo dispositivo de quórum.

7 Escriba el nombre del dispositivo de quórum que va a agregar.

Elija el nombre que quiera para el dispositivo de quórum. Este nombre usa sólo para procesar futuros comandos administrativos.

La utilidad `clsetup` solicita que 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 cluster 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 cluster.

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 6-3 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 181
- “Eliminación del último dispositivo de quórum de un cluster” en la página 182
- “Sustitución de un dispositivo de quórum” en la página 184

▼ Eliminación de un dispositivo de quórum

Estos procedimientos también se pueden llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Cuando se elimina un dispositivo de quórum, éste ya no participa en el voto para establecer el quórum. Todos los clusters de dos nodos necesitan como mínimo un dispositivo de quórum configurado. Si éste es el último dispositivo de quórum de un cluster, `clquorum(1CL)` no podrá eliminar 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 cluster, consulte el procedimiento “[Eliminación del último dispositivo de quórum de un cluster](#)” en la página 182.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Determine el dispositivo de quórum que se va a eliminar.**
`# clquorum list -v`
- 3 **Ejecute la utilidad `clsetup(1CL)`.**
`# clsetup`

Aparece el menú principal.

- 4 **Escriba el número correspondiente a la opción de quórum.**
- 5 **Escriba el número correspondiente a la opción para eliminar un dispositivo de quórum.**
Responda a las preguntas que aparecen durante el proceso de eliminación.
- 6 **Salga de clsetup.**
- 7 **Compruebe que se haya eliminado el dispositivo de quórum.**
`# clquorum list -v`

Ejemplo 6–4 Eliminación de un dispositivo de quórum

En este ejemplo se muestra cómo quitar un dispositivo de quórum de un cluster 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 cluster 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 clusters del servidor de quórum” en la página 197.](#)

▼ **Eliminación del último dispositivo de quórum de un cluster**

Este procedimiento elimina el último dispositivo del quórum de un cluster 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 clusters, siga los pasos descritos en [“Eliminación de un dispositivo de quórum” en la página 181](#).

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. 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 172](#) y la página del comando `man clquorum(1CL)`.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Quite el dispositivo de quórum con el comando `clquorum`. Si se produjo un error en el dispositivo de quórum, use la opción `-F` (Force) para eliminar el dispositivo con errores.**

```
# clquorum remove -F qd1
```

Nota – También puede situar el nodo que se va a eliminar en estado de mantenimiento y, a continuación, quitar el dispositivo de quórum con el comando `clquorum remove quorum`. Mientras el cluster está en modo de instalación, las opciones del menú de administración del cluster `clsetup(1CL)` no están disponibles. Consulte [“Cómo poner un nodo en estado de mantenimiento” en la página 251](#) para obtener más información.

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

```
# clquorum list -v
```

Ejemplo 6-5 Eliminación del último dispositivo de quórum

En este ejemplo se muestra cómo poner el cluster en modo de mantenimiento y quitar el último dispositivo de quórum de la configuración del cluster.

```
[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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 172](#) para agregar un nuevo dispositivo de quórum al cluster.

2 Quite el dispositivo que va a sustituir como dispositivo de quórum.

Consulte [“Eliminación de un dispositivo de quórum” en la página 181](#) 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 contenedor de discos en [Oracle Solaris Cluster 3.3 3/13 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 185](#)
- [“Colocación de un dispositivo de quórum en estado de mantenimiento” en la página 187](#)

- “Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento” en la página 189
- “Enumeración de una lista con la configuración de quórum” en la página 190
- “Reparación de un dispositivo de quórum” en la página 191
- “Cambio del tiempo de espera predeterminado del quórum” en la página 192

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

Puede emplear la utilidad `clsetup(1CL)` para agregar un nodo o para eliminar un nodo de la lista de nodos de un dispositivo de 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 cluster. Cuando se agrega un dispositivo de quórum, `clquorum(1CL)` configura automáticamente las rutas del nodo al disco de todos los nodos conectados al disco.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 Determine el nombre del dispositivo de quórum que va a modificar.**
`# clquorum list -v`
- 3 Inicie la utilidad `clsetup`.**
`# clsetup`
Aparece el menú principal.
- 4 Escriba el número correspondiente a la opción de quórum.**
Aparece el menú Quórum.
- 5 Escriba el número correspondiente a la opción para eliminar un dispositivo de quórum.**
Siga las instrucciones. Se preguntará el nombre del disco que se va a eliminar.
- 6 Agregue o elimine las conexiones del nodo con el dispositivo de quórum.**

7 Escriba el número correspondiente a la opción para agregar un dispositivo de quórum.

Siga las instrucciones. Se solicitará el nombre del disco que se va a usar como dispositivo de quórum.

8 Compruebe que se haya agregado el dispositivo de quórum.

```
# clquorum list -v
```

Ejemplo 6–6 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(1CL)` para poner un dispositivo de quórum en estado de mantenimiento. Actualmente, la utilidad `clsetup(1CL)` no tiene esta capacidad. Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Ponga el dispositivo de quórum en estado de mantenimiento cuando vaya a apartar del servicio el dispositivo de quórum durante un período de tiempo prolongado. De esta forma, el número de votos de quórum del dispositivo de quórum se establece en cero y no aporta nada al número de quórum mientras se efectúan las tareas de mantenimiento en el dispositivo. La información de configuración del dispositivo de quórum se conserva durante el estado de mantenimiento.

Nota – Todos los clusters de dos nodos deben tener configurado al menos un dispositivo de quórum. Si éste es el último dispositivo de quórum de un cluster de dos nodos, `clquorum` el dispositivo no puede ponerse en estado de mantenimiento.

Para poner un nodo de un cluster en estado de mantenimiento, consulte [“Cómo poner un nodo en estado de mantenimiento” en la página 251.](#)

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 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 6-7 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 189.](#)

Para poner un nodo en estado de mantenimiento, consulte [“Cómo poner un nodo en estado de mantenimiento” en la página 251.](#)

▼ 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 cluster.

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 cluster y sus dispositivos de quórum asociados del estado de mantenimiento, consulte [“Cómo sacar un nodo del estado de mantenimiento” en la página 252.](#)
- 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 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, re arranque el nodo.**
- 4 **Compruebe el número de votos de quórum.**

```
# clquorum show +
```

Ejemplo 6-8 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

Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Para enumerar la configuración de quórum no hace falta ser superusuario. Se puede asumir cualquier rol que proporcione la autorización `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 clusters, 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 6-9 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 172](#) para agregar un nuevo dispositivo de quórum.

Consulte [“Eliminación de un dispositivo de quórum” en la página 181](#) 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 de hardware del contenedor de discos en [Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#).

3 Agregue el disco sustituido como nuevo dispositivo de quórum.

Consulte [“Adición de un dispositivo de quórum” en la página 172](#) 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 181](#) 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 cluster. 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 cluster 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 193
- “Inicio de un servidor de quórum” en la página 193
- “Detención de un servidor de quórum” en la página 194
- “Visualización de información sobre el servidor de quórum” en la página 195
- “Limpieza de la información caducada sobre clusters del servidor de quórum” en la página 197

Si desea obtener información sobre cómo instalar y configurar servidores de quórum de Oracle Solaris Cluster, consulte [“Instalación y configuración del software Servidor de quórum” de Guía de instalación del software 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 predeterminado, 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 Servidor de quórum” 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 192](#).

▼ 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 6–10 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 6–11 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.

quorumserver 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 6–12 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 6–13 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 clusters donde el servidor de quórum se configuró como dispositivo de quórum, este comando muestra el nombre del cluster correspondiente, el ID de cluster, 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
```

quorumserver 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 6–14 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 clusters

que tienen configurado el servidor de quórum como dispositivo de quórum. Esta información incluye el nombre del cluster 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 cluster `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 6-15 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 6-16 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 clusters 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 181](#). 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 cluster 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 clusters del servidor de quórum perderá su validez en los casos siguientes:

- Cuando un cluster se anula sin que antes se haya eliminado dispositivo de quórum del cluster mediante el comando `clquorum remove`.
- Cuando un dispositivo de quórum del tipo `quorum_server` se elimina de un cluster 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 cluster, utilizar este procedimiento para limpiar un servidor de quórum válido puede afectar negativamente al quórum del cluster.



Limpieza de la información de configuración del servidor de quórum

Antes de empezar

Quite el dispositivo de quórum del cluster de servidor de quórum como se describe en [“Eliminación de un dispositivo de quórum” en la página 181](#).



Precaución – Si el cluster sigue utilizando este servidor de quórum, efectuar este procedimiento afectará negativamente al quórum del cluster.

- 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 clustername` El nombre del cluster que el servidor de quórum utilizaba anteriormente como dispositivo de quórum.

Puede obtener el nombre del cluster si ejecuta `cluster show` en un nodo del cluster.

`-I clusterID` ID del cluster.

	El ID del cluster es un número hexadecimal de ocho dígitos. Puede obtener el ID del cluster si ejecuta <code>cluster show</code> en un nodo del cluster.
<i>quorumserver</i>	Un identificador de uno o más servidores de quórum.
	El servidor de quórum se puede identificar mediante un número de puerto o un nombre de instancia. Los nodos del cluster usan el número del puerto para comunicarse con el servidor de quórum. El nombre de instancia aparece especificado en el archivo de configuración del servidor de quórum, <code>/etc/scqsd/scqsd.conf</code> .
-y	Fuerce el comando <code>clquorumserver clear</code> para limpiar la información del cluster 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 clusters 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 6–17 Limpieza de la información obsoleta de clusters de la configuración del servidor de quórum

En este ejemplo se elimina información correspondiente al cluster 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 clusters 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 clusters y las redes públicas implica procedimientos de software y de hardware. Durante la instalación y configuración inicial del cluster se suelen configurar las interconexiones de los clusters y las redes públicas, incluidos los grupos de varias rutas IP. El método de ruta múltiple se instala automáticamente con el sistema operativo Oracle Solaris 10 y se debe habilitar para utilizarlo. Si después debe modificarse la configuración de las interconexiones entre redes y cluster, puede usar los procedimientos de software descritos en este capítulo. Si desea obtener información sobre cómo configurar grupos de varias rutas IP en un cluster, consulte la sección [“Administración de redes públicas” en la página 214](#).

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

- [“Administración de interconexiones de clusters” en la página 199](#)
- [“Administración de redes públicas” en la página 214](#)

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 cluster.

Administración de interconexiones de clusters

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

Casi siempre se puede emplear la utilidad `clsetup` para administrar el transporte del cluster en la interconexión de cluster. Para obtener más información, consulte la página del comando `man`

`clsetup(1CL)`. Todos los comandos relacionados con las interconexiones de los clusters deben ejecutarse en el nodo de votación del cluster global.

Si desea información sobre los procedimientos de instalación del software de cluster, 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 clusters, consulte el *Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual*.

Nota – Durante los procedimientos de interconexión de cluster, 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 cluster

Tarea	Instrucciones
Administrar el transporte del cluster con <code>clsetup(1CL)</code>	“Obtención de acceso a las utilidades de configuración del cluster” en la página 25
Comprobar el estado de la interconexión de los clusters con <code>clinterconnect status</code>	“Comprobación del estado de la interconexión de cluster” en la página 201
Agregar un cable de transporte de cluster, un adaptador de transporte o un conmutador mediante <code>clsetup</code>	“Adición de dispositivos de cable de transporte de cluster, adaptadores o conmutadores de transporte” en la página 202
Quitar un cable de transporte de cluster, un adaptador de transporte o un conmutador de transporte mediante <code>clsetup</code>	“Eliminación de cable de transporte de cluster, adaptadores de transporte y conmutadores de transporte” en la página 205
Habilitar un cable de transporte de cluster mediante <code>clsetup</code>	“Habilitación de un cable de transporte de cluster” en la página 208
Inhabilitar un cable de transporte de cluster mediante <code>clsetup</code>	“Inhabilitación de un cable de transporte de cluster” en la página 209
Determinar el número de instancia de un adaptador de transporte	“Determinación del número de instancia de un adaptador de transporte” en la página 211
Cambiar la dirección IP o el intervalo de direcciones de un cluster	“Modificación de la dirección de red privada o del intervalo de direcciones de un cluster” en la página 212

Reconfiguración dinámica con interconexiones de clusters

Al completar operaciones de reconfiguración dinámica en las interconexiones de clusters 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 cluster por completo un adaptador activo, con el fin de desarrollar DR en una interconexión de cluster activa. Use el menú `clsetup` o los comandos correspondientes.



Precaución – El software Oracle Solaris Cluster necesita que cada nodo del cluster disponga al menos de una ruta que funcione y se comunique con cualquier otro nodo del cluster. No inhabilite una interfaz de interconexión privada que proporcione la última ruta a cualquier nodo del cluster.

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 216
2. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	

▼ Comprobación del estado de la interconexión de cluster

También puede lograr este procedimiento con la interfaz gráfica de usuario de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.
% 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 cluster

En el ejemplo siguiente se muestra el estado de una interconexión de cluster en funcionamiento.

```
% clinterconnect status
-- Cluster Transport Paths --
      Endpoint                Endpoint                Status
-----
Transport path: phys-schost-1:qfe1 phys-schost-2:qfe1 Path online
Transport path: phys-schost-1:qfe0 phys-schost-2:qfe0 Path online
Transport path: phys-schost-1:qfe1 phys-schost-3:qfe1 Path online
Transport path: phys-schost-1:qfe0 phys-schost-3:qfe0 Path online
Transport path: phys-schost-2:qfe1 phys-schost-3:qfe1 Path online
Transport path: phys-schost-2:qfe0 phys-schost-3:qfe0 Path online
```

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

Para obtener información sobre los requisitos para el transporte privado del cluster, consulte [“Interconnect Requirements and Restrictions” de Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#).

Este procedimiento también se puede llevar a cabo mediante la GUI de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

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

Este procedimiento proporciona las formas largas de los comandos de 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 clusters.**
Para conocer el procedimiento de instalación de un cable de transporte de cluster, consulte el [Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#).
- 2 Conviértase en superusuario en un nodo de cluster.**
- 3 Inicie la utilidad `clsetup`.**
`# clsetup`
Aparece el menú principal.
- 4 Escriba el número correspondiente a la opción de mostrar el menú de interconexión del cluster.**
- 5 Escriba el número correspondiente a la opción de agregar un cable de transporte.**
Siga las instrucciones y escriba la información que se solicite.
- 6 Escriba el número correspondiente a la opción de agregar el adaptador de transporte a un nodo.**
Siga las instrucciones y escriba la información que se solicite.

Si tiene pensado usar alguno de los adaptadores siguientes para la interconexión de cluster, agregue la entrada pertinente al archivo `/etc/system` en cada nodo del cluster. La entrada surtirá efecto la próxima vez que se arranque el sistema.

Adaptador	Entrada
ce	establecer <code>ce:ce_taskq_disable=1</code>
ipge	establecer <code>ipge:ipge_taskq_disable=1</code>
ixge	establecer <code>ixge:ixge_taskq_disable=1</code>

- 7 Escriba el número correspondiente a la opción de agregar el conmutador de transporte.**
Siga las instrucciones y escriba la información que se solicite.

8 Compruebe que se haya agregado el cable de transporte de cluster, 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 cluster, un adaptador de transporte o un conmutador de transporte

En el ejemplo siguiente se muestra cómo agregar un cable de transporte, un adaptador o un conmutador de transporte a un nodo mediante la utilidad `clsetup`.

```
[Ensure that the physical cable is installed.]
[Start the clsetup utility:]
# clsetup
[Select Cluster interconnect]

[Select either Add a transport cable,
Add a transport adapter to a node,
or Add a transport switch.]
[Answer the questions when prompted.]
[You Will Need: ]
[Information:      Example:]
node names          phys-schost-1
adapter names       qfe2
switch names        hub2
transport type       dlpi
[Verify that the clinterconnect
command completed successfully:]Command completed successfully.
Quit the clsetup Cluster Interconnect Menu and Main Menu.
[Verify that the cable, adapter, and switch are added:]
# clinterconnect show phys-schost-1:qfe2,hub2
===Transport Cables===
Transport Cable:          phys-schost-1:qfe2@0,hub2
Endpoint1:                phys-schost-2:qfe0@0
Endpoint2:                ethernet-1@2 ????. Should this be hub2?
State:                    Enabled

# clinterconnect show phys-schost-1:qfe2
=== Transport Adepters for qfe2
Transport Adapter:          qfe2
Adapter State:              Enabled
Adapter Transport Type:     dlpi
Adapter Property (device_name): ce
Adapter Property (device_instance): 0
Adapter Property (lazy_free): 1
Adapter Property (dlpi_heartbeat_timeout): 10000
Adpater Property (dlpi_heartbeat_quantum): 1000
Adapter Property (nw_bandwidth): 80
Adapter Property (bandwidth): 70
Adapter Property (ip_address): 172.16.0.129
Adapter Property (netmask): 255.255.255.128
Adapter Port Names:        0
Adapter Port 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 cluster, consulte [“Comprobación del estado de la interconexión de cluster” en la página 201.](#)

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

También puede lograr este procedimiento con la interfaz gráfica de usuario de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Complete el procedimiento descrito a continuación para quitar cables, adaptadores y conmutadores de transporte de la configuración de un nodo. Cuando se inhabilita un cable, sus dos extremos continúan configurados. Un adaptador no se puede quitar si sigue en uso como extremo de un cable de transporte.



Precaución – Cada nodo del cluster debe contar al menos una ruta de transporte operativa que lo comuniquen con todos los demás nodos del cluster. No debe haber ningún par de nodos que estén aislados entre sí. Antes de inhabilitar un cable, compruebe siempre el estado de interconexión de cluster 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 cluster.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Compruebe el estado de la ruta de transporte restante del cluster.**

```
# clinterconnect status
```



Precaución – Si recibe un mensaje de error del tipo "ruta errónea" al intentar eliminar un nodo de un cluster 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 cluster; como consecuencia, podría darse una reconfiguración del cluster.

3 Inicie la utilidad `c1setup`.

`c1setup`

Aparece el menú principal.

4 Escriba el número correspondiente a la opción de acceder al menú de interconexión de cluster.

5 Escriba el número correspondiente a la opción de desactivar el cable de transporte.

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

6 Escriba el número correspondiente a la opción de quitar el cable de transporte.

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

Nota – Si va a quitar un cable, desconecte el cable entre el puerto y el dispositivo de destino.

7 Escriba el número correspondiente a la opción de quitar de un nodo el adaptador de transporte.

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

Nota – Si va a quitar un adaptador físico de un nodo, consulte el [Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#) para conocer los procedimientos de las tareas de mantenimiento de hardware.

8 Escriba el número correspondiente a la opción de quitar un conmutador de transporte.

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

Nota – No es posible eliminar un conmutador si alguno de los puertos se sigue usando como extremo de algún cable de transporte.

9 Compruebe que se haya quitado el cable, adaptador o conmutador.

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

El cable o adaptador de transporte eliminado del nodo correspondiente no debe aparecer en la salida de este comando.

Ejemplo 7-3 Eliminación de un cable de transporte, un adaptador de transporte o un conmutador de transporte

En el ejemplo siguiente se muestra cómo quitar un cable, un adaptador o un conmutador de transporte mediante el comando `clsetup`.

```
[Become superuser on any node in the cluster.]
[Start the utility:]
# clsetup
[Select Cluster interconnect.]
[Select either Remove a transport cable,
Remove a transport adapter to a node,
or Remove a transport switch.]
[Answer the questions when prompted.]
  You Will Need:
    Information      Example:
    node names       phys-schost-1
    adapter names     qfe1
    switch names      hub1
[Verify that the clinterconnect
command was completed successfully:]
Command completed successfully.
[Quit the clsetup utility Cluster Interconnect Menu and Main Menu.]
[Verify that the cable, adapter, or switch is removed:]
# clinterconnect show phys-schost-1:qfe2,hub2
===Transport Cables===
Transport Cable:                phys-schost-2:qfe2@0,hub2
  Cable Endpoint1:                phys-schost-2:qfe0@0
  Cable Endpoint2:                ethernet-1@2 ??? Should this be hub2???
  Cable State:                    Enabled

# clinterconnect show phys-schost-1:qfe2
=== Transport Adepters for qfe2
Transport Adapter:                qfe2
  Adapter State:                  Enabled
  Adapter Transport Type:         dlpi
  Adapter Property (device_name): ce
  Adapter Property (device_instance): 0
  Adapter Property (lazy_free):   1
  Adapter Property (dlpi_heartbeat_timeout): 10000
  Adapter Property (dlpi_heartbeat_quantum): 1000
  Adapter Property (nw_bandwidth): 80
  Adapter Property (bandwidth):   70
  Adapter Property (ip_address):  172.16.0.129
  Adapter Property (netmask):     255.255.255.128
  Adapter Port Names:             0
  Adapter Port State (0):         Enabled
```

```
# clinterconnect show phys-schost-1:hub2
=== Transport Switches ===
Transport Switch:                hub2
Switch State:                    Enabled
Switch Type:                     switch
Switch Port Names:              1 2
Switch Port State(1):           Enabled
Switch Port State(2):           Enabled
```

▼ **Habilitación de un cable de transporte de cluster**

También puede lograr este procedimiento con la interfaz gráfica de usuario de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Esta opción se utiliza para activar un cable de transporte de cluster ya existente.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.

2 Inicie la utilidad clsetup.

```
# clsetup
```

Aparece el menú principal.

3 Escriba el número correspondiente a la opción de acceder al menú de interconexión de cluster y pulse la tecla de retorno.

4 Escriba el número correspondiente a la opción de activar el cable de transporte y pulse la tecla de retorno.

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 cluster

En este ejemplo, se muestra cómo activar un cable de transporte de cluster en el adaptador `qfe-1`, ubicado en el nodo `phys - schost - 2`.

```
[Become superuser on any node.]
[Start the clsetup utility:]
# clsetup
[Select Cluster interconnect>Enable a transport cable.]

[Answer the questions when prompted.]
[You will need the following information.]
  You Will Need:
Information:                                Example:
  node names                                phys-schost-2
  adapter names                             qfe1
  switch names                              hub1
[Verify that the scinterconnect
  command was completed successfully:]

clinterconnect enable phys-schost-2:qfe1

Command completed successfully.
[Quit the clsetup Cluster Interconnect Menu and Main Menu.]
[Verify that the cable is enabled:]
# clinterconnect show phys-schost-1:qfe2,hub2
Transport cable:  phys-schost-2:qfe1@0 ethernet-1@2    Enabled
Transport cable:  phys-schost-3:qfe0@1 ethernet-1@3    Enabled
Transport cable:  phys-schost-1:qfe0@0 ethernet-1@1    Enabled
```

▼ Inhabilitación de un cable de transporte de cluster

También puede lograr este procedimiento con la interfaz gráfica de usuario de Oracle Solaris Cluster Manager. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Es posible que necesite desactivar un cable de transporte de cluster para cerrar temporalmente una ruta de interconexión de cluster. Un cierre temporal es útil al buscar la solución a un problema de la interconexión de cluster o al sustituir hardware de interconexión de cluster.

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 cluster debe contar al menos una ruta de transporte operativa que lo comunique con todos los demás nodos del cluster. No debe haber ningún par de nodos que estén aislados entre sí. Antes de inhabilitar un cable, compruebe siempre el estado de interconexión de cluster 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 cluster.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**
- 2 **Compruebe el estado de la interconexión de cluster 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 cluster 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 cluster; como consecuencia, podría darse una reconfiguración del cluster.

- 3 **Inicie la utilidad `clsetup`.**

```
# clsetup
```

Aparece el menú principal.

- 4 **Escriba el número correspondiente a la opción de acceder al menú de interconexión de cluster y pulse la tecla de retorno.**
- 5 **Escriba el número correspondiente a la opción de desactivar el cable de transporte y pulse la tecla de retorno.**

Siga las instrucciones y escriba la información que se solicite. Se inhabilitarán todos los componentes de la interconexión de este cluster. 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 cluster

En este ejemplo, se muestra cómo desactivar un cable de transporte de cluster en el adaptador `qfe-1`, ubicado en el nodo `phys-schost-2`.

```
[Become superuser on any node.]
[Start the clsetup utility:]
# clsetup
[Select Cluster interconnect>Disable a transport cable.]

[Answer the questions when prompted.]
[You will need the following information.]
[ You Will Need:]
```

```

Information:
node names
adapter names
switch names

Example:
phys-schost-2
qfe1
hub1

[Verify that the clinterconnect
command was completed successfully:]
Command completed successfully.
[Quit the clsetup Cluster Interconnect Menu and Main Menu.]
[Verify that the cable is disabled:]
# clinterconnect show -p phys-schost-1:qfe2,hub2
Transport cable:  phys-schost-2:qfe1@0  ethernet-1@2  Disabled
Transport cable:  phys-schost-3:qfe0@1  ethernet-1@3  Enabled
Transport cable:  phys-schost-1:qfe0@0  ethernet-1@1  Enabled

```

▼ Determinación del número de instancia de un adaptador de transporte

Debe determinarse el número de instancia de un adaptador de transporte para asegurarse de que se agregue y quite el adaptador correcto mediante el comando `clsetup`. El nombre del adaptador es una combinación del tipo de adaptador y del número de instancia de dicho adaptador.

1 Según el número de ranura, busque el nombre del adaptador.

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```

# prtdiag
...
===== IO Cards =====
                        Bus Max
                        Type
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 cluster

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 superusuario al shell remoto (`rsh(1M)`) o el shell seguro (`ssh(1)`) esté activado en todos los nodos del cluster. Para obtener más información, consulte las páginas del comando `man rsh(1M)` o `ssh(1)`.

- 1 **Rearranque todos los nodos de cluster en un modo que no sea de cluster. Para ello, aplique los subpasos siguientes en cada nodo del cluster:**

- a. **Conviértase en superusuario o asuma una función que proporcione la autorización RBAC `solaris.cluster.admin` en el cluster que se va a iniciar en un modo que no es de cluster.**

- 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 cambia todos los grupos de recursos desde los nodos con voto o sin voto del nodo especificado hasta el siguiente nodo por orden de preferencia, sea con voto o sin voto.

```
# 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 cluster, la utilidad `clsetup` muestra el menú principal para operaciones de un modo que no sea de cluster.

- 3 **Escriba el número que corresponde a la opción para cambiar la asignación de direcciones de red y rangos para el transporte del cluster y pulse la tecla `Intro`.**

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

- 4 **Para cambiar la dirección IP de red privada o el intervalo de direcciones de red IP, escriba `yes` y pulse la tecla `Intro`.**

La utilidad `clsetup` muestra la dirección IP de red privada, `172.16.0.0`, y pregunta si desea aceptarla de forma predeterminada.

5 Cambie o acepte la dirección IP de red privada.

- **Para aceptar la dirección IP de red privada predeterminada y cambiar el rango de direcciones IP, escriba *yes* y pulse la tecla *Intro*.**

La utilidad `clsetup` le preguntará si está de acuerdo con aceptar la máscara de red predeterminada. Vaya al paso siguiente para introducir su respuesta.

- **Para cambiar la dirección IP de red privada, realice estos pasos secundarios.**

- 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 clusters de zona y 10 redes privadas en el cluster.

- **Para aceptar el rango de direcciones IP predeterminado, escriba *yes* y pulse la tecla *Intro*.**

Continúe directamente con el paso siguiente.

- **Para cambiar el rango de direcciones IP, realice los siguientes pasos secundarios.**

- 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 clusters de zona que tiene previsto configurar en el cluster.

- b. **Especifique el número de nodos, redes privadas y clusters de zona que tiene previsto configurar en el cluster.**

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 clusters de zona que haya especificado.
- La segunda máscara de red admite el doble de nodos, redes privadas y clusters 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 clusters de zona.
- 7 Escriba **yes** como respuesta a la pregunta de la utilidad **c1setup** sobre si desea continuar con la actualización.
 - 8 Cuando haya finalizado, salga de la utilidad **c1setup**.
 - 9 Rearranque todos los nodos del cluster de nuevo en modo de cluster; 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. Aparece el menú de GRUB, que es similar al siguiente:


```
GNU GRUB version 0.95 (631K lower / 2095488K upper memory)
+-----+
| Solaris 10 /sol_10_x86                               |
| Solaris failsafe                                       |
|                                                       |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press enter to boot the selected OS, 'e' to edit the
commands before booting, or 'c' for a command-line.
```
 - 10 Compruebe que el nodo haya arrancado sin errores y esté en línea.
cluster status -t node

Administración de redes públicas

El software 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 clusters y sin clusters. El método de ruta múltiple se instala automáticamente al instalar el sistema operativo Oracle Solaris 10 y se debe habilitar para utilizarlo. La administración de varias rutas se trata en la documentación correspondiente del sistema operativo Oracle Solaris. No obstante, consulte las indicaciones siguientes antes de administrar IPMP en un entorno de Oracle Solaris Cluster.

Administración de grupos de varias rutas de red IP en un cluster

Antes de realizar procedimientos relacionados con IPMP en un cluster, 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 cluster. 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 rutas múltiples 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 cluster. Sin una conexión de red pública no es posible tener acceso al cluster.

- Para ver el estado de los grupos de varias rutas de red IP de un cluster, use el comando `clinterconnect status`.

Si desea obtener más información sobre varias rutas de red IP, consulte la documentación pertinente en el conjunto de documentación de administración de sistemas del sistema operativo Oracle Solaris.

TABLA 7-3 Mapa de tareas: administrar la red pública

Versión del sistema operativo Oracle Solaris	Instrucciones
Sistema operativo Oracle Solaris 10	“Temas sobre rutas múltiples de redes IP” en <i>Administración de Oracle Solaris: servicios IP</i> .

Si desea información sobre los procedimientos de instalación del software de cluster, 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 3.3 3/13 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 cluster.

- 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 activa, cambie las direcciones IP del adaptador que se va a quitar a otro adaptador presente en el grupo de varias rutas con el comando `if_mpadm(1M)`.
- Si intenta eliminar una tarjeta de interfaz de red pública sin haberla inhabilitado como interfaz de red pública activa, Oracle Solaris Cluster rechaza la operación e identifica la interfaz que habría sido afectada por la operación.



Precaución – En el caso de grupos de varias rutas con dos adaptadores, si el adaptador de red restante sufre un error mientras se elimina la reconfiguración dinámica en el adaptador de red inhabilitado, la disponibilidad se verá afectada. El adaptador que se conserva no tiene ningún elemento al que migrar tras error mientras dure la operación de reconfiguración dinámica.

Aplice los procedimientos siguientes en el orden indicado al realizar operaciones de reconfiguración dinámica en interfaces de redes públicas.

TABLA 7-4 Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas

Tarea	Instrucciones
1. Cambiar las direcciones IP del adaptador que se va a quitar a otro adaptador del grupo de rutas múltiples mediante <code>if_mpadm</code>	Página del comando <code>man if_mpadm(1M)</code> . Parte V, “IPMP” de Administración de Oracle Solaris: servicios IP.
2. Quitar el adaptador del grupo de rutas múltiples mediante el comando <code>ifconfig</code>	Página del comando <code>man ifconfig(1M)</code> Parte V, “IPMP” de Administración de Oracle Solaris: servicios IP.
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 cluster:

- [“Adición de nodos a un cluster” en la página 219](#)
- [“Eliminación de nodos de un cluster” en la página 225](#)

Para obtener información acerca de las tareas de mantenimiento del cluster, consulte el [Capítulo 9, “Administración del cluster”](#).

Adición de nodos a un cluster

Esta sección describe cómo agregar un nodo a un cluster global o a un cluster de zona. Puede crear un nuevo nodo del cluster de zona sobre un nodo del cluster global que aloje como host el cluster de zona, siempre y cuando el nodo del cluster global no aloje ya un nodo de ese cluster de zona en concreto. No es posible convertir en un nodo de cluster de zona un nodo sin voto ya existente y que reside en un cluster global.

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

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

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

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

En este capítulo, `phys - schost#` refleja una solicitud de cluster 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 cluster ya existente. Efectúe las tareas en el orden en que se muestran.

Tarea	Instrucciones
Instalar el adaptador de host en el nodo y comprobar que las interconexiones ya existentes del cluster sean compatibles con el nuevo nodo.	Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual
Agregar almacenamiento compartido.	Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual
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 cluster.	Capítulo 2, “Instalación del software en los nodos del cluster global” de Guía de instalación del software de Oracle Solaris Cluster
Agregar el nuevo nodo a un cluster existente.	“Cómo agregar un nodo a un cluster existente” en la página 220
Si el cluster 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 Oracle Solaris Cluster Geographic Edition System Administration Guide

▼ Cómo agregar un nodo a un cluster existente

Antes de agregar una máquina virtual o un host de Oracle Solaris a un cluster global o un cluster 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 cluster.

Para obtener información sobre la instalación de hardware, consulte el [Oracle Solaris Cluster 3.3 3/13 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 cluster al agregar su nombre de nodo a la lista de nodos autorizados en ese cluster.

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

Este procedimiento proporciona las formas largas de los comandos de 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 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](#).
- 2 Asegúrese de que los conmutadores, adaptadores y cables hayan sido creados con la utilidad `scinstall`.
- 3 En el nodo que se está agregando, use la utilidad `scinstall`, elija opciones de menú para agregar un nodo de cluster y siga los indicadores.
- 4 Para agregar manualmente un nodo a un cluster de zona, siga estos pasos adicionales:
 - a. Especifique el nombre del nodo virtual y el host de Oracle Solaris. También debe especificar un recurso de red que se utilizará para la comunicación con red pública en cada nodo. En el ejemplo siguiente, el nombre de zona es `sczone`, y `bge0` es el adaptador de red pública en ambas máquinas.


```
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-1
clzc:sczone:node>set hostname=hostname1
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname1
clzc:sczone:node:net>set physical=bge0
clzc:sczone:node:net>end
clzc:sczone:node>end
clzc:sczone>add node
clzc:sczone:node>set physical-host=phys-cluster-2
clzc:sczone:node>set hostname=hostname2
clzc:sczone:node>add net
clzc:sczone:node:net>set address=hostname2
clzc:sczone:node:net>set physical=bge0
clzc:sczone:node:net>end
clzc:sczone:node>end
```
 - b. Después de configurar el nodo, reinicie el nodo en modo de cluster e instale el cluster de zona en el nodo.

```
# clzc install zone-clustername
```

Para obtener instrucciones sobre cómo configurar el nodo de cluster de zona, consulte “Configuración de un cluster de zona” de *Guía de instalación del software de Oracle Solaris Cluster*.

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

Siga los indicadores de `clsetup`. Esta opción indica al cluster que omita todas las solicitudes llegadas a través de la red pública procedentes de cualquier equipo nuevo que intente agregarse a sí mismo al cluster. Cuando haya terminado, salga de la utilidad `clsetup`.

Ejemplo 8–1 Adición de nodos de un cluster 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 cluster 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 cluster, consulte la [Tabla 8–1](#), "Mapa de tareas: agregar un nodo del cluster".

Para agregar un nodo a un grupo de recursos ya existente, consulte [Oracle Solaris Cluster Data Services Planning and Administration Guide](#).

Creación de un nodo sin voto (zona) en un cluster global

En esta sección, se incluyen los procedimientos y la información que se indican a continuación para crear un nodo sin voto, denominado *zona*, en un nodo de cluster global.

▼ Creación de un nodo sin voto en un cluster global

- 1 **Conviértase en superusuario en el nodo del cluster global en el que vaya a crear un nodo que no sea de votación.**

Debe encontrarse en la zona global.

2 Compruebe en todos los nodos que los servicios multiusuario para la Utilidad de gestión de servicios (SMF) estén en línea.

Si los servicios todavía no están en línea para un nodo, espere hasta que cambie el estado y aparezca como en línea antes de continuar con el siguiente paso.

```
phys-schost# svcs multi-user-server node
STATE          STIME      FMRI
online          17:52:55  svc:/milestone/multi-user-server:default
```

3 Configure, instale e inicie la nueva zona.

Nota – Debe establecer la propiedad autoboot en true para admitir el uso de las funciones de grupos de recursos en un nodo del cluster global que no sea de votación.

Siga los procedimientos descritos en la documentación de Oracle Solaris:

- a. Lleve a cabo los procedimientos establecidos en el **Capítulo 18, “Planificación y configuración de zonas no globales (tareas)”** de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.
- b. Lleve a cabo los procedimientos establecidos en **“Instalación e inicio de zonas”** de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.
- c. Lleve a cabo los procedimientos establecidos en **“Cómo iniciar una zona”** de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

4 Asegúrese de que la zona presente el estado ready (listo).

```
phys-schost# zoneadm list -v
ID  NAME      STATUS    PATH
0   global    running   /
1   my-zone   ready     /zone-path
```

5 (Opcional) En una zona de direcciones IP compartidas, asigne una dirección IP privada y un nombre de host privado a la zona.

El siguiente comando permite seleccionar y asignar una dirección IP disponible del intervalo de direcciones IP privadas del cluster. También permite asignar el nombre de host privado (o alias de host) especificado a la zona y a la dirección IP privada indicada.

```
phys-schost# clnode set -p zprivatehostname=hostalias node:zone
```

- p Especifica una propiedad.
- zprivatehostname=alias_host Especifica el nombre de host privado o el alias de host de la zona.
- nodo El nombre del nodo.

zona

El nombre del nodo del cluster global que no es de votación.

6 Realice una configuración inicial de zonas internas.

Siga los procedimientos descritos en “[Configuración inicial de la zona interna](#)” de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*. Seleccione uno de los siguientes métodos:

- Inicie una sesión en la zona.
- Utilice el archivo `/etc/sysidcfg`.

7 En el nodo que no es un nodo con voto, modifique el archivo `nsswitch.conf`.

Estos cambios permiten a la zona realizar búsquedas de direcciones IP y nombres de host específicos del cluster.

a. Inicie una sesión en la zona.

```
phys-schost# zlogin -c zonename
```

b. Abra el archivo `/etc/nsswitch.conf` para realizar la edición.

```
sczone# vi /etc/nsswitch.conf
```

c. Agregue el conmutador `cluster` al principio de las consultas de las entradas `hosts` y `netmasks`, seguido del conmutador `files`.

Las entradas modificadas deben ser similares a las siguientes:

```
...
hosts:      cluster files nis [NOTFOUND=return]
...
netmasks:   cluster files nis [NOTFOUND=return]
...
```

d. En las demás entradas, asegúrese de que el conmutador `files` sea el primer conmutador que aparece en la entrada.**e. Salga de la zona.****8 Si ha creado una zona de direcciones IP exclusivas, configure los grupos IPMP en cada archivo `/etc/hostname.interfaz` que se encuentre en esa zona.**

Debe configurar un grupo de IPMP para cada adaptador de red pública que se utilice para el tráfico de servicios de datos en la zona. Esta información no se hereda de la zona global.

Consulte “[Redes públicas](#)” de *Guía de instalación del software de Oracle Solaris Cluster* para obtener más información sobre cómo configurar grupos IPMP en un cluster.

- 9 Configure las asignaciones de nombre y dirección para todos los nombres de host lógicos que utilice la zona.
- a. Agregue asignaciones de nombre y dirección al archivo `/etc/inet/hosts` en la zona.
- Esta información no se hereda de la zona global.
- b. Si utiliza un servidor de nombres, agregue la asignación de nombre y dirección.

Eliminación de nodos de un cluster

Esta sección ofrece instrucciones sobre cómo eliminar un nodo de un cluster global o de un cluster de zona. También es posible eliminar un cluster de zona específico de un cluster global. La tabla siguiente muestra una lista con las tareas que se deben realizar para eliminar un nodo de un cluster 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 clusters de aplicación real de Oracle](#). Una vez que finalice ese proceso, 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 cluster de zona, iniciar sesión en el cluster de zona y evacuar el nodo de cluster de zona que se encuentra en el nodo físico que se está desinstalando. Luego, eliminar el nodo del cluster de zona antes de desactivar el nodo físico.	<code>clnode evacuate nodo</code> “Eliminación de un nodo de un cluster de zona” en la página 226
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 cluster.	
Eliminar el nodo de todos los grupos de dispositivos.	“Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)” en la página 136

TABLA 8-2 Mapa de tareas: eliminar un nodo (Continuación)

Tarea	Instrucciones
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 cluster compuesto por dos nodos. “Eliminación de un dispositivo de quórum” en la página 181 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 cluster” en la página 182
Poner el nodo que se va a eliminar en un modo que no sea de cluster.	“Cómo poner un nodo en estado de mantenimiento” en la página 251
Eliminar un nodo de un cluster de zona.	“Eliminación de un nodo de un cluster de zona” en la página 226
Eliminar un nodo de la configuración de software del cluster.	“Eliminación de un nodo de la configuración de software del cluster” en la página 227
(Opcional) Desinstalar el software de Oracle Solaris Cluster de un nodo de cluster.	“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de cluster” en la página 255

▼ Eliminación de un nodo de un cluster de zona

Puede eliminar un nodo de un cluster de zona deteniendo el nodo, desinstalándolo y eliminándolo de la configuración. Si posteriormente decidiese volver a agregar el nodo al cluster de zona, siga las instrucciones que se indican en la [Tabla 8-1](#). La mayoría de estos pasos se realizan desde el nodo del cluster global.

Antes de empezar Consulte la [Tabla 8-2](#) para ver si hay pasos que debe realizar antes de llevar a cabo este procedimiento.

- 1 **Conviértase en superusuario en un nodo del cluster global.**
- 2 **Cierre el nodo del cluster de zona que desee eliminar; para ello, especifique el nodo y su cluster de zona.**

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

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

3 Elimine el nodo de todos los grupos de recursos en el cluster de zona.

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

4 Desinstale el nodo del cluster de zona.

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

5 Elimine el nodo del cluster de zona de la configuración.

Use los comandos siguientes:

```
phys-schost# clzonecluster configure zoneclustername
```

```
clzc:sczone> remove node physical-host=node
```

```
clzc:sczone> exit
```

6 Compruebe que el nodo se haya eliminado del cluster de zona.

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

▼ Eliminación de un nodo de la configuración de software del cluster

Siga este procedimiento para eliminar un nodo del cluster global.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global.****3 Inicie el nodo del cluster global que desee eliminar en un modo que no sea el de cluster. Para un nodo de un cluster de zona, siga las instrucciones indicadas en [“Eliminación de un nodo de un cluster de zona” en la página 226](#) 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.**

Aparece el menú de GRUB, que es similar al siguiente:

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

- 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.**

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

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

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

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

```
grub edit> kernel /platform/i86pc/multiboot -x
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a)                        |
| kernel /platform/i86pc/multiboot -x   |
+-----+
```

```
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

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

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 cluster. Para arrancarlo en un modo que no sea de cluster, 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 iniciar, ejecute el comando siguiente en cualquier nodo del cluster activo: **`clnode clear`**
`-F<node-to-be-removed>`. Verifique la eliminación del nodo mediante la ejecución de **`clnode status <nodename>`**.

4 Elimine el nodo del cluster.

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 cluster, éste debe establecerse en un modo que no sea de cluster de forma que no quede ningún nodo activo en el cluster.

5 Compruebe la eliminación del nodo desde otro nodo del cluster.

```
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 cluster” en la página 255](#).

- Si no tiene previsto desinstalar el software de Oracle Solaris Cluster del nodo eliminado, puede eliminar físicamente el nodo del cluster mediante la eliminación de las conexiones de hardware como se describe en el [Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#).

Ejemplo 8–2 Eliminación de un nodo de la configuración de software del cluster

Este ejemplo muestra cómo eliminar un nodo (phys-schost-2) desde un cluster. El comando `clnode remove` se ejecuta en un modo que no sea de cluster desde el nodo que desea eliminar del cluster (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 cluster” en la página 255.

Para obtener información sobre los procedimientos de hardware, consulte el [Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual](#).

Para obtener una lista completa de las tareas necesarias para eliminar un nodo de cluster, consulte la [Tabla 8–2](#).

Para agregar un nodo a un cluster existente, consulte “Cómo agregar un nodo a un cluster existente” en la página 220.

▼ **Eliminación de un nodo sin voto (zona) de un cluster global**

- 1 **Conviértase en superusuario en el nodo de cluster global donde creó el nodo sin voto.**
- 2 **Suprima el nodo sin voto del sistema.**
Siga los procedimientos en “Supresión de una zona no global del sistema” de *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

▼ Eliminación de la conectividad entre una matriz y un único nodo, en un cluster con conectividad superior a dos nodos

Use este procedimiento para desconectar una matriz de almacenamiento de un nodo único de un cluster, en un cluster que tenga conectividad de tres o cuatro nodos.

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

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

- 1 **Realice copias de seguridad de todas las tablas de bases de datos, los servicios de datos y los volúmenes que estén asociados con la matriz de almacenamiento que vaya a eliminar.**
- 2 **Determine los grupos de recursos y grupos de dispositivos que se estén ejecutando en el nodo que se va a desconectar.**

```
phys-schost# clresourcegroup status
phys-schost# cldevicegroup status
```

- 3 **Si es necesario, traslade todos los grupos de recursos y grupos de dispositivos fuera del nodo que se vaya a desconectar.**



Caution (SPARC only) – Si el cluster 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 cambia todos los grupos de recursos desde los nodos con voto o sin voto del nodo especificado hasta el siguiente nodo por orden de preferencia, sea con voto o sin voto.

- 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 251](#).

5 Elimine el nodo de los grupos de dispositivos.

Si usa un disco sin procesar, utilice el comando `cldevicegroup` para eliminar los grupos de dispositivos. Para obtener más información, consulte la página del comando `man cldevicegroup(1CL)`.

6 Para cada grupo de recursos que contenga un recurso HASStoragePlus, elimine el nodo de la lista de nodos del grupo de recursos.

```
phys-schost# clresourcegroup remove-node -z zone -n node + | resourcegroup  
node      El nombre del nodo.
```

`zona` El nombre del nodo sin voto que puede controlar el grupo de recursos. Especifique `zone` sólo si ha proporcionado un nodo sin voto al crear el grupo de recursos.

Consulte *Oracle Solaris Cluster Data Services Planning and Administration Guide* si desea más información sobre cómo cambiar la lista de nodos de un grupo de recursos.

Nota – Los nombres de los tipos, los grupos y las propiedades de recursos distinguen mayúsculas y minúsculas cuando se ejecuta `clresourcegroup`.

7 Si la matriz de almacenamiento que va a eliminar es la última que está conectada al nodo, desconecte el cable de fibra óptica que comunica el nodo y el concentrador o conmutador que está conectado a esta matriz de almacenamiento. Si no es así, omita este paso.**8 Si va a quitar el adaptador de host del nodo que va a desconectar, cierre el nodo. Si va a eliminar el adaptador de host del nodo que vaya a desconectar, pase directamente al Paso 11.****9 Elimine el adaptador de host del nodo.**

Para el procedimiento de eliminar los adaptadores de host, consulte la documentación relativa al nodo.

10 Encienda el nodo pero no lo haga arrancar.**11 Si se ha instalado software Oracle RAC, elimine el paquete de software Oracle RAC del nodo que va a desconectar.**

```
phys-schost# pkgrm SUNWscum
```



Caution (SPARC only) – Si no elimina el software Oracle RAC del nodo que desconectó, dicho nodo experimentará un error grave al volver a introducirlo en el cluster, lo que podría provocar una pérdida de la disponibilidad de los datos.

12 Arranque el nodo en modo de cluster.

- 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. Aparece el menú de GRUB, que es similar al siguiente:

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

- 13 En el nodo, actualice el espacio de nombre del dispositivo mediante la actualización de las entradas `/devices` y `/dev`.

```
phys-schost# devfsadm -C
cldevice refresh
```

- 14 Vuelva a poner en línea los grupos de dispositivos.

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 252](#).

▼ 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 cluster, siga el procedimiento detallado a continuación.

- 1 Intente volver a unir el nodo al cluster global. Realice este procedimiento sólo en un cluster global.

```
phys-schost# boot
```

- 2 ¿Se ha vuelto a unir correctamente el nodo al cluster?

- 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 cluster, 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 135](#).
 - 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 cluster” en la página 255](#) y repita el procedimiento.
- 3 Si no se ha podido volver a unir el nodo al cluster, 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 cluster” en la página 255](#) y repita el procedimiento.

Administración del cluster

Este capítulo presenta los procedimientos administrativos que afectan a todo un cluster global o a un cluster de zona:

- “Información general sobre la administración del cluster” en la página 235
- “Cómo realizar tareas administrativas del cluster de zona” en la página 267
- “Resolución de problemas” en la página 275

Si desea información sobre cómo agregar o quitar un nodo del cluster, consulte el [Capítulo 8](#), “Adición y eliminación de un nodo”.

Información general sobre la administración del cluster

Esta sección describe cómo realizar tareas administrativas para todo el cluster global o de zona. La tabla siguiente enumera estas tareas administrativas y los procedimientos asociados. Las tareas administrativas del cluster suelen efectuarse en la zona global. Para administrar un cluster de zona, al menos un equipo que almacenará como host el cluster de zona debe estar activado en modo de cluster. No todos los nodos cluster de zona deben estar activados y en funcionamiento; Oracle Solaris Cluster reproduce cualquier cambio de configuración si el nodo que está fuera del cluster se vuelve a unir a éste.

Nota – De manera predeterminada, la administración de energía está desactivada para que no interfiera con el cluster. Si habilita la administración de energía en un cluster de un solo nodo, el cluster 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 cluster global. El indicador de solicitud de shell interactivo de `clzonecluster` es `clzc: schost>`.

TABLA 9-1 Lista de tareas: administrar el cluster

Tarea	Instrucciones
Agregar o eliminar un nodo de un cluster	Capítulo 8, “Adición y eliminación de un nodo”
Cambiar el nombre del cluster	“Cambio del nombre del cluster” en la página 236
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 238
Permitir o denegar que se agreguen nodos nuevos al cluster	“Uso de la autenticación del nodo del cluster nuevo” en la página 238
Cambiar la hora de un cluster mediante el protocolo de hora de red (NTP)	“Restablecimiento de la hora del día en un cluster” en la página 240
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 242
Agregar o cambiar el nombre de host privado	“Agregación de un nombre de host privado para un nodo sin voto en un cluster global” en la página 246 “Cómo cambiar un nombre de host privado de nodo” en la página 243
Poner un nodo de cluster en estado de mantenimiento	“Cómo poner un nodo en estado de mantenimiento” en la página 251
Cambiar el nombre de un nodo	“Cómo cambiar el nombre de un nodo” en la página 248
Sacar un nodo del cluster fuera del estado de mantenimiento	“Cómo sacar un nodo del estado de mantenimiento” en la página 252
Desinstalar el software del cluster de un nodo del cluster	“Cómo instalar el software de Oracle Solaris Cluster desde un nodo de cluster” en la página 255
Agregar y administrar una MIB de eventos de SNMP	“Cómo habilitar una MIB de eventos de SNMP” en la página 260 “Cómo agregar un usuario de SNMP a un nodo” en la página 263
Configurar los límites de carga de un nodo	“Cómo configurar límites de carga en un nodo” en la página 266
Mover un cluster de zona, preparar un cluster de zona para aplicaciones, eliminar un cluster de zona	“Cómo realizar tareas administrativas del cluster de zona” en la página 267

▼ Cambio del nombre del cluster

Si es necesario, el nombre del cluster puede cambiarse tras la instalación inicial.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global.**
- 2 Inicie la utilidad `clsetup`.**
`phys-schost# clsetup`
 Aparece el menú principal.
- 3 Para cambiar el nombre del cluster, escriba el número correspondiente a la opción **Other Cluster Properties (Otras propiedades del cluster)**.**
 Aparece el menú **Other Cluster Properties (Otras propiedades del cluster)**.
- 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 cluster nuevo, elimine la etiqueta de Oracle Solaris Cluster y reinicie el cluster. Para eliminar la instancia de la etiqueta de servicio de Oracle Solaris Cluster, realice los subpasos siguientes en todos los nodos del cluster.**
 - 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 cluster.**
`phys-schost# reboot`

Ejemplo 9–1 Cambio del nombre del cluster

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

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

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

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

Durante la instalación de Oracle Solaris Cluster, se les asigna automáticamente a todos los nodos un número exclusivo de ID de nodo. El número de ID de nodo se asigna a un nodo en el orden en que se une al cluster 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 cluster 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 cluster global o de zona no hace falta ser superusuario. Un paso de este procedimiento se realiza desde un nodo de cluster global. El otro paso se efectúa desde un nodo de cluster de zona.

- 1 **Utilice el comando `clnode show` para mostrar la información de configuración del cluster para el cluster 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 cluster de zona. El nodo de cluster de zona tiene el mismo ID que el nodo de cluster 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 cluster 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 cluster nuevo

Oracle Solaris Cluster permite determinar si se pueden agregar nodos nuevos al cluster global y el tipo de autenticación que puede utilizar. Puede permitir que cualquier nodo nuevo se una al cluster por la red pública, denegar que los nodos nuevos se unan al cluster o indicar que un nodo en particular se una al cluster. 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 del comando [man keyserv\(1M\)](#) y [publickey\(4\)](#).

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global.**
- 2 Inicie la utilidad `clsetup`.**
`phys-schost# clsetup`
Aparece el menú principal.
- 3 Para trabajar con la autenticación del cluster, escriba el número correspondiente a la opción de nodos nuevos.**
Aparece el menú Nuevos nodos.
- 4 Seleccione desde el menú y siga las instrucciones que aparecen en la pantalla.**

Ejemplo 9-3 Procedimiento para evitar que un equipo nuevo se agregue al cluster global

La utilidad `clsetup` genera el comando `claccess`. El ejemplo siguiente muestra el comando `claccess`, que evita que equipos nuevos se agreguen al cluster.

```
phys-schost# claccess deny -h hostname
```

Ejemplo 9-4 Procedimiento para permitir que todos los equipos nuevos se agreguen al cluster 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 cluster.

```
phys-schost# claccess allow-all
```

Ejemplo 9-5 Procedimiento para especificar que se agregue un equipo nuevo al cluster 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 cluster.

```
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 cluster.

```
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 cluster.

```
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 cluster. Para obtener más información, consulte las páginas del comando `man keyserv(1M)` y `publickey(4)`.

▼ Restablecimiento de la hora del día en un cluster

El software de Oracle Solaris Cluster usa el protocolo de hora de red (NTP) para mantener la sincronización temporal entre los nodos del cluster. Los ajustes del cluster 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://docs.oracle.com/cd/E19065-01/servers.10k/>.



Precaución – Si utiliza NTP, no intente ajustar la hora del cluster mientras se encuentre activo y en funcionamiento. No ajuste la hora con los comandos `date`, `rdate` o `svcadm` de forma interactiva o dentro de las secuencias de comandos `cron`. Para obtener más información, consulte las páginas del comando `man date(1)`, `rdate(1M)`, `xntpd(1M)`, `svcadm(1M)` o `cron(1M)`.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global.**
- 2 **Cierre el cluster 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 cluster.

- 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.

Aparece el menú de GRUB, que es similar al siguiente:

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

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

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

- 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.

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a)                  |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.

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

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

```
grub edit> kernel /platform/i86pc/multiboot -x
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a)                                |
| kernel /platform/i86pc/multiboot -x           |
| module /platform/i86pc/boot_archive           |
+-----+
```

```
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

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

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 cluster. Para arrancarlo en el modo que no es de cluster, 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 a ese nodo con el comando `rdate`.**

```
phys-schost# rdate hostname
```

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

- 7 Arranque cada nodo para reiniciar el cluster.**

```
phys-schost# reboot
```

- 8 Compruebe que el cambio se haya hecho en todos los nodos del cluster.**

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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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
```

`tc_name` Especifica el nombre del concentrador de terminales.

`tc_port_number` 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 cluster instalado en un sitio.

2 Cierre el nodo del cluster mediante el comando `clnode evacuate` y, a continuación, mediante el comando `shutdown`. El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. El comando también cambia todos los grupos de recursos desde el nodo con voto o sin voto especificado del cluster global hasta el siguiente nodo por orden de preferencia, sea con voto o sin voto.

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



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

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 cluster una vez finalizada la instalación.

Durante la instalación inicial del cluster 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 cluster los asigna.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster, 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 del protocolo de hora de red (NTP) en cada nodo del cluster.

Utilice el comando `svcadm` para cerrar el daemon del protocolo de hora de red (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(1CL)` para cambiar el nombre de host privado del nodo adecuado.

Ejecute la utilidad desde uno solo de los nodos del cluster.

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

4 Escriba el número correspondiente a la opción de nombre de host privado.

5 Escriba el número correspondiente a 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 cluster. El vaciado evita que las aplicaciones del cluster 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 el archivo de configuración NTP, actualice el archivo de configuración NTP (`ntp.conf` o `ntp.conf.cluster`) 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. La plantilla predeterminada está preconfigurada con 16 nodos. Por lo general, el archivo `ntp.conf.cluster` es el mismo en todos los nodos del cluster.

b. Compruebe que pueda realizar un ping correctamente en el nombre de host privado desde todos los nodos del cluster.

c. Reinicie el daemon de NTP.

Realice este paso en cada nodo del cluster.

Utilice el comando `svcadm` para reiniciar el daemon de NTP.

```
# svcadm enable ntp
```

8 Habilite todos los recursos de servicio de datos y otras aplicaciones 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 ejemplo siguiente, se cambia el nombre de host privado de `clusternode2-priv` a `clusternode4-priv`, en el nodo `phys-schost-2`.

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

```
peer clusternode4-priv
peer clusternode3-priv
phys-schost-1# ping clusternode4-priv
phys-schost-1# /etc/init.d/xntpd start
[Enable all applications and data services disabled at the beginning of the procedure.]
```

▼ Agregación de un nombre de host privado para un nodo sin voto en un cluster global

Utilice este procedimiento para agregar un nombre de host privado a un nodo sin voto en un cluster global una vez finalizada la instalación. En los procedimientos tratados en este capítulo, `phys-schost#` refleja una solicitud de cluster global. Realice este procedimiento sólo en un cluster global.

- 1 Ejecute la utilidad `clsetup(1CL)` para agregar un nombre de host privado en la zona adecuada.
`phys-schost# clsetup`
- 2 Escriba el número correspondiente a la opción de nombres de host privados y pulse la tecla de retorno.
- 3 Escriba el número correspondiente a la opción de agregar un nombre de host privado de zona y pulse la tecla de retorno.

Responda las preguntas cuando se lo solicite. No existe un nombre predeterminado para el sistema privado de un nodo sin votación del cluster global. Necesitará proporcionar un nombre de host.

▼ Cambio de nombre de host privado en un nodo sin voto de un cluster global

Utilice este procedimiento para cambiar el nombre de host privado de un nodo sin voto una vez finalizada la instalación.

Durante la instalación inicial del cluster se asignan nombre de host privados. El nombre de host privado usa el formato `clusternode< nodeid>-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 cluster los asigna.

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

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

1 En todos los nodos del cluster global, desactive los recursos de servicio de datos u otras aplicaciones que puedan almacenar en caché los nombres de host privados.

`phys-schost# clresource disable resource1, resource2`

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 Ejecute la utilidad `clsetup` para cambiar el nombre de host privado del nodo sin voto adecuado en el cluster global.

`phys-schost# clsetup`

Debe realizar este paso únicamente desde uno de los nodos del cluster. 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 cluster.

3 Escriba el número correspondiente a la opción de nombres de host privados y pulse la tecla de retorno.

4 Escriba el número correspondiente a la opción de agregar un nombre de host privado de zona y pulse la tecla de retorno.

No existe opción predeterminada para un nodo sin votación de un nombre de host privado del cluster global. Se debe proporcionar un nombre de host.

5 Escriba el número correspondiente a la opción de cambiar un nombre de host privado de zona.

Responda las preguntas cuando se lo solicite. Se le solicitará el nombre del nodo sin voto cuyo nombre de host privado esté cambiando (`clusternode< nodeid> -priv`) y el nuevo nombre de host privado.

6 Purgue la antememoria del servicio de nombres.

Efectúe este paso en todos los nodos del cluster. El vaciado evita que las aplicaciones del cluster y los servicios de datos intenten acceder al nombre de host privado anterior.

```
phys-schost# nscd -i hosts
```

7 Active todos los recursos de servicio de datos y otras aplicaciones desactivados en el [Paso 1](#).

▼ **Supresión de un nombre de host privado para un nodo sin voto en un cluster global**

Utilice este procedimiento para suprimir un nombre de host privado para un nodo sin voto en un cluster global. Realice este procedimiento sólo en un cluster global.

- 1 Ejecute la utilidad clsetup para suprimir un nombre de host privado en la zona adecuada.**
Para obtener más información, consulte la página del comando `man clsetup(1CL)`.
- 2 Escriba el número correspondiente a la opción de nombre de host privado de zona.**
- 3 Escriba el número correspondiente a la opción de cambiar un nombre de host privado de zona.**
- 4 Escriba el nombre de host privado del nodo sin voto que va a suprimir.**

▼ **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 cluster global.

- 1 En el cluster 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 cluster Geographic Edition de Oracle Solaris Cluster que está en sociedad con una configuración de Oracle Solaris, debe realizar pasos adicionales. Si el cluster 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 cluster secundario durante el procedimiento de cambio de nombre. Para**

obtener más información sobre los clusters y los nodos de Geographic Edition, consulte el [Capítulo 5, “Administering Cluster Partnerships” de Oracle Solaris Cluster Geographic Edition System Administration Guide](#).

- 3 Cambie los nombres de host de Oracle Solaris siguiendo los pasos indicados en [“Cómo cambiar el nombre de host de un sistema” de Guía de administración del sistema: Administración avanzada](#), pero *no* reinicie el sistema al final del procedimiento. En su lugar, cierre el cluster una vez completados los pasos.
- 4 Inicie todos los nodos del cluster en un modo que no sea de cluster.

```
ok> boot -x
```
- 5 En un modo que no sea de cluster 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 cluster.
- 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 cluster.

```
# sync;sync;sync;/etc/reboot
```
- 9 Verifique que el nodo muestre el nuevo nombre.

```
# clnode status -v
```
- 10 Si va a cambiar el nombre de un nodo de cluster de Geographic Edition y el cluster asociado del cluster 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 cluster 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 cluster 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 250](#) 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 248](#). Este paso es opcional.

- 1 En el cluster 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 cluster.

- a. En el modo de cluster, 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.

```
# clr disable apache-lh-res
```

- c. Proporcione la nueva lista de nombres de host.

```
# clr 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

```
# clr enable apache-lh-res
```

- f. Ponga en línea los grupos de recursos de Apache.

```
# clrg online -eM apache-rg
```

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

```
# clr status apache-rs
```

▼ Cómo poner un nodo en estado de mantenimiento

Ponga un nodo del cluster 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 cluster.

Cuando un cluster 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 cluster en estado de mantenimiento. Para obtener más información, consulte la página del comando `man clquorum(1CL)`.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global que pone en estado de mantenimiento.**
- 2 **Evacue todos los grupos de recursos y de dispositivos del nodo. El comando `clnode evacuate` cambia todos los grupos de recursos y de dispositivos, incluidos todos los nodos sin voto, del nodo especificado al siguiente nodo por orden de preferencia.**
`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 cluster 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 cluster 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 cluster global en estado de mantenimiento

En el siguiente ejemplo se pone un nodo de cluster 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 252.](#)

▼ Cómo sacar un nodo del estado de mantenimiento

Utilice el siguiente procedimiento para volver a poner en línea un nodo del cluster global y restablecer el recuento de votos de quórum al valor predeterminado. En los nodos del cluster, 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 cluster 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 cluster.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global distinto del que está en estado de mantenimiento.**
- 2 **Según el número de nodos que tenga en la configuración del cluster global, realice uno de los siguientes pasos:**
 - Si tiene dos nodos en la configuración del cluster, vaya al [Paso 4](#).
 - Si tiene más de dos nodos en la configuración del cluster, 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 cluster 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 cluster 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 cluster 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. Aparece el menú de GRUB, que es similar al siguiente:

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

`phys-schost-1# clquorum status`

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

      Needed   Present   Possible
      -----   -
      4         6         6

--- Quorum Votes by Node ---

Node Name      Present   Possible   Status
-----
phys-schost-2   1         1         Online
phys-schost-3   1         1         Online

--- Quorum Votes by Device ---

Device Name      Present   Possible   Status
-----
/dev/did/rdisk/d3s2  1         1         Online
/dev/did/rdisk/d17s2 0         1         Online
/dev/did/rdisk/d31s2 1         1         Online
```

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

Realice este procedimiento para desinstalar el software de Oracle Solaris Cluster de un nodo del cluster global antes de desconectarlo de una configuración de cluster completamente establecida. Puede seguir este procedimiento para desinstalar el software desde el último nodo de un cluster.

Nota – No siga este procedimiento para desinstalar el software Oracle Solaris Cluster desde un nodo que todavía no se haya unido al cluster o que aún esté en modo de instalación. En cambio, vaya a “Desinstalación del software de Oracle Solaris Cluster para solucionar problemas de instalación” en la [Guía de instalación del software de Oracle Solaris Cluster](#).

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster.**

Consulte la [Tabla 8–2](#).

Nota – Compruebe que haya eliminado el nodo de la configuración del cluster mediante `clnode remove` antes de continuar con este procedimiento.

- 2 Conviértase en superusuario en un miembro activo del cluster global que *no sea* el nodo del cluster global que está desinstalando. Realice este procedimiento desde un nodo del cluster global.**
- 3 En el miembro de cluster activo, agregue el nodo que desea desinstalar a la lista de autenticación de nodos del cluster.**

```
phys-schost# claccess allow -h hostname
```

-h Especifica el nombre del nodo que se va a agregar a la lista de autenticación del nodo.

También puede usar la utilidad `clsetup`. Consulte la página del comando `man clsetup(1CL)` y “[Cómo agregar un nodo a un cluster existente](#)” en la [página 220](#) para conocer los procedimientos.

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

5 Si tiene un cluster de zona, desinstálelo.

```
phys-schost# clzonecluster uninstall -F zoneclustername
```

Para ver los pasos específicos, consulte [“Cómo eliminar un cluster de zona” en la página 270.](#)

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

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

```
# shutdown -g0 -y -i0ok 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
```

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

8 Si tiene previsto volver a instalar el software de Oracle Solaris Cluster en este nodo, elimine la entrada de Oracle Solaris Cluster del registro de productos de Sun Java Enterprise System (Java ES).

Si el registro de productos de Java ES contiene un registro de que el software de Oracle Solaris Cluster ya se instaló, el instalador de Java ES muestra el componente de Oracle Solaris Cluster atenuado y no permite la reinstalación.

a. Inicie el desinstalador de Java ES.

Ejecute el siguiente comando, donde *ver* es la versión de la distribución de Java ES desde donde instaló el software de Oracle Solaris Cluster.

```
# /var/sadm/prod/SUNWentsys $ver$ /uninstall
```

b. Siga las indicaciones para seleccionar Oracle Solaris Cluster para su desinstalación.

Para obtener más información sobre el uso del comando `uninstall`, consulte el [Capítulo 8, “Uninstalling” de *Sun Java Enterprise System 5 Update 1 Installation Guide for UNIX*.](#)

- 9 Si no tiene previsto volver a instalar el software Oracle Solaris Cluster en este cluster, desconecte los cables y el conmutador de transporte, si los hubiera, de los otros dispositivos del cluster.
 - 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 – Si utiliza un dispositivo de archivo de bucle de retorno (lofi), el desinstalador de Java ES quita automáticamente el archivo lofi, denominado `.globaldevices`. Para obtener más información acerca de la migración de un espacio de nombre de dispositivos globales a lofi, consulte [“Migración del espacio de nombre de dispositivos globales” en la página 127](#).

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 cluster no eliminadas

Los siguientes mensajes de error indican que el nodo del cluster global que ha eliminado todavía tiene sistemas de archivos del cluster 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 cluster” en la página 255](#) y repita el procedimiento. Compruebe que haya realizado correctamente el [Paso 7](#) 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.
```

Secuencia de comandos de desinstalación faltante

Si no utilizó el programa `installer` para instalar o actualizar el software de Sun Cluster u Oracle Solaris Cluster que ahora desea quitar, no hay ninguna secuencia de comandos de desinstalación que se pueda utilizar para esa versión de software. En cambio, siga los pasos indicados a continuación para desinstalar el software.

▼ Desinstalación del software de Sun Cluster 3.1 y 3.2 sin una secuencia de comandos de desinstalación

- 1 Conviértase en superusuario o asuma una función que proporcione la autorización de RBAC `solaris.cluster.modify`.
- 2 Cambie a un directorio que no esté asociado con ningún paquete de Sun Cluster.
`# cd /directory`
- 3 Desinstale el software de Sun Cluster desde el nodo.
`# scinstall -r`
- 4 Cambie el nombre del archivo `productregistry` para permitir una nueva instalación del software en el futuro.
`# mv /var/sadm/install/productregistry /var/sadm/install/productregistry.sav`

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 cluster en tiempo real. Si se habilita, el administrador de SNMP envía automáticamente notificaciones de captura a todos los sistemas definidos en el comando `clsnmp host`. La MIB mantiene una tabla de sólo lectura con los 50 eventos más actuales. Debido a que los clusters 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
Activar una MIB de eventos de SNMP.	“Cómo habilitar una MIB de eventos de SNMP” en la página 260
Desactivar una MIB de eventos de SNMP.	“Cómo deshabilitar una MIB de eventos de SNMP” en la página 260
Cambiar una MIB de eventos de SNMP.	“Cómo cambiar una MIB de eventos de SNMP” en la página 261
Agregar 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 262
Eliminar un host de SNMP.	“Cómo deshabilitar un host de SNMP para que no reciba capturas de SNMP en un nodo” en la página 262

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

Tarea	Instrucciones
Agregar un usuario de SNMP.	“Cómo agregar un usuario de SNMP a un nodo” en la página 263
Eliminar un usuario de SNMP.	“Cómo eliminar un usuario de SNMP de un nodo” en la página 264

▼ **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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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.**
- 260

Guía de administración del sistema de Oracle Solaris Cluster • Marzo de 2013, E39395-01

2 Deshabilite la MIB de eventos de SNMP.

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

-n node 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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=value MIB
```

-n node 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=value 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 `SNMPcommunity` 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 node`

Especifica el nombre del *nodo* del host de SNMP al que se ha proporcionado acceso a las MIB de SNMP en el cluster. 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 cluster.

▼ **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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 *SNMPcommunity**

Especifica el nombre de la comunidad de SNMP de la que se ha eliminado el host de SNMP.

*-n *node**

Especifica el nombre del *node* 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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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>authentication</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>password</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 cluster global. Siga este procedimiento en un cluster global.

Este procedimiento proporciona las formas largas de los comandos de 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 *nodo* 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.

usuario Especifica el nombre del usuario de SNMP que desea eliminar.

Configuración de límites de carga

Puede establecer límites de carga para activar la distribución automática de la carga del grupo de recursos entre nodos o zonas. Puede configurar un conjunto de límites de carga para cada nodo del cluster. 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 cluster global o de un cluster 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 cluster 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 cluster 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 cluster de zona zc1 y el grupo de recursos rg2 del cluster de zona zc2.

```
# clresourcegroup set -p RG_affinities=++zc2:rg2 zc1:rg1
```

9 Compruebe el estado de todos los nodos del cluster global y del cluster de zona del cluster.

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

El resultado incluirá toda configuración del límite de carga definida en el nodo o en las zonas no globales.

Cómo realizar tareas administrativas del cluster de zona

En un cluster de zona puede efectuar otras tareas administrativas, por ejemplo, mover la ruta de zona, preparar un cluster de zona para que ejecute aplicaciones y clonar un cluster de zona. Todos estos comandos deben ejecutarse desde el nodo de votación del cluster global.

Puede crear un nuevo cluster de zona o agregar un sistema de archivos o dispositivo de almacenamiento mediante la utilidad clsetup para iniciar el asistente de configuración de cluster de zona. Las zonas en un cluster de zona se configuran cuando ejecuta clzonecluster install -c para configurar los perfiles. Consulte [“Establecimiento del cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster](#) para obtener instrucciones sobre cómo usar la utilidad clsetup o la opción -c config_profile.

Nota – Los comandos de Oracle Solaris Cluster que sólo ejecuta desde el nodo de votación en el cluster global no son válidos para usarlos con los clusters 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 cluster de zona

Tarea	Instrucciones
Mover la ruta de zona a una nueva ruta de zona	clzonecluster move -f zonepath zoneclustername
Preparar el cluster de zona para ejecutar aplicaciones	clzonecluster ready -n nodename zoneclustername

TABLA 9-3 Otras tareas del cluster de zona (Continuación)	
Tarea	Instrucciones
Clonar un cluster de zona	<code>clzonecluster clone -Z target- zoneclustername [-m copymethod] source-zoneclustername</code> Detenga el cluster de zona de origen antes de usar el subcomando <code>clone</code> . El cluster de zona de destino ya debe estar configurado.
Agregar una dirección de red a un cluster de zona	“Cómo agregar una dirección de red a un cluster de zona” en la página 268
Eliminar un cluster de zona	“Cómo eliminar un cluster de zona” en la página 270
Eliminar un sistema de archivos de un cluster de zona	“Cómo eliminar un sistema de archivos de un cluster de zona” en la página 270
Eliminar un dispositivo de almacenamiento de un cluster de zona	“Cómo eliminar un dispositivo de almacenamiento de un cluster de zona” en la página 273
Solucionar problema de desinstalación de nodo	“Resolución de problemas de desinstalación de nodos” en la página 257
Crear, configurar y gestionar 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 259

▼ Cómo agregar una dirección de red a un cluster de zona

Este procedimiento agrega una dirección de red para que la utilice un cluster de zona existente. Una dirección de red se utiliza para configurar el host lógico o los recursos de dirección IP compartida en el cluster de zona. Puede ejecutar la utilidad `clsetup` varias veces para agregar tantas direcciones de red como sea necesario.

- 1 **Conviértase en superusuario en un nodo del cluster global que aloje el cluster de zona.**
- 2 **En el cluster global, configure el sistema de archivos compartidos de cluster que desee utilizar en el cluster de zona.**
Inicie la utilidad `clsetup`.
`phys-schost# clsetup`
Aparece el menú principal.
- 3 **Seleccione la opción de menú Zone Cluster (Cluster de zona).**

- 4 **Seleccione la opción de menú Add Network Address to a Zone Cluster (Agregar una dirección de red a un cluster de zona).**

- 5 **Elija el cluster de zona donde desea agregar la dirección de red.**

- 6 **Seleccione la propiedad para especificar la dirección de red que desea agregar.**

`address=value` Especifica la dirección de red que se utiliza para configurar recursos de dirección IP compartida o host lógico en el cluster de zona. Por ejemplo: 192.168.100.101.

Se admiten los siguientes tipos de direcciones de red:

- Una dirección IPv4 válida, seguida, de forma opcional, por / y una longitud de prefijo.
- Una dirección IPv6 válida, que debe ir seguida de / y una longitud de prefijo.
- Un nombre de host que se resuelve en una dirección IPv4. No se admiten los nombres de host que se resuelven en direcciones IPv6.

Consulte la página del comando `man zonecfg(1M)` para obtener más información sobre las direcciones de red.

- 7 **Para agregar otra dirección de red, escriba a.**

- 8 **Escriba c para guardar los cambios en la configuración.**

Se muestran los resultados de su cambio de configuración. Por ejemplo:

```
>>> Result of Configuration Change to the Zone Cluster(sczone) <<<

Adding network address to the zone cluster...

The zone cluster is being created with the following configuration

/usr/cluster/bin/clzonecluster configure sczone
add net
set address=phys-schost-1
end

All network address added successfully to sczone.
```

- 9 **Cuando haya finalizado, salga de la utilidad `clsetup`.**

▼ Cómo eliminar un cluster de zona

Puede eliminar un cluster de zona específico o usar un comodín para eliminar todos los clusters de zona configurados en el cluster global. El cluster 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 cluster global. Siga todos los pasos de este procedimiento desde un nodo del cluster global.**

- 2 **Elimine todos los grupos de recursos y sus recursos del cluster de zona.**

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

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

- 3 **Detenga el cluster de zona.**

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

- 4 **Desinstale el cluster de zona.**

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

- 5 **Anule la configuración del cluster de zona.**

```
phys-schost# clzonecluster delete zoneclustername
```

Ejemplo 9–11 Eliminación de un cluster de zona de un cluster 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 cluster de zona

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

Los clusters de zona admiten montajes directos para lo siguiente:

- Sistema local de archivos UFS
- Sistema de archivos de cluster Oracle Automatic Storage Management (Oracle ACFS)
- Sistema de archivos independientes QFS
- Sistema de archivos compartidos QFS, cuando se utiliza únicamente para compatibilidad con Oracle RAC
- ZFS (exportado como un conjunto de datos)
- NFS desde dispositivos NAS admitidos

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

- Sistema local de archivos UFS
- Sistema de archivos independientes QFS
- Sistema de archivos compartidos QFS, cuando se utiliza únicamente para compatibilidad con Oracle RAC
- Sistema de archivos de cluster de UFS

Puede configurar un recurso `HASStoragePlus` o `ScalMountPoint` para gestionar el montaje del sistema de archivos. Para obtener instrucciones sobre cómo agregar un sistema de archivos a un cluster de zona, consulte [“Adición de sistemas de archivos a un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster](#).

`phys-schost#` refleja un indicador de cluster global. Este procedimiento proporciona las formas largas de los comandos de 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 cluster global que aloje el cluster de zona. Algunos pasos de este procedimiento se realizan desde un nodo del cluster global. Otros pasos se efectúan desde un nodo del cluster 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 cluster 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 de tipo SUNW.qfs configurados en el cluster 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 cluster. 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 cluster 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 de alta disponibilidad en un cluster 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 cluster 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/dsl/dsk/d0
  raw:                                             /dev/md/dsl/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 cluster 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 cluster 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 cluster de zona

Puede eliminar dispositivos de almacenamiento, como conjuntos de discos de SVM y dispositivos de DID, de un cluster de zona. Siga este procedimiento para eliminar un dispositivo de almacenamiento desde un cluster de zona.

- 1 **Conviértase en superusuario en un nodo del cluster global que aloje el cluster de zona.** Algunos pasos de este procedimiento se realizan desde un nodo del cluster global. Otros pasos se efectúan desde un nodo del cluster 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 cluster 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 cluster de zona.

```
phys-schost# clzonecluster configure zoneclustername
clzc:zoneclustername> remove device match=<devices_match>
clzc:zoneclustername> commit
clzc:zoneclustername> end
```

5 Rearranque el cluster 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 cluster de zona

En este ejemplo, se muestra cómo eliminar un conjunto de discos de SVM denominado `apachedg` y configurado en un cluster de zona denominado `sczone`. El número establecido del conjunto de discos `apachedg` es 3. El recurso `zc_rs` configurado en el cluster 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 cluster de zona

En este ejemplo se muestra cómo eliminar dispositivos de DID `d10` y `d11`, configurados en un cluster de zona denominado `sczone`. El recurso `zc_rs` configurado en el cluster 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 cluster global

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

Siga este procedimiento para ejecutar una aplicación fuera del cluster 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 cluster.

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

- b. Quite el dispositivo de quórum anterior.

```
phys-schost# clqorum 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.

El siguiente comando busca las claves de emulación de reserva de grupo persistente (PGRE). Si no hay claves en el disco, aparece el mensaje *errno=22*.

```
# /usr/cluster/lib/sc/pgre -c pgre_inkeys -d /dev/did/rdsk/dids2
```

Una vez que encuentra las claves, elimine las claves de PGRE.

```
# /usr/cluster/lib/sc/pgre -c pgre_scrub -d /dev/did/rdisk/dids2
```



Precaución – Si elimina las claves del dispositivo de quórum activo del disco, en la próxima reconfiguración, el cluster generará un aviso grave con el mensaje *Lost operational quorum* (Quórum operativo perdido).

- 2 **Evacue el nodo del cluster global que desee iniciar en un modo que no sea el de cluster.**

```
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 cluster 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 cluster.**

```
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 -eM resourcegroupname`

Restauración de 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 cluster 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.

▼ Guardado de la configuración del 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 [“Recreación de la configuración del software de Solaris Volume Manager” en la página 279](#).

- 1 Guarde la tabla de particiones para cada disco del conjunto de discos.
`# /usr/sbin/prtvto /dev/global/rdisk/diskname > /etc/lvm/diskname.vtoc`
- 2 Guarde la configuración de software de Solaris Volume Manager.
`# /bin/cp /etc/lvm/md.tab /etc/lvm/md.tab_ORIGINAL`
`# /usr/sbin/metastat -p -s setname >> /etc/lvm/md.tab`

Nota – Otros archivos de configuración, como el archivo `/etc/vfstab`, pueden hacer referencia al software 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 `prtvtoc` 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 un conjunto de discos de varios propietarios se depura mientras los nodos se inician fuera del modo de cluster, es posible que tenga que eliminar la información de los archivos de configuración DCS.

```
# /usr/cluster/lib/sc/dcs_config -c remove -s setname
```

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

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 surge un problema, elimine la información de los archivos de configuración DCS.

```
# /usr/cluster/lib/sc/dcs_config -c remove -s setname
```

Si las opciones de depuración fallan, verifique si tiene instaladas las últimas actualizaciones para el núcleo y el metadispositivo, y póngase en contacto con [My Oracle Support](#).

▼ **Recreación de la configuración del 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 clusters 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 cluster.

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/metadevice
```

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 cluster 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 281
- “Configuración del control de la CPU” en la página 283

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 zonas, proyectos, agrupaciones de recursos, conjuntos de procesadores y clases de programación, consulte la *Guía de administración de sistemas: administración de recursos y contenedores de Oracle Solaris y zonas de Oracle Solaris*.

En el sistema operativo Oracle Solaris se pueden realizar las tareas siguientes:

- Asignar recursos compartidos de CPU a grupos de recursos.
- Asignar procesadores a grupos de recursos.

Elección de una situación hipotética

Según las opciones de configuración y la versión del sistema operativo que elija, puede tener niveles distintos de control de la CPU. Todos los aspectos de control de la CPU que se describen en este capítulo dependen de la propiedad del grupo de recursos `RG_SLM_TYPE` establecida en `automated`.

La [Tabla 10–1](#) describe las diferentes situaciones hipotéticas de configuración posibles.

TABLA 10-1 Situaciones de control de la CPU

Descripción	Instrucciones
El grupo de recursos se ejecuta en el nodo de votación del cluster global. Asigne recursos compartidos de CPU a zonas y grupos de recursos proporcionando valores para <code>project.cpu-shares</code> y <code>zone.cpu-shares</code>. Puede realizar este procedimiento sin importar si hay nodos sin voto del cluster global configurados.	“Control del uso de la CPU en el nodo de votación de un cluster global” en la página 283
El grupo de recursos se ejecuta en una zona sin voto del cluster global utilizando el conjunto de procesadores predeterminado. Asigne recursos compartidos de CPU a zonas y grupos de recursos proporcionando valores para <code>project.cpu-shares</code> y <code>zone.cpu-shares</code>. Realice este procedimiento si no es necesario controlar el tamaño del conjunto de procesadores.	“Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores predeterminado” en la página 285
El grupo de recursos se ejecuta en un nodo sin voto del cluster global con un conjunto de procesadores dedicado. Asigne recursos compartidos de CPU a grupos de recursos proporcionando valores para <code>project.cpu-shares</code>, <code>zone.cpu-shares</code> y el número máximo de procesadores en un conjunto de procesadores dedicado. Defina el número mínimo de conjuntos de procesadores en un conjunto de procesadores dedicado. Realice este procedimiento si desea controlar los recursos compartidos de la CPU y el tamaño de un conjunto de procesadores. Puede ejercer este control sólo en un nodo sin voto del cluster global usando un conjunto de procesadores dedicado.	“Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores dedicado” en la página 288

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 cluster global” en la página 283
- “Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores predeterminado” en la página 285
- “Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores dedicado” en la página 288

▼ Control del uso de la CPU en el nodo de votación de un cluster 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 cluster 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 cluster 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.

```
# disadmin -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 `dispadm`, se asegura de que FSS se convierta en el programador predeterminado inmediatamente y de que permanezca después del rearranque. Para obtener más información sobre cómo configurar una clase de programación, consulte las páginas del comando `man dispadm(1M)` y `priocntl(1)`.

Nota – Si FSS no es el programador predeterminado, la asignación de recursos compartidos de CPU no se lleva a cabo.

2 En todos los nodos que usen el control de la CPU, configure el número de recursos compartidos para los nodos de votación del cluster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.

El establecimiento de estos parámetros ayuda a evitar que los procesos que se ejecutan en los nodos con voto compitan por las CPU con los procesos que se ejecutan en los nodos sin voto. Si no asigna un valor a las propiedades `globalzoneshares` y `defaulttpsetmin`, éstas asumen sus valores predeterminados.

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

`-p defaulttpsetmin= defaulttpsetmininteger` Establece el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado. El valor predeterminado es 1.

`-p globalzoneshares= integer` Configura el número de recursos compartidos asignados al nodo de votación. El valor predeterminado es 1.

`node` Especifica los nodos donde van a configurarse las propiedades.

Al configurar estas propiedades, también se establecen las del nodo de votación. Si no establece estas propiedades, no puede beneficiarse de la propiedad `RG_SLM_PSET_TYPE` en los nodos sin voto.

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
```

<code>-p RG_SLM_TYPE=automated</code>	Permite controlar el uso de la CPU y automatiza algunos pasos para configurar el sistema operativo Oracle Solaris con el fin de administrar recursos del sistema.
<code>-p RG_SLM_CPU_SHARES= valor</code>	Especifica el número de recursos compartidos de CPU asignados al proyecto específico del grupo de recursos, <code>project.cpu-shares</code> , y determina el número de recursos compartidos CPU asignados al nodo de votación <code>zone.cpu-shares</code> .
<code>nombre_grupo_recursos</code>	Especifica el nombre del grupo de recursos.

En este procedimiento, no se configura la propiedad `RG_SLM_PSET_TYPE`. En el nodo de votación, esta propiedad asume el valor `default`.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos.

5 Active el cambio de configuración.

```
# clresourcegroup online -eM resource_group_name
```

`resource_group_name` 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)`.

▼ Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores predeterminado

Realice este procedimiento si desea asignar recursos compartidos de la CPU para grupos de recursos en un nodo sin voto del cluster global, pero no necesita crear un conjunto de procesadores dedicado.

Si se asignan recursos compartidos de la CPU a un grupo de recursos, el software de Oracle Solaris Cluster realiza las siguientes tareas al iniciar un recurso de ese grupo de recursos en un nodo sin voto:

- Crea una agrupación denominada `SCSLM_resource_group_name` si aún no se ha realizado.
- Asocia la agrupación `SCSLM_pool_zone_name` al conjunto de procesadores predeterminado.

- Enlaza dinámicamente el nodo sin voto a la agrupación `SCSLM_poolzone_name`.
- Aumenta el número de recursos compartidos de CPU asignados al nodo sin voto (`zone.cpu-shares`) con el número especificado de recursos compartidos de CPU si todavía no se ha hecho.
- Crea un proyecto denominado `SCSLM_resourcegroup_name` en el nodo sin voto si todavía no se ha hecho. Este proyecto es específico del grupo de recursos y se le asigna el número especificado de recursos compartidos de CPU (`project.cpu-shares`).
- 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 la programación predeterminada inmediatamente y de que permanezca después del reinicio. Para obtener más información sobre cómo configurar una clase de programación, consulte las páginas del 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, para usar el control de la CPU, configure el número de recursos compartidos para el nodo con voto del cluster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.

El establecimiento de estos parámetros ayuda a evitar que los procesos que se ejecutan en el nodo con voto compitan por las CPU con los procesos que se ejecutan en los nodos sin voto del cluster global. Si no asigna un valor a las propiedades `globalzoneshares` y `defaultpsetmin`, éstas asumen sus valores predeterminados.

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

```
-p globalzoneshares=integer
```

Configura el número de recursos compartidos asignados al nodo de votación. El valor predeterminado es 1.

<code>-p defaultpsetmin=defaultpsetmininteger</code>	Establece el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado. El valor predeterminado es 1.
<code>node</code>	Identifica 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=value` Especifica el número de recursos compartidos de CPU asignados al proyecto específico del grupo de recursos (`project.cpu-shares`) y determina el número de recursos compartidos de CPU asignados al nodo sin voto del cluster global (`zone.cpu-shares`).

`resource_group_name` Especifica el nombre del grupo de recursos.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos.

No puede establecer `RG_SLM_TYPE` en `automated` en un nodo sin voto si una agrupación que no sea la agrupación predeterminada está en la configuración de la zona o si la zona está enlazada dinámicamente a una agrupación que no sea la agrupación predeterminada. Consulte las páginas del comando `man zonecfg(1M)` y `poolbind(1M)` para obtener información sobre la configuración de zonas y el enlace de agrupaciones respectivamente. Vea la configuración de la zona de la siguiente forma:

```
# zonecfg -z zone_name info pool
```

Nota – Un recurso, como, por ejemplo, `HASStoragePlus` o `LogicalHostname`, configurado para iniciarse en un nodo sin voto, pero con la propiedad `GLOBAL_ZONE` establecida en `TRUE`, se inicia en el nodo con voto. Aunque establezca la propiedad `RG_SLM_TYPE` en `automated`, este recurso no se beneficia de la configuración de los recursos compartidos de la CPU y se lo considera como si estuviera en un grupo de recursos con `RG_SLM_TYPE` establecida en `manual`.

En este procedimiento, no se configura la propiedad `RG_SLM_PSET_TYPE`. Oracle Solaris Cluster usa el conjunto de procesadores predeterminado.

5 Active el cambio de configuración.

```
# clresourcegroup online -eM resource_group_name
```

resource_group_name Especifica el nombre del grupo de recursos.

Si establece `RG_SLM_PSET_TYPE` en `default`, Oracle Solaris Cluster crea una agrupación, `SCSLM_pool_zone_name`, pero no crea un conjunto de procesadores. En este caso, `SCSLM_pool_zone_name` está asociado al conjunto de procesadores predeterminado.

Si los grupos de recursos en línea ya no están configurados para el control de la CPU en un nodo sin voto, el valor de recursos compartidos de la CPU para el nodo sin voto toma el valor de `zone.cpu-shares` en la configuración de la zona. Este parámetro tiene un valor de 1, de manera predeterminada. Para obtener más información sobre la configuración de zonas, consulte la página del comando `man zonecfg(1M)`.

Nota – No elimine ni modifique el proyecto `SCSLM_nombre_grupo_recursos`. Puede agregar más control de recursos de forma manual al proyecto, por ejemplo, configurando la propiedad `project.max-lwps`. Para obtener más información, consulte la página del comando `man projmod(1M)`.

▼ Control del uso de la CPU en un nodo sin voto del cluster global con el conjunto de procesadores dedicado

Realice este procedimiento si desea que el grupo de recursos se ejecute en un conjunto de procesadores dedicado.

Si un grupo de recursos se configura para ejecutarse en un conjunto de procesadores dedicado, el software de Oracle Solaris Cluster realiza las siguientes tareas cuando inicia un recurso del grupo de recursos en un nodo sin voto del cluster global:

- Crea una agrupación denominada `SCSLM_pool_zone_name` si aún no se ha realizado.

- Crea un conjunto de procesadores dedicado. El tamaño del conjunto de procesadores se determina mediante el uso de las propiedades `RG_SLM_CPU_SHARES` y `RG_SLM_PSET_MIN`.
- Asocia la agrupación `SCSLM_pool_zone_name` al conjunto de procesadores creado.
- Enlaza dinámicamente el nodo sin voto a la agrupación `SCSLM_pool_zone_name`.
- Aumenta el número de recursos compartidos de CPU asignados al nodo sin voto con el número especificado de recursos compartidos de CPU si todavía no se ha hecho.
- Crea un proyecto denominado `SCSLM_resourcegroup_name` en el nodo sin voto si todavía no se ha hecho. Este proyecto es específico del grupo de recursos y se le asigna el número especificado de recursos compartidos de CPU (`project.cpu-shares`).
- Inicia el recurso en el proyecto `SCSLM_nombre_grupo_recursos`.

1 Configure el programador para el sistema como el programador de reparto justo (FSS).

```
# 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 la programación predeterminada inmediatamente y de que permanezca después del reinicio. Para obtener más información sobre cómo configurar una clase de programación, consulte las páginas del 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, para usar el control de la CPU, configure el número de recursos compartidos para el nodo con voto del cluster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.

El establecimiento de estos parámetros ayuda a evitar que los procesos que se ejecutan en el nodo con voto compitan por las CPU con los procesos que se ejecutan en los nodos sin voto. Si no asigna un valor a las propiedades `globalzoneshares` y `defaultpsetmin`, éstas asumen sus valores predeterminados.

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

```
-p defaultpsetmin=defaultpsetmininteger
```

Establece el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado. El valor predeterminado es 1.

<code>-p globalzoneshares= integer</code>	Configura el número de recursos compartidos asignados al nodo de votación. El valor predeterminado es 1.
<code>node</code>	Identifica 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] \
  -p -y RG_SLM_PSET_TYPE=value \
  [-p RG_SLM_PSET_MIN=value] resource_group_name
```

<code>-p RG_SLM_TYPE=automated</code>	Permite controlar el uso de la CPU y automatiza algunos pasos para configurar el sistema operativo Oracle Solaris con el fin de gestionar recursos del sistema.
<code>-p RG_SLM_CPU_SHARES=value</code>	Especifica el número de recursos compartidos de CPU asignados al proyecto específico del grupo de recursos (<code>project.cpu-shares</code>) y determina el número de recursos compartidos de la CPU asignados al nodo sin voto (<code>zone.cpu-shares</code>) y el número máximo de procesadores en un conjunto de procesadores.
<code>-p RG_SLM_PSET_TYPE= value</code>	Permite la creación de un conjunto de procesadores dedicado. Para tener un conjunto de procesadores dedicado, puede establecer esta propiedad en <code>strong</code> o <code>weak</code> . Los valores <code>strong</code> y <code>weak</code> se excluyen mutuamente. Es decir, no se pueden configurar grupos de recursos en la misma zona de forma que algunos sean <code>strong</code> y otros <code>weak</code> .
<code>-p RG_SLM_PSET_MIN= value</code>	Determina el número mínimo de procesadores en el conjunto de procesadores.
<code>nombre_grupo_recursos</code>	Especifica el nombre del grupo de recursos.

Este paso crea un grupo de recursos. Si lo desea, utilice el comando `clresourcegroup set` para modificar un grupo de recursos existente.

No puede establecer `RG_SLM_TYPE` en `automated` en un nodo sin voto si una agrupación que no sea la agrupación predeterminada está en la configuración de la zona o si la zona está enlazada dinámicamente a una agrupación que no sea la agrupación predeterminada. Consulte las páginas del comando `man zonecfg(1M)` y `poolbind(1M)` para obtener información sobre la configuración de zonas y el enlace de agrupaciones respectivamente. Vea la configuración de la zona de la siguiente forma:

```
# zonecfg -z zone_name info pool
```

Nota – Un recurso, como, por ejemplo, `HASStoragePlus` o `LogicalHostname`, configurado para iniciarse en un nodo sin voto, pero con la propiedad `GLOBAL_ZONE` establecida en `TRUE`, se inicia en el nodo con voto. Aunque establezca la propiedad `RG_SLM_TYPE` en `automated`, este recurso no se beneficia de la configuración de los recursos compartidos de la CPU y del conjunto de procesadores dedicado, y se lo considera como si estuviera en un grupo de recursos con `RG_SLM_TYPE` establecida en `manual`.

5 Active el cambio de configuración.

```
# clresourcegroup online -eM resource_group_name
```

resource_group_name Especifica el nombre del grupo de recursos.

Nota – No elimine ni modifique el proyecto `SCSLM_nombre_grupo_recursos`. Puede agregar más control de recursos de forma manual al proyecto, por ejemplo, configurando la propiedad `project.max-lwps`. Para obtener más información, consulte la página del comando `man projmod(1M)`.

Los cambios realizados en `RG_SLM_CPU_SHARES` y `RG_SLM_PSET_MIN` mientras el grupo de recursos está en línea se tienen en cuenta de forma dinámica. Sin embargo, si `RG_SLM_PSET_TYPE` se establece en `strong` y si no hay suficientes CPU disponibles para adaptarse al cambio, el cambio solicitado para `RG_SLM_PSET_MIN` no se aplica. En este caso, se muestra un mensaje de advertencia. En el próximo switchover, se pueden producir errores debido a CPU insuficientes si no hay suficientes CPU disponibles para confirmar los valores que ha configurado para `RG_SLM_PSET_MIN`.

Si un grupo de recursos en línea ya no está configurado para el control de la CPU en el nodo sin voto, el valor de recursos compartidos de CPU para el nodo sin voto toma el valor de `zone.cpu-shares`. Este parámetro tiene un valor de 1, de manera predeterminada.

Aplicación de parches de software y firmware de Oracle Solaris Cluster

Este capítulo describe los procedimientos para agregar y quitar parches en una configuración de Oracle Solaris Cluster. Estos procedimientos se incluyen en las secciones siguientes.

- “Descripción general de la aplicación de parches de Oracle Solaris Cluster” en la página 293
- “Aplicación de parches de software de Oracle Solaris Cluster” en la página 295

Descripción general de la aplicación de parches de Oracle Solaris Cluster

Debido a la naturaleza de un cluster, todos los nodos que forman parte del cluster deben tener el mismo nivel de parche para un correcto funcionamiento del cluster. En ocasiones, al aplicar parches en un nodo con un parche de Oracle Solaris Cluster, es posible que se deba eliminar temporalmente la pertenencia de un nodo al cluster o detener el cluster completo antes de instalar el parche. En esta sección, se describen esos pasos.

Antes de aplicar un parche de Oracle Solaris Cluster, consulte el archivo README del parche. Además, puede consultar los requisitos de actualización de los dispositivos de almacenamiento para determinar qué método de parches necesitan.

Nota – Para los parches de Oracle Solaris Cluster, siempre consulte el archivo README del parche y SunSolve para obtener instrucciones que sustituyan los procedimientos descritos en este capítulo.

La instalación de parches en todos los nodos del cluster puede describirse mediante uno de los siguientes escenarios:

Reiniciar parche (nodo)

Un nodo debe iniciarse en modo de un solo usuario con el comando `boot -sx` o `shutdown -g -y -i0`, antes de que pueda aplicarse el parche o el firmware y, luego, se debe reiniciar para

	unirse al cluster. Primero, debe poner el nodo en un estado “silencioso”. Para ello, cambie los grupos de recursos o de dispositivos del nodo al cual se aplicará el parche a otro miembro del cluster. Además, aplique el parche o el firmware a un nodo del cluster a la vez para evitar cerrar todo el cluster. El cluster en sí permanece disponible durante este tipo de aplicación de parches, aunque los nodos individuales no están disponibles temporalmente. Un nodo con parches puede volver a unirse a un cluster como un nodo miembro, aunque otros nodos no estén aún en el mismo nivel de parche.
Reiniciar parche (cluster)	El cluster se debe detener y se debe iniciar cada nodo en modo de un solo usuario, con el comando <code>boot -sx</code> o <code>shutdown -g -y -i0</code>, para aplicar el parche de software o firmware. A continuación, reinicie los nodos para volver a unirse al cluster. Para este tipo de parche, el cluster no está disponible durante la aplicación de parches.
Parche de no reinicio	Un nodo no tiene que estar en estado “silencioso” (puede seguir controlando grupos de recursos o dispositivos) ni debe reiniciarse al aplicar el parche. No obstante, se debe aplicar el parche a un nodo por vez y verificar que funciona antes de aplicar el parche a otro nodo.

Nota – Los protocolos subyacentes del cluster no cambian debido a un parche.

Use el comando `patchadd` para aplicar un parche al cluster y `patchrm` para eliminar un parche (si es posible).

Sugerencias de parches de Oracle Solaris Cluster

Las siguientes sugerencias lo ayudan a administrar los parches de Oracle Solaris Cluster más eficazmente:

- Siempre lea el archivo `README` del parche antes de aplicar el parche.
- Consulte los requisitos de actualización de los dispositivos de almacenamiento para determinar qué método de parches necesitan.
- Aplique todos los parches (necesarios y recomendados) antes de ejecutar el cluster en un entorno de producción.
- Compruebe los niveles de hardware del firmware e instale todas las actualizaciones de firmware que puedan ser necesarias.

- Todos los nodos que actúan como miembros del cluster deben tener los mismos parches.
- Mantenga actualizados los parches de los subsistemas del cluster. Entre otros, estos parches incluyen la administración de volúmenes, el firmware de dispositivos de almacenamiento y el transporte del cluster.
- Revise los informes de parches frecuentemente, por ejemplo, una vez cada trimestre, y aplique una configuración de Oracle Solaris Cluster mediante el conjunto de parches recomendado.
- Aplique parches seleccionados según las recomendaciones de Enterprise Services.
- Pruebe la conmutación por error tras haber aplicado actualizaciones de parches importantes. Prepárese para retirar el parche si disminuye el rendimiento del cluster o si no funciona correctamente.

Aplicación de parches de software de Oracle Solaris Cluster

TABLA 11-1 Mapa de tareas: aplicación de parches al cluster

Tarea	Instrucciones
Aplicar un parche de no reinicio de Oracle Solaris Cluster en un nodo a la vez sin detener el nodo	“Aplicación de un parche de no reinicio de Oracle Solaris Cluster” en la página 303
Aplicar un parche de reinicio de Oracle Solaris Cluster tras poner el miembro del cluster en modo sin cluster	“Aplicación de un parche de reinicio (nodo)” en la página 295 “Aplicación de un parche de reinicio (cluster)” en la página 300
Aplicar un parche en modo de un solo usuario a los nodos con zonas de conmutación por error	“Aplicación de parches en modo de un solo usuario a nodos con zonas de conmutación por error” en la página 305
Eliminar un parche de Oracle Solaris Cluster	“Cambio de un parche de Oracle Solaris Cluster” en la página 309

▼ Aplicación de un parche de reinicio (nodo)

Aplique el parche a un nodo del cluster a la vez para mantener el cluster en funcionamiento durante el proceso de aplicación de parches. Con este procedimiento, primero debe detener el nodo en el cluster e iniciarlo en modo de un solo usuario con el comando `boot -sx` o `shutdown -g -y -i0` antes de aplicar el parche.

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

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

- 1 **Antes de aplicar el parche, consulte el sitio web de productos Oracle Solaris Cluster para obtener instrucciones especiales previas o posteriores a la instalación.**
- 2 **Conviértase en superusuario o asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin` en el nodo al cual aplicará el parche.**

- 3 **Enumere los grupos de recursos y dispositivos en el nodo al cual se está aplicando el parche.**

```
# clresourcegroup status -Z all -n node[,...]
```

node El nombre del nodo del cluster global o el nodo del cluster de zona que reside en el nodo al que se está aplicando el parche.

```
# cldevicegroup status -n node
```

node El nombre del nodo del cluster global al que se está aplicando el parche.

Nota – Los grupos de dispositivos no están asociados a un cluster de zona.

- 4 **Cambie todos los grupos de recursos, recursos y grupos de dispositivos del nodo al que se aplicará el parche a otros miembros del cluster.**

```
# clnode evacuate -n node
```

evacuate Evacúa todos grupos de dispositivos y de recursos, incluidos todos los nodos sin voto del cluster global.

-n node Especifica el nodo desde el que está cambiando los grupos de recursos y de dispositivos.

- 5 **Cierre el nodo.**

```
# shutdown -g0 [-y]
[-i0]
```

- 6 **Inicie el nodo en modo de un solo usuario sin cluster.**

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

```
ok boot -sx
```

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

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

Press any key to continue

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

Aparece el menú de GRUB, que es similar al siguiente:

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

- 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.

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

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

- c. Agregue **-sx** al comando para especificar que el sistema se inicie en modo sin cluster.

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

```
grub edit> kernel /platform/i86pc/multiboot -sx
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
+-----+
| root (hd0,0,a)                                         |
| kernel /platform/i86pc/multiboot -sx                 |
+-----+
```

```
| module /platform/i86pc/boot_archive |
+-----+
Use the ^ and v keys to select which entry is highlighted.
Press 'b' to boot, 'e' to edit the selected command in the
boot sequence, 'c' for a command-line, 'o' to open a new line
after ('O' for before) the selected line, 'd' to remove the
selected line, or escape to go back to the main menu.-
```

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

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 cluster. Si, por el contrario, desea iniciar en modo sin cluster, siga estos pasos para volver a agregar la opción -sx al comando del parámetro de inicio del núcleo.

7 Aplique el parche de software o firmware.

```
# patchadd -M patch-dir patch-id
```

patch-dir Especifica la ubicación del directorio del parche.

patch-id Especifica el número de parche de un parche determinado.

Nota – Siempre consulte las instrucciones que se incluyen en el directorio del parche, que sustituyen los procedimientos descritos en este capítulo.

8 Compruebe que el parche se haya instalado correctamente.

```
# showrev -p | grep patch-id
```

9 Reinicie el nodo en el cluster.

```
# reboot
```

10 Verifique que el parche funcione y que el nodo y el cluster funcionen con normalidad.

11 Repita del Paso 2 al Paso 10 para todos los nodos del cluster restantes.

12 Cambie los grupos de recursos y de dispositivos, según sea necesario.

Después de reiniciar todos los nodos, el último nodo reiniciado no tendrá los grupos de recursos ni de dispositivos en línea.

```
# cldevicegroup switch -n node + | devicegroup ...
# clresourcegroup switch -n node[:zone][,...] + | resource-group ...
```

node El nombre del nodo al que va a cambiar los grupos de recursos y de dispositivos.

zona El nombre del nodo sin voto de cluster global (node) que puede controlar el grupo de recursos. Especifique la zona solamente si ha proporcionado un nodo sin voto al crear

el grupo de recursos.

```
# clresourcegroup switch -Z zoneclustername -n zcnode[...] + | resource-group ...
zoneclustername    El nombre del cluster de zona al que va a cambiar los grupos de recursos.
zcnode             El nombre del nodo del cluster de zona que puede controlar el grupo de
                   recursos.
```

Nota – Los grupos de dispositivos no están asociados a un cluster de zona.

13 Verifique si debe confirmar el software del parche con el comando `scversions`.

```
# /usr/cluster/bin/scversions
```

Verá uno de los siguientes resultados:

```
Upgrade commit is needed.
Upgrade commit is NOT needed. All versions match.
```

14 Si se necesita confirmación, confirme el software del parche.

```
# scversions -c
```

Nota – Al ejecutar `scversions`, se iniciarán una o más reconfiguraciones de CMM, según la situación.

Ejemplo 11–1 Aplicación de un parche de reinicio (nodo)

En el ejemplo siguiente, se muestra la aplicación de un parche de reinicio de Oracle Solaris Cluster a un nodo.

```
# clresourcegroup status -n rg1
...Resource Group      Resource
-----
rg1                     rs-2
rg1                     rs-3
...
# cldevicegroup status -n nodedg-schost-1
...
Device Group Name:      dg-schost-1
...
# clnode evacuate phys-schost-2
# shutdown -g0 -y -i0
...
```

Inicie el nodo en modo de un solo usuario sin cluster.

- SPARC: Escriba:

```

ok boot -sx
■ x86: Inicie el nodo en modo de un solo usuario sin cluster. Consulte los pasos de inicio en el
procedimiento siguiente.

# patchadd -M /var/tmp/patches 234567-05
...
# showrev -p | grep 234567-05

...
# reboot
...
# cldevicegroup switch -n phys-schost-1 dg-schost-1
# clresourcegroup switch -n phys-schost-1 schost-sa-1
# scversions
Upgrade commit is needed.
# scversions -c

```

Véase también Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 309.](#)

▼ Aplicación de un parche de reinicio (cluster)

Con este procedimiento, primero debe detener el cluster e iniciar cada nodo en modo de un solo usuario con el comando `boot -sx` o `shtudown -g -y -i0` antes de aplicar el parche.

- 1 **Antes de aplicar el parche, consulte el sitio web de productos Oracle Solaris Cluster para obtener instrucciones especiales previas o posteriores a la instalación.**

- 2 **Conviértase en superusuario en un nodo de cluster.**

- 3 **Cierre el cluster.**

```
# cluster shutdown -y -g grace-period "message"
```

<code>-y</code>	Especifica que se responda <i>yes</i> al indicador de confirmación.
<code>-g <i>grace-period</i></code>	Especifica, en segundos, el período que se debe esperar antes de apagar el equipo. El período de gracia predeterminado es de 60 segundos.
<code><i>message</i></code>	Especifica el mensaje de advertencia que se transmitirá. Utilice comillas si <i>message</i> contiene varias palabras.

- 4 **Inicie cada nodo en modo de un solo usuario sin cluster.**

En la consola de cada nodo, ejecute los comandos siguientes.

- En los sistemas basados en SPARC, ejecute el comando siguiente.
- ```
ok boot -sx
```

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

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

Press any key to continue

- 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.**

Aparece el menú de GRUB, que es similar al siguiente:

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

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

Use the ^ and v keys to select which entry is highlighted.

Press enter to boot the selected OS, 'e' to edit the commands before booting, or 'c' for a command-line.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

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

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot |
| module /platform/i86pc/boot_archive |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.

Press 'b' to boot, 'e' to edit the selected command in the boot sequence, 'c' for a command-line, 'o' to open a new line after ('O' for before) the selected line, 'd' to remove the selected line, or escape to go back to the main menu.

- Agregue -sx al comando para especificar que el sistema se inicie en modo sin cluster.**

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

```
grub edit> kernel /platform/i86pc/multiboot -sx
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
```

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

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

---

**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 cluster. Si, por el contrario, desea iniciar en modo sin cluster, siga estos pasos para volver a agregar la opción -sx al comando del parámetro de inicio del núcleo.

---

**5 Aplique el parche de software o firmware.**

Ejecute el comando siguiente, nodo por nodo.

```
patchadd -M patch-dir patch-id
```

*patch-dir*                      Especifica la ubicación del directorio del parche.

*patch-id*                      Especifica el número de parche de un parche determinado.

---

**Nota** – Siempre consulte las instrucciones que se incluyen en el directorio del parche, que sustituyen los procedimientos descritos en este capítulo.

---

**6 Compruebe que el parche se haya instalado correctamente en todos los nodos.**

```
showrev -p | grep patch-id
```

**7 Después de aplicar el parche en todos los nodos, reinicie los nodos en el cluster.**

Ejecute el comando siguiente en todos los nodos.

```
reboot
```

**8 Verifique si debe confirmar el software del parche con el comando scversions.**

```
/usr/cluster/bin/scversions
```

Verá uno de los siguientes resultados:

Upgrade commit is needed.

Upgrade commit is NOT needed. All versions match.

**9 Si se necesita confirmación, confirme el software del parche.**

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, se iniciarán una o más reconfiguraciones de CMM, según la situación.

---

**10 Verifique que el parche funcione y que los nodos y el cluster funcionen con normalidad.**

**Ejemplo 11-2 Aplicación de un parche de reinicio (cluster)**

En el ejemplo siguiente, se muestra la aplicación de un parche de reinicio de Oracle Solaris Cluster a un cluster.

```
cluster shutdown -g0 -y
...
```

Inicie el cluster en modo de un solo usuario sin cluster.

- SPARC: Escriba:
 

```
ok boot -sx
```
- x86: Inicie cada nodo en modo de un solo usuario sin cluster. Consulte el procedimiento siguiente para ver los pasos.

```
...
patchadd -M /var/tmp/patches 234567-05
(Apply patch to other cluster nodes)
...
showrev -p | grep 234567-05
reboot
scversions
Upgrade commit is needed.
scversions -c
```

**Véase también** Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 309.](#)

## ▼ Aplicación de un parche de no reinicio de Oracle Solaris Cluster

Aplique el parche en un nodo del cluster a la vez. Al aplicar un parche de no reinicio, no es necesario detener primero el nodo que está recibiendo el parche.

- 1 Antes de aplicar el parche, consulte la página web de productos Oracle Solaris Cluster para obtener instrucciones especiales previas o posteriores a la instalación.**

**2 Aplique el parche en un solo nodo.**

```
patchadd -M patch-dir patch-id
```

*patch-dir*                      Especifica la ubicación del directorio del parche.

*patch-id*                      Especifica el número de parche de un parche determinado.

**3 Compruebe que el parche se haya instalado correctamente.**

```
showrev -p | grep patch-id
```

**4 Verifique que el parche funcione y que el nodo y el cluster funcionen con normalidad.**

**5 Repita del [Paso 2](#) al [Paso 4](#) para todos los nodos del cluster restantes.**

**6 Verifique si debe confirmar el software del parche con el comando `scversions`.**

```
/usr/cluster/bin/scversions
```

Verá uno de los siguientes resultados:

```
Upgrade commit is needed.
```

```
Upgrade commit is NOT needed. All versions match.
```

**7 Si se necesita confirmación, confirme el software del parche.**

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, se iniciarán una o más reconfiguraciones de CMM, según la situación.

---

**Ejemplo 11–3**    Aplicación de un parche de no reinicio de Oracle Solaris Cluster

```
patchadd -M /tmp/patches 234567-05
```

```
...
```

```
showrev -p | grep 234567-05
```

```
scversions
```

```
Upgrade commit is needed.
```

```
scversions -c
```

**Véase también**    Si necesita anular un parche, consulte [“Cambio de un parche de Oracle Solaris Cluster” en la página 309.](#)

## ▼ Aplicación de parches en modo de un solo usuario a nodos con zonas de conmutación por error

Realice esta tarea para aplicar parches en modo de un solo usuario con zonas de conmutación por error. Este método de parche es necesario si utiliza el servicio de datos de Oracle Solaris Cluster para Solaris Containers en una configuración de conmutación por error con software de Oracle Solaris Cluster.

- 1 Verifique que el dispositivo de quórum no esté configurado para uno de los LUN utilizados como almacenamiento compartido que pertenece a los conjuntos de discos que contienen la ruta de zona adoptada manualmente en este procedimiento.

- a. Determine si el dispositivo de quórum se utiliza en los conjuntos de discos que contienen las rutas de zonas y si usa reservas de SCSI2 y SCSI3.

```
clquorum show
```

- b. Si el dispositivo de quórum está dentro de un LUN de los conjuntos de discos, agregue un nuevo LUN como un dispositivo de quórum que no pertenece a ningún conjunto de discos que contiene la ruta de zona.

```
clquorum add new-didname
```

- c. Quite el dispositivo de quórum anterior.

```
clquorum remove old-didname
```

- d. Si se usan reservas de SCSI2 para el dispositivo de quórum anterior, elimine las reservas de SCSI2 del quórum anterior y verifique que no queden reservas SCSI2.

El siguiente comando busca las claves de emulación de reserva de grupo persistente (PGRE). Si no hay claves en el disco, aparece el mensaje *errno=22*.

```
/usr/cluster/lib/sc/pgre -c pgre_inkeys -d /dev/did/rdisk/dids2
```

Una vez que encuentra las claves, elimine las claves de PGRE.

```
/usr/cluster/lib/sc/pgre -c pgre_scrub -d /dev/did/rdisk/dids2
```



**Precaución** – Si elimina las claves del dispositivo de quórum activo del disco, en la próxima reconfiguración, el cluster generará un aviso grave con el mensaje *Lost operational quorum* (Quórum operativo perdido).

- 2 Evacúe el nodo al que desea aplicar el parche.

```
clresourcegroup evacuate -n node1
```

- 3 **Desconecte los grupos de recursos que contengan recursos de Solaris Container de alta disponibilidad.**

```
clresourcegroup offline resourcegroupname
```

- 4 **Deshabilite todos los recursos de los grupos de recursos que desconectó.**

```
clresource disable resourcename
```

- 5 **Anule la administración de los grupos de recursos que desconectó.**

```
clresourcegroup unmanage resourcegroupname
```

- 6 **Desconecte el grupo o los grupos de dispositivos correspondientes.**

```
cldevicegroup offline cldevicegroupname
```

---

**Nota** – Si está aplicando un parche en una zona de conmutación por error que tiene zpools para la ruta de zona, omita este paso y el [Paso 7](#).

---

- 7 **Desactive los grupos de dispositivos que desconectó.**

```
cldevicegroup disable devicegroupname
```

- 8 **Inicie el nodo pasivo fuera del cluster.**

```
reboot -- -x
```

---

**Nota** – Utilice el siguiente comando si está aplicando un parche a una zona de conmutación por error que tiene zpools para la ruta de zona.

---

```
reboot -- -xs
```

- 9 **Verifique que los métodos de inicio de SMF se hayan completado en el nodo pasivo antes de continuar.**

```
svcs -x
```

---

**Nota** – Si está aplicando un parche en una zona de conmutación por error que tiene zpools para la ruta de zona, omita este paso.

---

- 10 **Verifique que se hayan completado los procesos de reconfiguración en el nodo activo.**

```
cluster status
```

- 11 **Determinar si existen reservas de SCSI-2 en el disco del conjunto y libere las claves. Siga estas instrucciones para determinar si existen reservas de SCSI-2 y, a continuación, libere las claves.**

- Para todos los discos del conjunto, ejecute el siguiente comando:  

```
/usr/cluster/lib/sc/scsi -c disfailfast -d /dev/did/rdisk/d#s2.
```

- Si se enumeran claves, libérelas con el siguiente comando: `/usr/cluster/lib/sc/scsi -c release -d /dev/did/rdisk/d#s2`.

Cuando termine de liberar las claves de reserva, omita el Paso 12 y continúe con el Paso 13.

## 12 Determine si existen reservas de SCSI-3 en los discos del conjunto.

- a. Ejecute el siguiente comando en todos los discos de los conjuntos.

```
/usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/didnames2
```

- b. Si se enumeran claves, elimínelas.

```
/usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/didnames2
```

## 13 Conviértase en propietario del metaconjunto en el nodo pasivo

```
metaset -s disksetname -C take -f
```

---

**Nota** – Utilice el siguiente comando si está aplicando un parche a una zona de conmutación por error que tiene zpools para la ruta de zona.

```
zpool import -R / pool_name
```

---

## 14 Monte los sistemas de archivos que contienen la ruta de zona en el nodo pasivo.

```
mount device mountpoint
```

---

**Nota** – Si está aplicando un parche en una zona de conmutación por error que tiene zpools para la ruta de zona, omita este paso y el [Paso 15](#).

---

## 15 Cambie al modo de un solo usuario en el nodo pasivo.

```
init s
```

## 16 Detenga todas las zonas iniciadas posibles que no estén bajo el servicio de datos de Oracle Solaris Cluster para el control de Solaris Container.

```
zoneadm -z zonename halt
```

## 17 (Opcional) Si instala varios parches, por motivos de rendimiento, puede elegir iniciar todas las zonas configuradas en modo de un solo usuario.

```
zoneadm -z zonename boot -s
```

## 18 Aplique los parches.

- 19 Reinicie el nodo y espere hasta que se hayan completado todos los métodos de inicio de SMF. Ejecute el comando `svcs -a` después de reiniciar el nodo.  

```
reboot
```

```
svcs -a
```

El primer nodo ahora está listo.
- 20 Evacúe el segundo nodo al que desea aplicar el parche.  

```
clresourcegroup evacuate -n node2
```
- 21 Repita los pasos 8 a 13 para el segundo nodo.
- 22 Desconecte las zonas a las cuales aplicó los parches. Si no desconecta estas zonas, el proceso de aplicación de parches fallará.  

```
zoneadm -z zonename detach
```
- 23 Cambie al modo de un solo usuario en el nodo pasivo.  

```
init s
```
- 24 Detenga todas las zonas iniciadas posibles que no estén bajo el servicio de datos de Oracle Solaris Cluster para el control de Solaris Containers.  

```
zoneadm -z zonename halt
```
- 25 (Opcional) Si instala varios parches, por motivos de rendimiento, puede elegir iniciar todas las zonas configuradas en modo de un solo usuario.  

```
zoneadm -z zonename boot -s
```
- 26 Aplique los parches.
- 27 Conecte las zonas que desconectó.  

```
zoneadm -z zonename attach -F
```
- 28 Reinicie el nodo en modo de cluster.  

```
reboot
```
- 29 Ponga en línea el grupo o los grupos de dispositivos.
- 30 Inicie los grupos de recursos.
- 31 Verifique si debe confirmar el software del parche con el comando `scversions`.  

```
/usr/cluster/bin/scversions
```

Verá uno de los siguientes resultados:

Upgrade commit is needed.

Upgrade commit is NOT needed. All versions match.

### 32 Si se necesita confirmación, confirme el software del parche.

```
scversions -c
```

---

**Nota** – Al ejecutar `scversions`, se iniciarán una o más reconfiguraciones de CMM, según la situación.

---

## Cambio de un parche de Oracle Solaris Cluster

Para eliminar un parche de Oracle Solaris Cluster aplicado al cluster, primero debe eliminar el nuevo parche de Oracle Solaris Cluster y, a continuación, volver a aplicar el parche anterior o versión de actualización. Para eliminar el nuevo parche de Oracle Solaris Cluster, consulte los siguientes procedimientos. Para volver a aplicar el parche de Oracle Solaris Cluster anterior, consulte uno de los siguientes procedimientos:

- [“Aplicación de un parche de reinicio \(nodo\)” en la página 295](#)
- [“Aplicación de un parche de reinicio \(cluster\)” en la página 300](#)
- [“Aplicación de un parche de no reinicio de Oracle Solaris Cluster” en la página 303](#)

---

**Nota** – Antes de aplicar un parche de Oracle Solaris Cluster, consulte el archivo README del parche.

---

### ▼ Eliminación de un parche de no reinicio de Oracle Solaris Cluster

- 1 Conviértase en superusuario en un nodo de cluster.

- 2 Elimine el parche de no reinicio.

```
patchrm patchid
```

### ▼ Eliminación de un parche de reinicio de Oracle Solaris Cluster

- 1 Conviértase en superusuario en un nodo de cluster.

- 2 Inicie el nodo del cluster en modo sin cluster. Para obtener información sobre el inicio de un nodo en modo sin cluster, consulte [“Rearranque de un nodo en un modo que no sea de cluster” en la página 82.](#)

- 3 Elimine el parche de reinicio.**  
`# patchrm patchid`
- 4 Reinicie el nodo del cluster en modo de cluster.**  
`# reboot`
- 5 Repita los pasos 2 a 4 para cada nodo del cluster.**

# Copias de seguridad y restauraciones de clusters

Este capítulo contiene las secciones siguientes.

- “Copia de seguridad de un cluster” en la página 311
- “Restauración de los archivos del cluster” en la página 319

## Copia de seguridad de un cluster

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

| Tarea                                                                                           | Instrucciones                                                                                        |
|-------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|
| Buscar los nombres de los sistemas de archivos de los que desea realizar una copia de seguridad | “Búsqueda de nombres de sistemas de archivos para realizar copias de seguridad” en la página 312     |
| Calcular cuántas cintas necesita para una copia de seguridad completa                           | “Determinación del número de cintas necesario para una copia de seguridad completa” en la página 312 |
| Realizar una copia de seguridad del sistema de archivos raíz                                    | “Realización de una copia de seguridad del sistema de archivos raíz (/)” en la página 313            |
| Realizar una copia de seguridad en línea para sistemas de archivos reflejados                   | “Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)” en la página 316          |
| Hacer una copia de seguridad de la configuración del cluster                                    | “Copias de seguridad de la configuración del cluster” en la página 319                               |
| Realizar una copia de seguridad de la configuración de partición del disco de almacenamiento    | Consulte la documentación del disco de almacenamiento                                                |

## ▼ Búsqueda de nombres de sistemas de archivos para realizar copias de seguridad

Utilice este procedimiento para determinar los nombres de los sistemas de archivos de los que desea realizar una copia de seguridad.

- 1 **Visualice el contenido del archivo `/etc/vfstab`.**  
No es necesario ser superusuario ni asumir un rol equivalente para ejecutar este comando.  
`# more /etc/vfstab`
- 2 **Busque en la columna de punto de montaje el nombre del sistema de archivos del que realizará la copia de seguridad.**  
Use este nombre cuando realice la copia de seguridad del sistema de archivos.  
`# more /etc/vfstab`

### Ejemplo 12-1 Búsqueda de nombres de sistemas de archivos para realizar copias de seguridad

En el ejemplo siguiente, se muestran los nombres de los sistemas de archivos disponibles que aparecen en el archivo `/etc/vfstab`.

```
more /etc/vfstab
#device device mount FS fsck mount mount
#to mount to fsck point type pass at boot options
#
#/dev/dsk/cl1d0s2 /dev/rdisk/cl1d0s2 /usr ufs 1 yes -
f - /dev/fd fd - no -
/proc - /proc proc - no -
/dev/dsk/cl1t6d0s1 - - swap - no -
/dev/dsk/cl1t6d0s0 /dev/rdisk/cl1t6d0s0 / ufs 1 no -
/dev/dsk/cl1t6d0s3 /dev/rdisk/cl1t6d0s3 /cache ufs 2 yes -
swap - /tmp tmpfs - yes -
```

## ▼ Determinación del número de cintas necesario para una copia de seguridad completa

Utilice este procedimiento para calcular el número de cintas que necesita para realizar la copia de seguridad de un sistema de archivos.

- 1 **Conviértase en superusuario o asuma una función equivalente en el nodo del cluster del que esté haciendo la copia de seguridad.**
- 2 **Calcule el tamaño de la copia de seguridad en bytes.**  
`# ufsdump S filesystem`

|                         |                                                                                             |
|-------------------------|---------------------------------------------------------------------------------------------|
| <code>S</code>          | Muestra el número estimado de bytes necesarios para realizar la copia de seguridad.         |
| <code>filesystem</code> | Especifica el nombre del sistema de archivos del que desea realizar una copia de seguridad. |

- 3 **Divida el tamaño estimado entre la capacidad de la cinta para saber cuántas cintas necesita.**

### Ejemplo 12-2 Determinación del número de cintas necesarias

En el ejemplo siguiente, el tamaño del sistema de archivos de 905.881.620 bytes entra fácilmente en una cinta de 4 GB ( $905.881.620 \div 4.000.000.000$ ).

```
ufsdump S /global/phys-schost-1
905881620
```

## ▼ Realización de una copia de seguridad del sistema de archivos raíz (/)

Utilice este procedimiento para realizar una copia de seguridad del sistema de archivos raíz (/) de un nodo del cluster. Asegúrese de que el cluster se esté ejecutando sin errores antes de llevar a cabo el procedimiento de copia de seguridad.

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

Este procedimiento proporciona las formas largas de los comandos de 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 RBAC `solaris.cluster.modify` en el nodo del cluster del que va a realizar la copia de seguridad.**
- 2 **Cambie cada servicio de datos en ejecución desde el nodo del que se va a realizar la copia de seguridad hasta otro nodo del cluster.**

```
clnode evacuate node
```

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

- 3 **Cierre el nodo.**  

```
shutdown -g0 -y -i0
```
- 4 **Reinicie el nodo en el modo sin cluster.**

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

```
ok boot -xs
```

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

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

Press any key to continue

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

Aparece el menú de GRUB, que es similar al siguiente:

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Cómo iniciar un sistema basado en x86 mediante GRUB \(mapa de tareas\)” de Administración de Oracle Solaris: administración básica.](#)

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

Aparece la pantalla de parámetros de inicio de GRUB, que es similar a la siguiente:

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

- Agregue -x al comando para especificar que el sistema se inicia en el modo sin cluster.**

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

```
grub edit> kernel /platform/i86pc/multiboot -x
```

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

La pantalla muestra el comando editado.

```
GNU GRUB version 0.95 (615K lower / 2095552K upper memory)
```

```
+-----+
| root (hd0,0,a) |
| kernel /platform/i86pc/multiboot -x |
| module /platform/i86pc/boot_archive |
+-----+
```

Use the ^ and v keys to select which entry is highlighted.  
Press 'b' to boot, 'e' to edit the selected command in the  
boot sequence, 'c' for a command-line, 'o' to open a new line  
after ('O' for before) the selected line, 'd' to remove the  
selected line, or escape to go back to the main menu.-

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

---

**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 cluster. Si, por el contrario, desea iniciar en modo sin cluster, siga estos pasos para volver a agregar la opción -x al comando del parámetro de inicio del núcleo.

---

## 5 Realice la copia de seguridad del sistema de archivos raíz (/) creando una instantánea de UFS.

- a. Asegúrese de que el sistema de archivos dispone de suficiente espacio en el disco para el archivo de copia de seguridad.

```
df -k
```

- b. Asegúrese de que no exista un archivo de copia de seguridad con el mismo nombre y en la misma ubicación.

```
ls /backing-store-file
```

- c. Cree la instantánea de UFS.

```
fssnap -F ufs -o bs=/backing-store-file /file-system
```

- d. Verifique que se haya creado la instantánea.

```
/usr/lib/fs/ufs/fssnap -i /file-system
```

## 6 Realice una copia de seguridad de la instantánea del sistema de archivos.

```
ufsdump 0ucf /dev/rmt/0 snapshot-name
```

Por ejemplo:

```
ufsdump 0ucf /dev/rmt/0 /dev/rfssnap/1
```

- 7 Verifique que se haya realizado la copia de seguridad de la instantánea.

```
ufsrestore ta /dev/rmt/0
```

- 8 Reinicie el nodo en modo de cluster.

```
init 6
```

### Ejemplo 12-3 Copia de seguridad del sistema de archivos raíz (/)

En el ejemplo siguiente, una instantánea del sistema de archivos raíz (/) se guarda en /scratch/usr.back.file, en el directorio /usr. ‘

```
fssnap -F ufs -o bs=/scratch/usr.back.file /usr
/dev/fssnap/1
```

## ▼ Copias de seguridad en línea para duplicaciones (Solaris Volume Manager)

Se puede realizar una copia de seguridad de un volumen de Solaris Volume Manager sin desmontarlo o poner la duplicación completa fuera de línea. Una de las subduplicaciones debe ponerse temporalmente fuera de línea; si bien se pierde el duplicado, puede ponerse en línea y volver a sincronizarse en cuanto la copia de seguridad esté terminada sin detener el sistema ni denegar el acceso del usuario a los datos. El uso de duplicaciones para realizar copias de seguridad en línea crea una copia de seguridad que es una "instantánea" de un sistema de archivos activo.

Si un programa escribe datos en el volumen justo antes de ejecutarse el comando `lockfs` puede causar problemas. Para evitarlo, detenga temporalmente todos los servicios que se estén ejecutando en este nodo. Asimismo, antes de efectuar la copia de seguridad, compruebe que el cluster se ejecute sin errores.

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster del que esté haciendo la copia de seguridad.
- 2 Utilice el comando `metaset(1M)` 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.

- 3 Utilice el comando **lockfs(1M)** con la opción -w para bloquear el sistema de archivos contra escritura.

```
lockfs -w mountpoint
```

---

**Nota** – Debe bloquear el sistema de archivos únicamente si un sistema de archivos UFS reside en el reflejo. Por ejemplo, si el volumen de Solaris Volume Manager está configurado como un dispositivo sin formato para software de administración de base de datos o alguna otra aplicación específica, no es necesario utilizar el comando `lockfs`. Sin embargo, es posible ejecutar la utilidad adecuada dependiente del proveedor para vaciar el búfer y bloquear el acceso.

---

- 4 Utilice el comando **metastat(1M)** para determinar los nombres de los subreflejos.

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

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

- 5 Utilice el comando **metadetach(1M)** para desconectar un subreflejo del reflejo.

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

---

**Nota** – Las lecturas se siguen efectuando desde otras subduplicaciones. Sin embargo, la subduplicación fuera de línea pierde la sincronización al realizarse la primera escritura en la duplicación. Esta incoherencia se corrige al poner la subduplicación en línea de nuevo. El comando `fsck` no debe ejecutarse.

---

- 6 Desbloquee los sistemas y permita que se pueda seguir escribiendo mediante el comando **lockfs** con la opción -u.

```
lockfs -u mountpoint
```

- 7 Realice una comprobación del sistema de archivos.

```
fsck /dev/md/diskset/rdisk/submirror
```

- 8 Haga una copia de seguridad de la subduplicación fuera de línea en una cinta o en otro medio.

Utilice el comando **ufsdump(1M)** o la utilidad de copia de seguridad que suele usar.

```
ufsdump 0ucf dump-device submirror
```

---

**Nota** – Use el nombre del dispositivo básico (`/rdsk`) para la subduplicación, en lugar del nombre del dispositivo de bloqueo (`/dsk`).

---

- 9 Utilice el comando **metattach(1M)** para volver a conectar el metadispositivo o el volumen.

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

Cuando se pone en línea un metadispositivo o un volumen, se vuelve a sincronizar automáticamente con la duplicación.

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

```
metastat -s setname mirror
```

## Ejemplo 12–4 Realización de copias de seguridad en línea para reflejos (Solaris Volume Manager)

En el ejemplo siguiente, el nodo del cluster `phys-schost-1` es el propietario del metaconjunto `schost-1`, por lo tanto, el procedimiento de copia de seguridad se realiza desde `phys-schost-1`. El reflejo `/dev/md/schost-1/dsk/d0` consta de los subreflejos `d10`, `d20` y `d30`.

```
[Determine the owner of the metaset:]
metaset -s schost-1
Set name = schost-1, Set number = 1
Host Owner
 phys-schost-1 Yes
...
[Lock the file system from writes:]
lockfs -w /global/schost-1
[List the submirrors:]
metastat -s schost-1 -p
schost-1/d0 -m schost-1/d10 schost-1/d20 schost-1/d30 1
schost-1/d10 1 1 d4s0
schost-1/d20 1 1 d6s0
schost-1/d30 1 1 d8s0
[Take a submirror offline:]
metadetach -s schost-1 d0 d30
[Unlock the file system:]
lockfs -u /
[Check the file system:]
fsck /dev/md/schost-1/rdisk/d30
[Copy the submirror to the backup device:]
ufsdump 0ucf /dev/rmt/0 /dev/md/schost-1/rdisk/d30
DUMP: Writing 63 Kilobyte records
DUMP: Date of this level 0 dump: Tue Apr 25 16:15:51 2000
DUMP: Date of last level 0 dump: the epoch
DUMP: Dumping /dev/md/schost-1/rdisk/d30 to /dev/rdisk/clt9d0s0.
...
DUMP: DUMP IS DONE
[Bring the submirror back online:]
metattach -s schost-1 d0 d30
schost-1/d0: submirror schost-1/d30 is attached
[Resynchronize the submirror:]
metastat -s schost-1 d0
schost-1/d0: Mirror
 Submirror 0: schost-0/d10
 State: Okay
 Submirror 1: schost-0/d20
 State: Okay
```

```

Submirror 2: schost-0/d30
State: Resyncing
Resync in progress: 42% done
Pass: 1
Read option: roundrobin (default)
...

```

## ▼ Copias de seguridad de la configuración del cluster

Para asegurarse de que se archive la configuración del cluster y para facilitar su recuperación, realice periódicamente copias de seguridad de la configuración del cluster. Oracle Solaris Cluster proporciona la capacidad de exportar la configuración del cluster a un archivo eXtensible Markup Language (XML).

- 1 **Inicie sesión en cualquier nodo del cluster 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 cluster a un archivo.**  

```
/usr/cluster/bin/cluster export -o configfile
```

*configfile* El nombre del archivo de configuración XML al que el comando del cluster va a exportar la información de configuración del cluster. Para obtener información sobre el archivo de configuración XML, consulte [clconfiguration\(5CL\)](#).
- 3 **Compruebe que la información de configuración del cluster se haya exportado correctamente al archivo XML.**  

```
vi configfile
```

## Restauración de los archivos del cluster

El comando [ufsrestore\(1M\)](#) copia archivos en el disco, relacionados con el directorio de trabajo actual, desde copias de seguridad que se crearon con el comando [ufsdump\(1M\)](#). Puede utilizar `ufs restore` para volver a cargar una jerarquía completa del sistema de archivos desde un volcado de nivel 0 y volcados incrementales que siguen, o para restaurar uno o más archivos individuales desde cualquier cinta de volcado. Si `ufs restore` se ejecuta como superusuario o con un rol equivalente, los archivos se restauran con su propietario original, hora de última modificación y modo (permisos).

Antes de empezar a restaurar los archivos o los sistemas de archivos, debe conocer la información siguiente:

- Las cintas que necesita
- El nombre del dispositivo básico en el que va a restaurar el sistema de archivos

- El tipo de unidad de cinta que va a utilizar
- El nombre del dispositivo (local o remoto) para la unidad de cinta
- El esquema de partición de cualquier disco donde haya habido un error, porque las particiones y los sistemas de archivos deben estar duplicados exactamente en el disco de reemplazo

TABLA 12-2 Mapa de tareas: restaurar los archivos del cluster

| Tarea                                                                     | Instrucciones                                                                                                                                                                                                                 |
|---------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Para Solaris Volume Manager, restaurar los archivos de manera interactiva | “Restauración interactiva de archivos individuales (Solaris Volume Manager)” en la página 320                                                                                                                                 |
| Para Solaris Volume Manager, restaurar el sistema de archivos raíz (/)    | “Restauración del sistema de archivos raíz (/) (Solaris Volume Manager)” en la página 320<br><br>“Restauración de un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager” en la página 323 |

## ▼ Restauración interactiva de archivos individuales (Solaris Volume Manager)

Utilice este procedimiento para restaurar uno o más archivos individuales. Asegúrese de que el cluster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

- 1 **Conviértase en superusuario o asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin` en el nodo del cluster que está restaurando.**
- 2 **Detenga todos los servicios de datos que usan los archivos que se van a restaurar.**  
`# clresourcegroup offline resource-group`
- 3 **Restaurar los archivos.**  
`# ufsrestore`

## ▼ Restauración del sistema de archivos raíz (/) (Solaris Volume Manager)

Utilice este procedimiento para restaurar los sistemas de archivos raíz (/) en un disco nuevo, por ejemplo, después de reemplazar un disco raíz defectuoso. El nodo que se está restaurando no se debe arrancar. Asegúrese de que el cluster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

---

**Nota** – Como debe particionar el disco nuevo con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

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

Este procedimiento proporciona las formas largas de los comandos de 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 cluster 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
```

|                             |                                                                            |
|-----------------------------|----------------------------------------------------------------------------|
| <code>-s setname</code>     | Especifica el nombre del conjunto de discos.                               |
| <code>-f</code>             | Elimina el último host del conjunto de discos.                             |
| <code>-d</code>             | Elimina del conjunto de discos.                                            |
| <code>-h lista_nodos</code> | Especifica el nombre del nodo que se va a eliminar del conjunto de discos. |

- 3 Restaure los sistemas de archivos de raíz (/) y /usr.**

Para restaurar los sistemas de archivos raíz y /usr, siga el procedimiento descrito en [“Restoring UFS Files and File Systems” de System Administration Guide: Devices and File Systems](#). Omita el paso de reinicio del sistema en el procedimiento del sistema operativo Oracle Solaris.

---

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

---

- 4 Rearranque el nodo en modo multiusuario.**

```
reboot
```

- 5 Sustituya el ID del dispositivo.**

```
cldevice repair rootdisk
```

**6 Utilice el comando `metadb(1M)` para volver a crear 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 *raw-disk-device* Dispositivo de disco básico en el que crear las réplicas.

-a Agrega réplicas.

**7 Desde un nodo del cluster que no sea el restaurado, agregue el nodo restaurado a todos los grupos de discos.**

```
phys-schost-2# metaset -s setname -a -h nodelist
```

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

El nodo se reanuda en modo de cluster. El cluster está preparado para utilizarse.

## **Ejemplo 12-5 Restauración del sistema de archivos raíz (/) (Solaris Volume Manager)**

En el ejemplo siguiente, se muestra el sistema de archivos raíz (/) restaurado en el nodo `phys-schost-1` desde el dispositivo de cinta `/dev/rmt/0`. El comando `metaset` se ejecuta desde otro nodo en el cluster, `phys-schost-2`, para eliminar y luego volver a agregar el nodo `phys-schost-1` al conjunto de discos `schost-1`. Todos los demás comandos se ejecutan desde `phys-schost-1`. Se crea un bloque de arranque nuevo en `/dev/rdsk/c0t0d0s0` y se vuelven a crear tres réplicas de la base de datos en `/dev/rdsk/c0t0d0s4`.

```
[Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on a cluster node
other than the node to be restored.]
[Remove the node from the metaset:]
phys-schost-2# metaset -s schost-1 -f -d -h phys-schost-1
[Replace the failed disk and boot the node:]
Restore the root (/) and /usr file system using the procedure in the Solaris system
administration documentation
[Reboot:]
reboot
[Replace the disk ID:]
cldevice repair /dev/dsk/c0t0d0
[Re-create state database replicas:]
metadb -c 3 -af /dev/rdsk/c0t0d0s4
[Add the node back to the metaset:]
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

## ▼ Restauración de un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager

Utilice este procedimiento para restaurar un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager al realizar las copias de seguridad. Realice este procedimiento cuando el disco raíz está dañado y se sustituye con un nuevo disco. El nodo que se está restaurando no se debe arrancar. Asegúrese de que el cluster se esté ejecutando sin errores antes de llevar a cabo el proceso de restauración.

---

**Nota** – Como debe particionar el disco nuevo con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

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

Este procedimiento proporciona las formas largas de los comandos de 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 RBAC `solaris.cluster.modify` en un nodo del cluster con acceso al conjunto de discos, que *no sea* el nodo que está restaurando.**

Use otro nodo *que no sea* el que va a restaurar.

- 2 **Elimine el nombre de host del nodo que se está restaurando de todos los conjuntos de discos a los que está conectado. Ejecute el siguiente comando una vez para cada conjunto de discos.**

```
metaset -s setname -d -h hostname
```

|                              |                                                                                            |
|------------------------------|--------------------------------------------------------------------------------------------|
| -s <i>setname</i>            | Especifica el nombre del metaconjunto.                                                     |
| -f                           | Elimina el último host del conjunto de discos.                                             |
| -d                           | Suprime desde el metaconjunto.                                                             |
| -h <i>lista_nodos</i>        | Especifica el nombre del nodo que se va a suprimir del metaconjunto.                       |
| -h <i>hostname</i>           | Especifica el nombre del host.                                                             |
| -m <i>mediator_host_list</i> | Especifica el nombre del host mediador que se agregará o suprimirá del conjunto de discos. |

- 3 Si el nodo es un host mediador de dos cadenas, elimine el mediador. Ejecute el siguiente comando una vez para cada conjunto de discos a los que está conectado el nodo.

```
metaset -ssetname-d -m hostname
```

- 4 Reemplace el disco defectuoso en el nodo en el que se va a restaurar el sistema de archivos raíz (/).

Consulte los procedimientos de reemplazo de disco en la documentación incluida con el servidor.

- 5 Inicie el nodo que va a restaurar. El nodo reparado se inicia en modo de un solo usuario desde el CD-ROM, de modo que Solaris Volume Manager no se está ejecutando en el nodo.

- Si está utilizando el CD del sistema operativo Oracle Solaris, tenga en cuenta lo siguiente:

- SPARC: Escriba:

```
ok boot cdrom -s
```

- x86: Inserte el CD en la unidad de CD del sistema e inicie el sistema (cierre el sistema y, luego, apáguelo y enciéndalo). En la pantalla Current Boot Parameters (Parámetros de inicio actuales), escriba b o i.

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

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

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

- Si utiliza un servidor Solaris JumpStart, tenga en cuenta lo siguiente:

- SPARC: Escriba:

```
ok boot net -s
```

- x86: Inserte el CD en la unidad de CD del sistema e inicie el sistema (cierre el sistema y, luego, apáguelo y enciéndalo). En la pantalla Current Boot Parameters (Parámetros de inicio actuales), escriba b o i.

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

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

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

- 6 Cree todas las particiones e intercambie el espacio del disco raíz mediante el comando `format`.**  
Vuelva a crear el esquema de partición original que estaba en el disco defectuoso.

- 7 Cree el sistema de archivos raíz (/) y otros sistemas de archivos según corresponda, mediante el comando `newfs`.**

Vuelva a crear los sistemas de archivos originales que estaban en el disco defectuoso.

---

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

---

- 8 Monte el sistema de archivos raíz (/) en un punto de montaje temporal.**

```
mount device temp-mountpoint
```

- 9 Use los comandos siguientes para restaurar el sistema de archivos raíz (/).**

```
cd temp-mountpoint
ufsrestore rvf dump-device
rm restoresymtable
```

- 10 Instale un nuevo bloque de inicio en el disco nuevo.**

```
/usr/sbin/installboot /usr/platform/'uname -i'/lib/fs/ufs/bootblk
raw-disk-device
```

- 11 Elimine las líneas en el archivo `/temp-mountpoint/etc/system` para información raíz MDD.**

```
* Begin MDD root info (do not edit)
forceload: misc/md_trans
forceload: misc/md_raid
forceload: misc/md_mirror
forceload: misc/md_hotspares
forceload: misc/md_stripe
forceload: drv/pcipsy
forceload: drv/glm
forceload: drv/sd
rootdev:/pseudo/md@0:0,10,blk
* End MDD root info (do not edit)
```

- 12 Edite el archivo `/temp-mountpoint/etc/vfstab` para cambiar la entrada raíz de un volumen de Solaris Volume Manager a un segmento normal correspondiente para cada sistema de archivos del disco raíz que forma parte del metadispositivo o volumen.**

Example:

Change from–

```
/dev/md/dsk/d10 /dev/md/rdisk/d10 / ufs 1 no -
```

Change to–

```
/dev/dsk/c0t0d0s0 /dev/rdisk/c0t0d0s0 / ufs 1 no -
```

**13 Desmonte el sistema de archivos temporales y compruebe el dispositivo de disco sin formato.**

```
cd /
umount temp-mountpoint
fsck raw-disk-device
```

**14 Rearranque el nodo en modo multiusuario.**

```
reboot
```

**15 Sustituya el ID del dispositivo.**

```
cldevice repair rootdisk
```

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

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

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

-af *raw-disk-device* Crea réplicas de bases de datos de estado iniciales en el dispositivo de disco sin formato designado.

**17 Desde un nodo del cluster que no sea el restaurado, agregue el nodo restaurado a todos los conjuntos de discos.**

```
phys-schost-2# metaset -s setname -a -h nodelist
```

-a Agrega (crea) el metaconjunto.

Configure el volumen/reflejo para la raíz (/) de acuerdo con la documentación.

El nodo se rearranca en modo de cluster.

**18 Si el nodo era un host mediador de dos cadenas, vuelva a agregar el mediador.**

```
phys-schost-2# metaset -s setname -a -m hostname
```

**Ejemplo 12-6 Restauración de un sistema de archivos raíz (/) que se encontraba en un volumen de Solaris Volume Manager**

En el ejemplo siguiente, se muestra el sistema de archivos raíz (/) restaurado en el nodo phys-schost-1 desde el dispositivo de cinta /dev/rmt/0. El comando metaset se ejecuta desde otro nodo en el cluster, phys-schost-2, para eliminar el nodo phys-schost-1 y, luego, volver a agregarlo al metaconjunto schost-1. Todos los demás comandos se ejecutan desde phys-schost-1. Se crea un bloque de inicio nuevo en /dev/rdisk/c0t0d0s0 y se vuelven a crear tres réplicas de la base de datos de estado en /dev/rdisk/c0t0d0s4.

[Become superuser or assume a role that provides solaris.cluster.modify RBAC authorization on a cluster node with access to the metaset, other than the node to be restored.]

[Remove the node from the metaset:]

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

[Replace the failed disk and boot the node:]

Inicie el nodo desde el CD del sistema operativo Oracle Solaris:

- SPARC: Escriba:

```
ok boot cdrom -s
```

- x86: Inserte el CD en la unidad de CD del sistema e inicie el sistema (cierre el sistema y, luego, apáguelo y enciéndalo). En la pantalla Current Boot Parameters (Parámetros de inicio actuales), escriba b o i.

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

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

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

```
[Use format and newfs to recreate partitions and file systems
.]
[Mount the root file system on a temporary mount point:]
mount /dev/dsk/c0t0d0s0 /a
[Restore the root file system:]
cd /a
ufsrestore rvf /dev/rmt/0
rm restoresymtable
[Install a new boot block:]
/usr/sbin/installboot /usr/platform/'uname \
-i' /lib/fs/ufs/bootblk /dev/rdisk/c0t0d0s0

[Remove the lines in / temp-mountpoint/etc/system file for MDD root information:
]
* Begin MDD root info (do not edit)
forceload: misc/md_trans
forceload: misc/md_raid
forceload: misc/md_mirror
forceload: misc/md_hotspares
forceload: misc/md_stripe
forceload: drv/pcipsy
forceload: drv/glm
forceload: drv/sd
rootdev:/pseudo/md@0:0,10,blk
* End MDD root info (do not edit)
[Edit the /temp-mountpoint/etc/vfstab file]
Example:
Change from-
/dev/md/dsk/d10 /dev/md/rdisk/d10 / ufs 1 no -

Change to-
/dev/dsk/c0t0d0s0 /dev/rdisk/c0t0d0s0 /usr ufs 1 no -
[Unmount the temporary file system and check the raw disk device:]
cd /
umount /a
```

```
fsck /dev/rdisk/c0t0d0s0
[Reboot:]
reboot
[Replace the disk ID:]
cldevice repair /dev/rdisk/c0t0d0
[Re-create state database replicas:]
metadb -c 3 -af /dev/rdisk/c0t0d0s4
[Add the node back to the metaset:]
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

# Administración de Oracle Solaris Cluster con las interfaces gráficas de usuario

---

Este capítulo proporciona descripciones de las herramientas de la interfaz gráfica de usuario de Oracle Solaris Cluster Manager, que podrá utilizar para administrar muchos aspectos de un cluster. También contiene procedimientos para configurar e iniciar Oracle Solaris Cluster Manager. La ayuda en línea incluida con la GUI de Oracle Solaris Cluster Manager contiene instrucciones para realizar distintas tareas administrativas de Oracle Solaris Cluster.

Este capítulo incluye lo siguiente:

- “Descripción general sobre Oracle Solaris Cluster Manager” en la página 329
- “Configuración de Oracle Solaris Cluster Manager” en la página 330
- “Inicio del software de Oracle Solaris Cluster Manager” en la página 333

## Descripción general sobre Oracle Solaris Cluster Manager

Oracle Solaris Cluster Manager es una GUI que permite visualizar información del cluster, comprobar el estado de los componentes del cluster y supervisar los cambios aplicados a la configuración. Oracle Solaris Cluster Manager también permite realizar numerosas tareas administrativas para los componentes siguientes de Oracle Solaris Cluster:

- Adaptadores
- Cables
- Servicios de datos
- Dispositivos globales
- Interconexiones
- Uniones
- Límites de carga de nodos
- Dispositivos NAS
- Nodos
- Dispositivos de quórum
- Grupos de recursos
- Recursos

En las ubicaciones siguientes encontrará información sobre la instalación y el uso de Oracle Solaris Cluster Manager:

- **Instalación de Oracle Solaris Cluster Manager:** Consulte la [Guía de instalación del software de Oracle Solaris Cluster](#).
- **Inicio de Oracle Solaris Cluster Manager:** consulte “Inicio del software de Oracle Solaris Cluster Manager” en la página 333.
- **Configuración de números de puerto, direcciones de servidor, certificados de seguridad y usuarios:** consulte “Configuración de Oracle Solaris Cluster Manager” en la página 330.
- **Instalación y administración de aspectos del cluster mediante Oracle Solaris Cluster Manager:** consulte la ayuda en línea que se proporciona con Oracle Solaris Cluster Manager.
- **Regeneración de las claves de seguridad de Oracle Solaris Cluster Manager:** consulte “Regeneración de las claves de seguridad de Common Agent Container” en la página 332.

---

**Nota** – Sin embargo, Oracle Solaris Cluster Manager no puede realizar de momento todas las tareas administrativas de Oracle Solaris Cluster. En algunas operaciones debe utilizarse la interfaz de línea de comandos.

---

## Configuración de Oracle Solaris Cluster Manager

Oracle Solaris Cluster Manager es una GUI válida para administrar y visualizar el estado de todos los aspectos de los dispositivos de quórum, los grupos de IPMP, los componentes de interconexión y los dispositivos globales. La GUI puede utilizarse en lugar de muchos comandos de la interfaz de línea de comandos de Oracle Solaris Cluster.

El procedimiento para instalar Oracle Solaris Cluster Manager en el cluster se explica en la [Guía de instalación del software de Oracle Solaris Cluster](#). La ayuda en línea de Oracle Solaris Cluster Manager proporciona instrucciones para realizar varias tareas con la GUI.

Esta sección contiene los procedimientos siguientes para volver a configurar Oracle Solaris Cluster Manager tras la instalación inicial.

- “Configuración de funciones de RBAC” en la página 330
- “Cambio de la dirección del servidor para Oracle Solaris Cluster Manager” en la página 332
- “Regeneración de las claves de seguridad de Common Agent Container” en la página 332

## Configuración de funciones de RBAC

Oracle Solaris Cluster Manager utiliza RBAC para determinar quién tiene derechos para administrar el cluster. Algunos perfiles de derechos de RBAC están incluidos en el software Oracle Solaris Cluster. Puede asignar estos perfiles de derechos a los usuarios o a las funciones

para proporcionar a los usuarios niveles diferentes de acceso a Oracle Solaris Cluster. Para obtener más información sobre cómo configurar y administrar RBAC para el software Oracle Solaris Cluster, consulte el [Capítulo 2, “Oracle Solaris Cluster y RBAC”](#).

## ▼ **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 cluster 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 cluster.**

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

- 2 Detenga Sun Java Web Console.**

```
/usr/sbin/smcwebserver stop
```

- 3 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 secuencias de comandos | <code>commandstream-adapter-port</code> |

- 4 Cambie un número de puerto.**

```
/opt/bin/cacaoadm set-param parameterName=parameterValue
```

- 5 Repita el [Paso 4](#) en cada nodo del cluster.**

- 6 Reinicie Sun Java Web Console.**

```
/usr/sbin/smcwebserver start
```

- 7 Reinicie el daemon de gestión de contenedor de agentes común en todos los nodos del cluster.

```
/opt/bin/cacaoadm start
```

## ▼ Cambio de la dirección del servidor para Oracle Solaris Cluster Manager

Si cambia el nombre de host de un nodo del cluster, debe cambiar la dirección desde la que se ejecuta Oracle Solaris Cluster Manager. El certificado de seguridad predeterminado se genera conforme al nombre de host del nodo al instalarse Oracle Solaris Cluster Manager. Para restablecer el nombre de host del nodo, suprima el archivo de certificado, `keystore`, y reinicie Oracle Solaris Cluster Manager. Oracle Solaris Cluster Manager crea automáticamente un nuevo archivo de certificado con el nombre de host nuevo. Aplique este procedimiento en todos los nodos cuyo nombre de host se haya cambiado.

- 1 Elimine el archivo de certificado, `keystore`, ubicado en `/etc/opt/webconsole`.

```
cd /etc/opt/webconsole
pkgrm keystore
```

- 2 Reinicie Oracle Solaris Cluster Manager.

```
/usr/sbin/smcwebserver restart
```

## ▼ Regeneración de las claves de seguridad de Common Agent Container

Oracle Solaris Cluster Manager usa potentes técnicas de cifrado para asegurar una comunicación segura entre el servidor web de Oracle Solaris Cluster Manager y todos los nodos del cluster.

Las claves utilizadas por Oracle Solaris Cluster Manager se almacenan en el directorio `/etc/opt/SUNWcacao/security` en todos los nodos. Deben ser idénticas en todos los nodos del cluster.

Durante el funcionamiento normal, estas claves se pueden quedar en la configuración predeterminada. Si cambia el nombre de host de un nodo del cluster, debe regenerar las claves de seguridad de Common Agent Container. También podría tener que regenerar las claves si hubiese un posible peligro para la clave, por ejemplo un peligro en la raíz o root del equipo. Para regenerar las claves de seguridad, siga este siguiente.

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

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

**2 Regenera las claves de seguridad en un nodo del cluster.**

```
phys-schost-1# /opt/bin/cacoadm create-keys --force
```

**3 Reinicie el daemon de administración de contenedor de agentes común en el nodo en que regeneró las claves de seguridad.**

```
phys-schost-1# /opt/bin/cacoadm start
```

**4 Cree un archivo tar del directorio /etc/cacao/instances/default .**

```
phys-schost-1# cd /etc/cacao/instances/default
phys-schost-1# tar cf /tmp/SECURITY.tar security
```

**5 Copie el archivo /tmp/Security . tar en todos los nodos del cluster.****6 Extraiga los archivos de seguridad de todos los nodos en los que haya copiado el archivo /tmp/SECURITY . tar.**

Se sobrescriben todos los archivos de seguridad que ya existan en el directorio /etc/opt/SUNWcacao/.

```
phys-schost-2# cd /etc/cacao/instances/default
phys-schost-2# tar xf /tmp/SECURITY.tar
```

**7 Elimine el archivo /tmp/SECURITY . tar de todos los nodos del cluster.**

Es preciso eliminar todas las copias del archivo tar para evitar riesgos de seguridad.

```
phys-schost-1# rm /tmp/SECURITY.tar
```

```
phys-schost-2# rm /tmp/SECURITY.tar
```

**8 Reinicie el daemon de administración contenedor de agentes común en todos los nodos.**

```
phys-schost-1# /opt/bin/cacoadm start
```

**9 Reinicie Oracle Solaris Cluster Manager.**

```
/usr/sbin/smcwebserver restart
```

## Inicio del software de Oracle Solaris Cluster Manager

La GUI de Oracle Solaris Cluster Manager permite administrar de forma sencilla los diferentes aspectos del software Oracle Solaris Cluster. Consulte la ayuda en línea de Oracle Solaris Cluster Manager para obtener más información.

Sun Java Web Console y Common Agent Container se inician de forma automática al iniciar el cluster. Si necesita comprobar que Sun Java Web Console y Common Agent Container estén en ejecución, consulte la sección Resolución de problemas inmediatamente después de este procedimiento.

## ▼ Inicio de Oracle Solaris Cluster Manager

Este procedimiento muestra cómo iniciar Oracle Solaris Cluster Manager en el cluster.

- 1 **Determine si tiene intención de acceder a Oracle Solaris Cluster Manager mediante el nombre y la contraseña del usuario root del nodo del cluster, o de configurar un nombre y una contraseña de usuario diferentes.**
  - Si va a acceder a Oracle Solaris Cluster Manager con el nombre de usuario root del nodo del cluster, vaya al [Paso 5](#).
  - Si pretende configurar un nombre y una contraseña de usuario diferentes, vaya al [Paso 3](#) para establecer las cuentas de usuario de Oracle Solaris Cluster Manager.
- 2 **Conviértase en superusuario en un nodo del cluster.**

- 3 **Cree una cuenta de usuario para acceder al cluster mediante Oracle Solaris Cluster Manager.**

Use el comando `useradd(1M)` para agregar una cuenta de usuario al sistema. Debe configurarse al menos una cuenta de usuario para acceder a Oracle Solaris Cluster Manager si no usa la cuenta del sistema root. Las cuentas de usuario de Oracle Solaris Cluster Manager sólo las utiliza Oracle Solaris Cluster Manager. Estas cuentas no se corresponden con ninguna cuenta de usuario del sistema del sistema operativo Oracle Solaris. En “[Creación y asignación de una función de RBAC con un perfil de derechos de administración de Oracle Solaris Cluster](#)” en la [página 55](#) se describe con más detalle cómo crear y asignar una función de RBAC a una cuenta de usuario.

---

**Nota** – Los usuarios que no tengan establecida una cuenta de usuario en un nodo determinado no pueden acceder al cluster mediante Oracle Solaris Cluster Manager desde ese nodo; tampoco pueden administrar el nodo a través de otro nodo del cluster al que los usuarios tengan acceso.

---

- 4 **(Opcional) Repita el [Paso 3](#) para configurar más cuentas de usuario.**
- 5 **Inicie un navegador desde la consola de administración u otro equipo fuera del cluster.**
- 6 **Compruebe que los tamaños del disco del navegador y de la antememoria estén configurados con un valor superior a 0.**
- 7 **Compruebe que Java y Javascript estén habilitados en el navegador.**
- 8 **Conéctese al puerto de Oracle Solaris Cluster Manager en un nodo del cluster desde el navegador.**

El número de puerto predeterminado es 6789.

**https://node:6789/**

**9 Acepte los certificados que presente el navegador web.**

Aparece la página de inicio de sesión de Java Web Console.

**10 Indique el nombre y la contraseña del usuario que quiera que tenga acceso a Oracle Solaris Cluster Manager.****11 Haga clic en el botón Log In (Iniciar sesión).**

Aparece la página de inicio de la aplicación de Java Web Console.

**12 Haga clic en el vínculo de Oracle Solaris Cluster Manager en la categoría Systems (Sistemas).****13 Acepte cualquier otro certificado que presente el navegador web.****14 Si no puede conectarse a Oracle Solaris Cluster Manager, efectúe los subpasos siguientes para determinar si se ha elegido un perfil de red restringido durante la instalación de Solaris y para restaurar el acceso externo al servicio de Java Web Console.**

Si elige un perfil de red restringido durante la instalación de Oracle Solaris, el acceso externo al servicio de Sun Java Web Console está restringido. Esta red es necesaria para usar la GUI de Oracle Solaris Cluster Manager.

**a. Determine si el servicio de Java Web Console está restringido.**

```
svcprop /system/webconsole:console | grep tcp_listen
```

Si el valor de la propiedad de `tcp_listen` no es `true`, el servicio de la consola web está restringido.

**b. Restaure el acceso externo al servicio de Java Web Console.**

```
svccfg
svc:> select system/webconsole
svc:/system/webconsole> setprop options/tcp_listen=true
svc:/system/webconsole> quit
/usr/sbin/smcwebserver restart
```

**c. Compruebe que el servicio esté disponible.**

```
netstat -a | grep 6789
```

Si el servicio está disponible, la salida del comando devuelve una entrada de 6789, que es el número de puerto que se utiliza para conectarse con Java Web Console.

**Errores más frecuentes**

- Si después de realizar este proceso no se puede conectar con Oracle Solaris Cluster Manager, determine si Sun Java Web Console está en funcionamiento mediante el comando `/usr/sbin/smcwebserver status`. Si Sun Java Web Console no está en funcionamiento, haga que se inicie de forma manual con el comando `/usr/sbin/smcwebserver start`. Si

todavía no puede conectarse con Oracle Solaris Cluster Manager, determine si Common Agent Container se está ejecutando mediante `usr/sbin/cacaoadm status`. Si Common Agent Container no se está ejecutando, inícielo de forma manual mediante `/usr/sbin/cacaoadm start`.

- Si recibe un mensaje de error del sistema al intentar ver información sobre un nodo distinto al nodo que ejecuta la GUI, compruebe si el parámetro de la dirección de enlace a la red del contenedor de agentes común está definido en el valor correcto de `0.0.0.0`.

Realice estos pasos en cada nodo del cluster.

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
```

## Ejemplo

---

### Configuración de replicación de datos basada en host con el software de 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 clusters. Consulte [“Comprensión de la replicación de datos” en la página 88](#).

El ejemplo de este apéndice muestra cómo configurar una replicación de datos basada en host entre clusters con el software de Sun StorageTek Availability Suite 4.0. El ejemplo muestra una configuración de cluster 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 cluster global. El ejemplo no incluye todos los pasos necesarios para otras aplicaciones ni las configuraciones de otros clusters.

Si utiliza el control de acceso basado en funciones (RBAC) en lugar de ser superusuario para acceder a los nodos del cluster, 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 la [System Administration Guide: Security Services](#) para obtener más información sobre el uso de roles de RBAC. Consulte las páginas del comando `man` de Oracle Solaris Cluster para saber la autorización RBAC que requiere cada subcomando de Oracle Solaris Cluster.

## Descripción del software de Availability Suite en un cluster

En esta sección, se presenta la tolerancia ante problemas graves y se describen los métodos de replicación de datos que utiliza el software de Availability Suite.

La tolerancia a desastres es la capacidad de restaurar una aplicación en un cluster alternativo cuando el cluster 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 cluster secundario estableciendo en línea uno o más grupos de recursos y grupos de dispositivos.

Si los datos se replican de manera síncrona entre el cluster 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 cluster secundario antes del fallo del sitio principal y, por lo tanto, se pierdan.

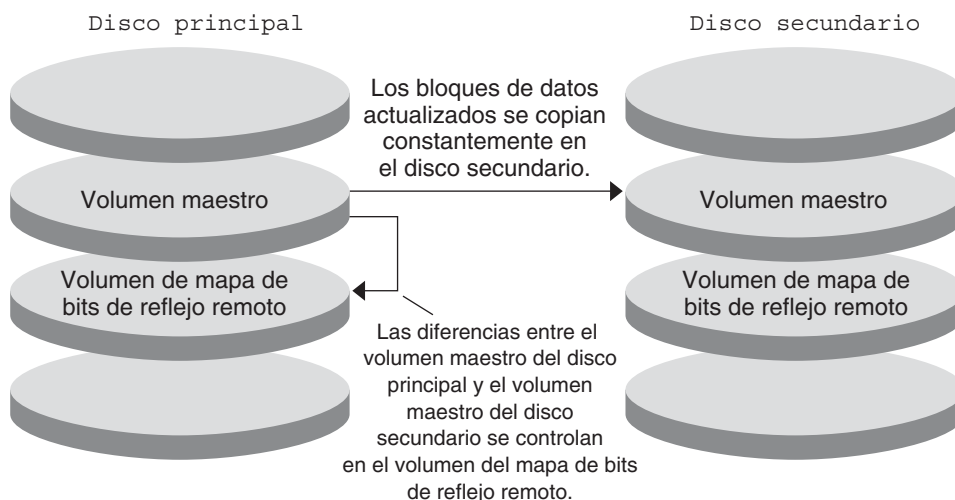
## Métodos de replicación de datos utilizados por el software de Availability Suite

En esta sección, se describen los métodos de replicación de reflejo remoto y de instantánea de un momento determinado utilizados por el software de Availability Suite. Este software utiliza los comandos `sndradm(1RPC)` e `iiadm(1II)` para replicar datos. Para obtener más información sobre estos comandos, consulte la [Availability Suite documentation](#).

### 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



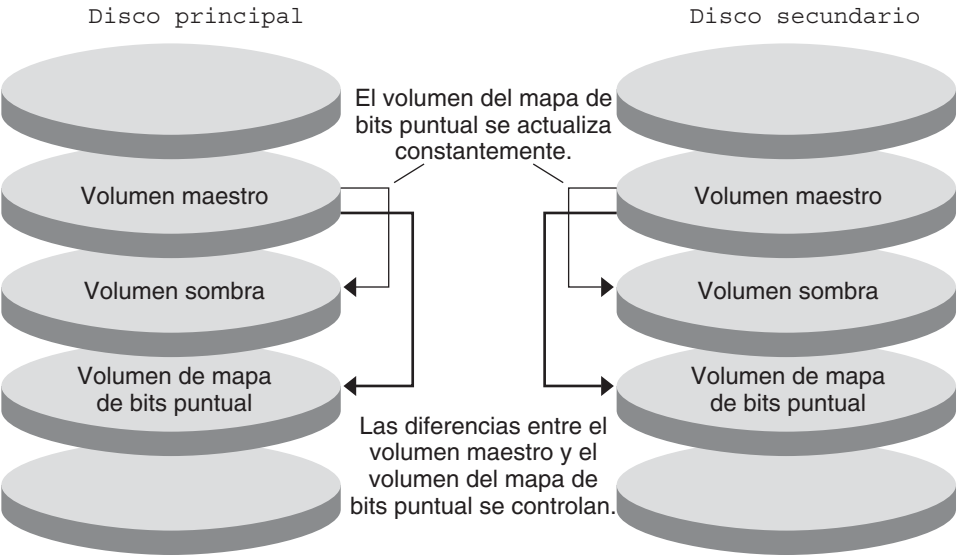
La replicación por duplicación remota se puede efectuar de manera sincrónica en tiempo real o de manera asincrónica. Cada volumen definido en cada cluster se puede configurar de manera individual, para replicación sincrónica o asincrónica.

- 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 instantánea de un momento determinado. Los datos del volumen primario de cada disco se copian en el volumen sombra del mismo disco. El mapa de bits instantáneo controla y detecta las diferencias entre el volumen primario y el volumen sombra. Cuando los datos se copian en el volumen sombra, el mapa de bits de un momento determinado se vuelve a configurar.

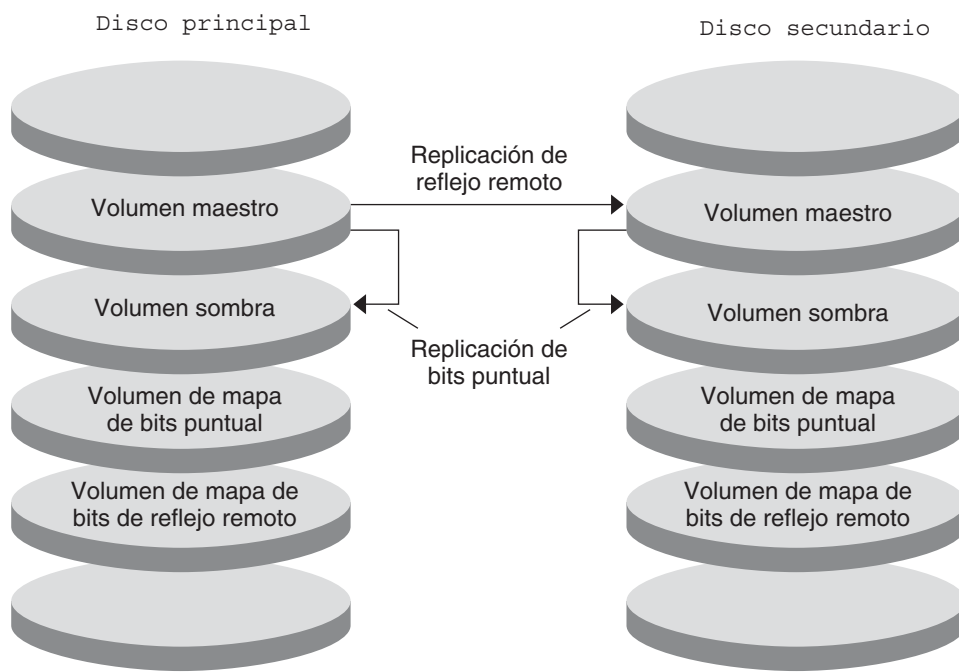
FIGURA A-2 Instantánea puntual



### 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 clusters

Esta sección proporciona directrices para la configuración de replicación de datos entre clusters. 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 cluster.

En esta sección se analizan los aspectos siguientes:

- [“Configuración de grupos de recursos de replications” en la página 342](#)
- [“Configuración de grupos de recursos de aplicaciones” en la página 342](#)
  - [“Configuración de grupos de recursos en una aplicación de migración tras error” en la página 343](#)
  - [“Configuración de grupos de recursos en una aplicación escalable” en la página 344](#)
- [“Directrices para la gestión de una recuperación” en la página 345](#)

## Configuración de grupos de recursos de replicaciones

Los grupos de recursos de replicación sitúan al grupo de dispositivos bajo el control del software de Availability Suite con el recurso de nombre de host lógico. Un grupo de recursos de replicaciones debe tener las características siguientes:

- Ser un grupo de recursos de migración tras error.

Un recurso de migración tras error sólo puede ejecutarse en un nodo a la vez. Cuando se produce la migración tras error, los recursos de recuperación participan en ella.

- Debe tener un recurso de nombre de host lógico.

El nombre de host lógico debe estar alojado por el cluster principal. Después de una conmutación por error, el nombre de host lógico debe estar alojado por el cluster secundario. El sistema DNS se utiliza para asociar el nombre de host lógico con un cluster.

- 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*.
- *ZPoolsSearchDir*. Esta propiedad de extensión es necesaria para utilizar el sistema de archivos ZFS.

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 cluster primario y en el secundario

## Configuración de grupos de recursos de aplicaciones

Para que esté totalmente disponible, una aplicación se debe administrar como un recurso en un grupo de recursos de aplicaciones. Un grupo de recursos de aplicaciones se puede configurar en una aplicación de migración tras error o en una aplicación escalable.

Los recursos de aplicaciones y los grupos de recursos de aplicaciones configurados en el cluster primario también se deben configurar en el cluster secundario. Asimismo, se deben replicar en el cluster 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 343](#)
- [“Configuración de grupos de recursos en una aplicación escalable” en la página 344](#)

## 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 cluster. Un grupo de recursos de una aplicación de migración tras error debe tener las características siguientes:

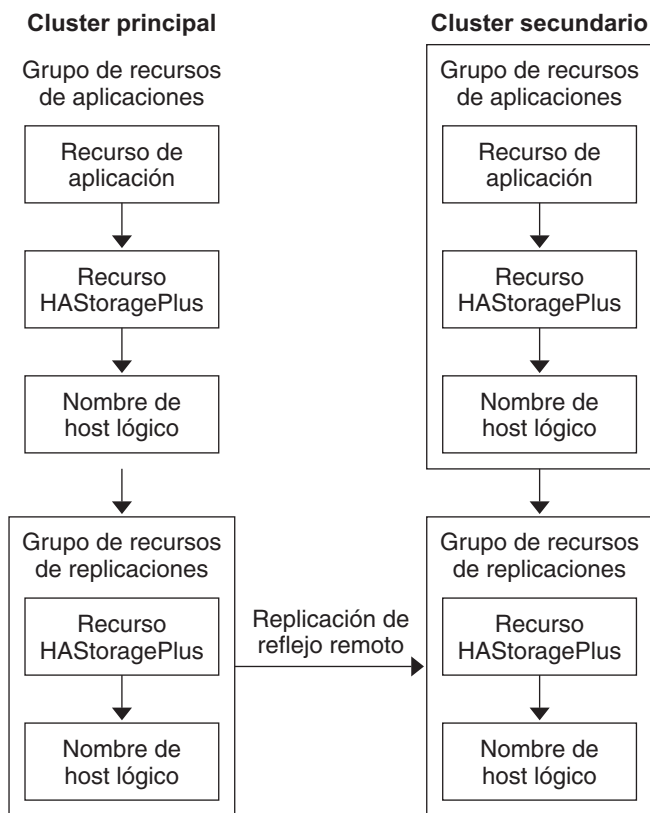
- Debe tener un recurso de `HAStoragePlus` para aplicar la conmutación por error del sistema de archivos o `zpool` cuando el grupo de recursos de aplicaciones se conmuta.  
El grupo de dispositivos se coloca con el grupo de recursos de replications y el grupo de recursos de aplicaciones. Por este motivo, 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 cluster principal y sin conexión en el cluster secundario.  
El grupo de recursos de aplicaciones debe estar en línea en el cluster secundario cuando éste hace las funciones de cluster 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.

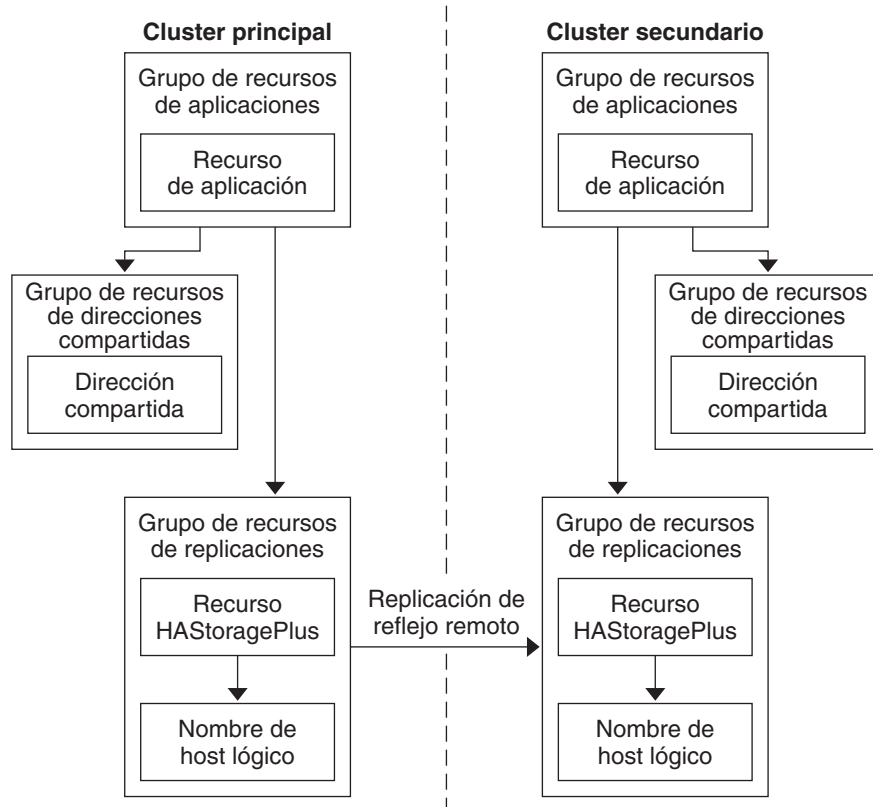
Cuando una aplicación escalable se administra como recurso en un grupo de recursos de aplicaciones, no es necesario acoplar el grupo de recursos de aplicaciones con el grupo de dispositivos. Por este motivo, no es necesario crear un recurso de HAStoragePlus para el grupo de recursos de aplicaciones.

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 cluster 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

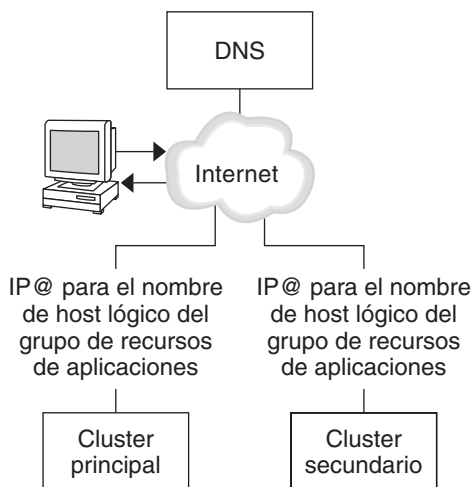


## Directrices para la gestión de una recuperación

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

Los clientes utilizan 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 cluster 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. En la [Figura A-6](#), se muestra cómo el DNS asigna un cliente a un cluster.

FIGURA A-6 Asignación del DNS de un cliente a un cluster



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 371](#).

Después de la reparación, el cluster principal se puede volver a colocar en línea. Para volver al cluster primario original, siga estos pasos:

1. Sincronice el cluster 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 cluster principal.
4. Actualice el DNS de modo que los clientes tengan acceso a la aplicación en el cluster primario.

## Mapa de tareas: ejemplo de configuración de una replicación de datos

En la [Tabla A-1](#), se enumeran las tareas de este ejemplo de configuración de replicación de datos para una aplicación NFS mediante el software de Availability Suite.

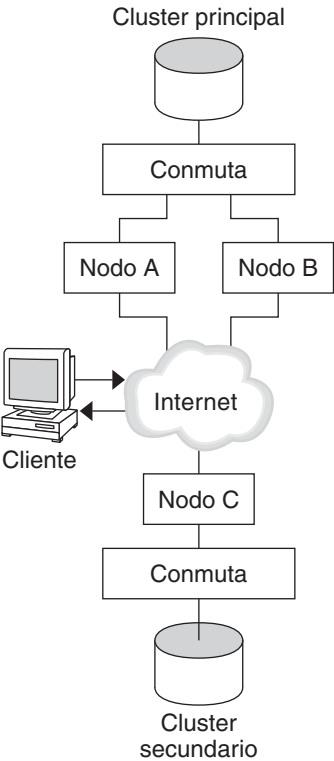
TABLA A-1 Mapa de tareas: ejemplo de configuración de una replicación de datos

| Tarea                                                                                                                                                        | Instrucciones                                                                                                                                          |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1. Conectar e instalar los clusters                                                                                                                          | “Conexión e instalación de clusters” en la página 347                                                                                                  |
| 2. Configurar los grupos de dispositivos, los sistemas de archivos para la aplicación de NFS y los grupos de recursos de los clusters principal y secundario | “Ejemplo de configuración de grupos de dispositivos y de recursos” en la página 349                                                                    |
| 3. Habilitar la replicación de datos en el cluster primario y en el secundario                                                                               | “Habilitación de la replicación en el cluster primario” en la página 363<br>“Habilitación de la replicación en el cluster secundario” en la página 365 |
| 4. Efectuar la replicación de datos                                                                                                                          | “Replicación por duplicación remota” en la página 366<br>“Instantánea de un momento determinado” en la página 368                                      |
| 5. Comprobar la configuración de la replicación de datos                                                                                                     | “Procedimiento para comprobar la configuración de la replicación” en la página 369                                                                     |

## Conexión e instalación de clusters

La [Figura A-7](#) ilustra la configuración de clusters utilizada en la configuración de ejemplo. El cluster secundario de la configuración de ejemplo contiene un nodo, pero se pueden usar otras configuraciones de cluster.

FIGURA A-7 Ejemplo de configuración de cluster



La [Tabla A-2](#) resume el hardware y el software necesarios para la configuración de ejemplo. El sistema operativo Oracle Solaris, el software de Oracle Solaris Cluster y el software de administrador de volúmenes deben instalarse en los nodos del cluster *antes* de instalar los parches y el software de Availability Suite.

TABLA A-2 Hardware y software necesarios

| Hardware o software | Requisito                                                                                                                                                                                                                                                                                   |
|---------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Hardware de nodo    | El software de Availability Suite es compatible con todos los servidores que utilizan el sistema operativo Oracle Solaris.<br><br>Para obtener información sobre el hardware que se debe usar, consulte el <a href="#">Oracle Solaris Cluster 3.3 3/13 Hardware Administration Manual</a> . |
| Espacio en el disco | Aproximadamente 15 MB.                                                                                                                                                                                                                                                                      |

TABLA A-2 Hardware y software necesarios (Continuación)

| Hardware o software                        | Requisito                                                                                                                                                                                                                                                                                                                                                         |
|--------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Sistema operativo Oracle Solaris           | <p>Las versiones del sistema operativo Oracle Solaris compatibles con Oracle Solaris Cluster.</p> <p>Todos los nodos deben utilizar la misma versión del sistema operativo Oracle Solaris.</p> <p>Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software de Oracle Solaris Cluster</a></p>                      |
| Software Oracle Solaris Cluster            | <p>Software Oracle Solaris Cluster 3.3.</p> <p>Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software de Oracle Solaris Cluster</a>.</p>                                                                                                                                                                        |
| Software Volume Manager                    | <p>Software de Solaris Volume Manager.</p> <p>Todos los nodos deben usar la misma versión del administrador de volúmenes.</p> <p>Para obtener información sobre la instalación, consulte el <a href="#">Capítulo 4, “Configuración del software de Solaris Volume Manager”</a> de <a href="#">Guía de instalación del software de Oracle Solaris Cluster</a>.</p> |
| Software de Availability Suite             | <p>Para obtener información sobre cómo instalar el software, consulte los manuales de instalación de su versión del software de Availability Suite:</p> <ul style="list-style-type: none"> <li>Documentación de Sun StorageTek Availability Suite 4.0: Sun StorageTek Availability</li> </ul>                                                                     |
| Parches del software de Availability Suite | <p>Para obtener información sobre los últimos parches, inicie sesión en <a href="#">My Oracle Support</a>.</p>                                                                                                                                                                                                                                                    |

## Ejemplo de configuración de grupos de dispositivos y de recursos

Esta sección describe cómo se configuran los grupos de dispositivos y los de recursos en una aplicación NFS. Para obtener más información, consulte “[Configuración de grupos de recursos de replicaciones](#)” en la página 342 y “[Configuración de grupos de recursos de aplicaciones](#)” en la página 342.

Esta sección incluye los procedimientos siguientes:

- “[Configuración de un grupo de dispositivos en el cluster primario](#)” en la página 351
- “[Configuración de un grupo de dispositivos en el cluster secundario](#)” en la página 352
- “[Configuración de sistemas de archivos en el cluster primario para la aplicación NFS](#)” en la página 353
- “[Configuración del sistema de archivos en el cluster secundario para la aplicación NFS](#)” en la página 354

- “Creación de un grupo de recursos de replicaciones en el cluster primario” en la página 356
- “Creación de un grupo de recursos de replicaciones en el cluster secundario” en la página 357
- “Creación de un grupo de recursos de aplicaciones NFS en el cluster primario” en la página 359
- “Creación de un grupo de recursos de aplicaciones NFS en el cluster secundario” en la página 361
- “Procedimiento para comprobar la configuración de la replicación” en la página 369

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

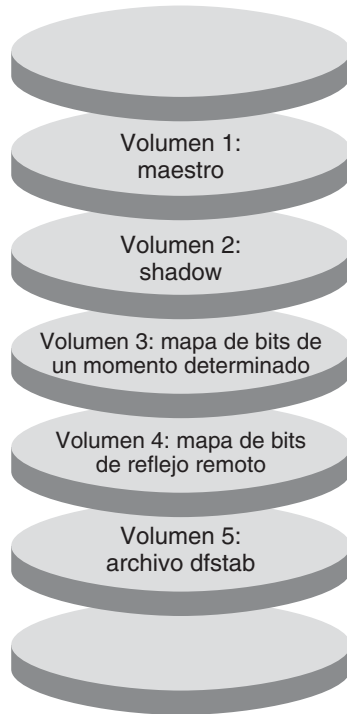
| Grupo o recurso                                      | Nombre                               | Descripción                                                                                                   |
|------------------------------------------------------|--------------------------------------|---------------------------------------------------------------------------------------------------------------|
| Grupo de dispositivos                                | devgrp                               | El grupo de dispositivos                                                                                      |
| El grupo de recursos de replicaciones y los recursos | devgrp-stor-rg                       | El grupo de recursos de replicaciones                                                                         |
|                                                      | lhost-reprg-prim,<br>lhost-reprg-sec | Los nombres de host lógicos para el grupo de recursos de replicaciones en el cluster primario y el secundario |
|                                                      | devgrp-stor                          | El recurso de HAStoragePlus para el grupo de recursos de replicaciones                                        |
| El grupo de recursos de aplicaciones y los recursos  | nfs-rg                               | El grupo de recursos de aplicaciones                                                                          |
|                                                      | lhost-nfsrg-prim,<br>lhost-nfsrg-sec | Los nombres de hosts lógicos para el grupo de recursos de aplicaciones en el cluster 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 replicaciones debe tener un nombre con el formato *nombre\_grupo\_dispositivos*-stor-rg.

En este ejemplo, la configuración utiliza el software de SVM. Para obtener más información sobre el software de Solaris Volume Manager, consulte el [Capítulo 4, “Configuración del software de Solaris Volume Manager” de Guía de instalación del software de Oracle Solaris Cluster](#).

En la siguiente figura, se muestran los volúmenes que se crean en el grupo de dispositivos.

FIGURA A-8 Volúmenes para el grupo de dispositivos



**Nota** – Los volúmenes que se definen en este procedimiento no deben incluir áreas privadas de etiqueta de disco, por ejemplo, cilindro 0. El software de SVM gestiona esta restricción automáticamente.

## ▼ Configuración de un grupo de dispositivos en el cluster primario

### Antes de empezar

Asegúrese de realizar las tareas siguientes:

- Lea las directrices y los requisitos de las secciones siguientes:
  - “Descripción del software de Availability Suite en un cluster” en la página 338
  - “Directrices para la configuración de replicación de datos basada en host entre clusters” en la página 341
- Configure los clusters primario y secundario como se describe en “Conexión e instalación de clusters” en la página 347.

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

El nodo nodeA es el primero del cluster primario. Si desea recordar qué nodo es nodeA, consulte la [Figura A-7](#).

- 2 **Cree un grupo de discos en nodeA que contenga desde el volumen 1, vol01, hasta el volumen 4, vol04.**

- 3 **Configure el grupo de discos para crear un grupo de dispositivos.**

```
nodeA# cldevicegroup create -t svm -n nodeA nodeB devgrp
```

El grupo de dispositivos se denomina devgrp.

- 4 **Cree el sistema de archivos para el grupo de dispositivos.**

```
nodeA# newfs /dev/vx/rdisk/devgrp/vol01 < /dev/null
nodeA# newfs /dev/vx/rdisk/devgrp/vol02 < /dev/null
```

No se necesita ningún sistema de archivos para vol03 o vol04, que se utilizan, en cambio, como volúmenes sin procesar.

**Pasos siguientes** Vaya a [“Configuración de un grupo de dispositivos en el cluster secundario”](#) en la página 352.

## ▼ Configuración de un grupo de dispositivos en el cluster secundario

**Antes de empezar**

Complete el procedimiento [“Configuración de un grupo de dispositivos en el cluster primario”](#) en la página 351.

- 1 **Obtenga acceso al nodo nodeC como superusuario o asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
- 2 **Cree un grupo de discos en nodeC que contenga cuatro volúmenes: desde el volumen 1, vol01, hasta el volumen 4, vol04.**

- 3 **Configure el grupo de discos para crear un grupo de dispositivos.**

```
nodeC# cldevicegroup create -t svm -n nodeC devgrp
```

El grupo de dispositivos se denomina devgrp.

- 4 **Cree el sistema de archivos para el grupo de dispositivos.**

```
nodeC# newfs /dev/vx/rdisk/devgrp/vol01 < /dev/null
nodeC# newfs /dev/vx/rdisk/devgrp/vol02 < /dev/null
```

No se necesita ningún sistema de archivos para vol03 o vol04, que se utilizan, en cambio, como volúmenes sin procesar.

**Pasos siguientes** Vaya a “Configuración de sistemas de archivos en el cluster primario para la aplicación NFS” en la página 353.

## ▼ Configuración de sistemas de archivos en el cluster primario para la aplicación NFS

**Antes de empezar** Complete el procedimiento “Configuración de un grupo de dispositivos en el cluster secundario” en la página 352.

- 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 se monte automáticamente en el punto de montaje.**  
 Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo `nodeA` y el nodo `nodeB`. El texto debe estar en una sola línea.  

```
/dev/vx/dsk/devgrp/vol01 /dev/vx/rdisk/devgrp/vol01 \
/global/mountpoint ufs 3 no global,logging
```

 Para ver los nombres y los números de los volúmenes que se utilizan en el grupo de dispositivos, consulte la [Figura A-8](#).
- 4 **En `nodeA`, cree un volumen para la información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.**  

```
nodeA# vxassist -g devgrp make vol05 120m disk1
```

 El volumen 5, `vol05`, contiene información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.
- 5 **En `nodeA`, vuelva a sincronizar el grupo de dispositivos con el software de Oracle Solaris Cluster.**  

```
nodeA# cldevicegroup sync devgrp
```
- 6 **En `nodeA`, cree el sistema de archivos para `vol05`.**  

```
nodeA# newfs /dev/vx/rdsk/devgrp/vol05
```
- 7 **En `nodeA` y `nodeB`, cree un punto de montaje para `vol05`.**  
 El ejemplo siguiente crea el punto de montaje `/global/etc`.  

```
nodeA# mkdir /global/etc
```

- 8 En **nodeA** y **nodeB**, configure **vol05** para que se monte automáticamente en el punto de montaje.

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo **nodeA** y el nodo **nodeB**. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol05 /dev/vx/rdisk/devgrp/vol05 \
/global/etc ufs 3 yes global,logging
```

- 9 Monte **vol05** en **nodeA**.

```
nodeA# mount /global/etc
```

- 10 Permita que los sistemas remotos puedan acceder a **vol05**.

- a. Cree un directorio con el nombre `/global/etc/SUNW.nfs` en el nodo **nodeA**.

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

- b. Cree el archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en el nodo **nodeA**.

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. Agregue la línea siguiente al archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en el nodo **nodeA**.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

**Pasos siguientes** Vaya a [“Configuración del sistema de archivos en el cluster secundario para la aplicación NFS” en la página 354.](#)

## ▼ Configuración del sistema de archivos en el cluster secundario para la aplicación NFS

**Antes de empezar** Complete el procedimiento [“Configuración de sistemas de archivos en el cluster primario para la aplicación NFS” en la página 353.](#)

- 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/vx/dsk/devgrp/vol01 /dev/vx/rdisk/devgrp/vol01 \
/global/mountpoint ufs 3 no global,logging
```

- 4 En **nodeC**, cree un volumen para la información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.

```
nodeC# vxassist -g devgrp make vol05 120m disk1
```

El volumen 5, **vol05**, contiene información del sistema de archivos utilizada por el servicio de datos de Oracle Solaris Cluster HA para NFS.

- 5 En **nodeC**, vuelva a sincronizar el grupo de dispositivos con el software de Oracle Solaris Cluster.

```
nodeC# cldevicegroup sync devgrp
```

- 6 En **nodeC**, cree el sistema de archivos para **vol05**.

```
nodeC# newfs /dev/vx/rdisk/devgrp/vol05
```

- 7 En **nodeC**, cree un punto de montaje para **vol05**.

El ejemplo siguiente crea el punto de montaje **/global/etc**.

```
nodeC# mkdir /global/etc
```

- 8 En **nodeC**, configure **vol05** para que se monte automáticamente en el punto de montaje.

Agregue o sustituya el texto siguiente en el archivo **/etc/vfstab** del nodo **nodeC**. El texto debe estar en una sola línea.

```
/dev/vx/dsk/devgrp/vol05 /dev/vx/rdisk/devgrp/vol05 \
/global/etc ufs 3 yes global,logging
```

- 9 Monte **vol05** en **nodeC**.

```
nodeC# mount /global/etc
```

- 10 Permita que los sistemas remotos puedan acceder a **vol05**.

- a. Cree un directorio que se denomine **/global/etc/SUNW.nfs** en **nodeC**.

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

- b. Cree el archivo **/global/etc/SUNW.nfs/dfstab.nfs-rs** en **nodeC**.

```
nodeC# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. Agregue la línea siguiente al archivo **/global/etc/SUNW.nfs/dfstab.nfs-rs** en **nodeC**:

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de replicaciones en el cluster primario” en la página 356.](#)

## ▼ Creación de un grupo de recursos de replicaciones en el cluster primario

### Antes de empezar

Complete el procedimiento “Configuración del sistema de archivos en el cluster secundario para la aplicación NFS” en la página 354.

- 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 replicaciones para el grupo de dispositivos.

```
nodeA# clresourcegroup create -n nodeA,nodeB devgrp-stor-rg
```

-n nodeA,nodeB      Especifica que los nodos del cluster `nodeA` y `nodeB` pueden controlar el grupo de recursos de replicaciones.

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.HAStoragePlus` al grupo de recursos de replicaciones.

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HAStoragePlus \
-p GlobalDevicePaths=devgrp \
-p AffinityOn=True \
devgrp-stor
```

-g      Especifica el grupo de recursos en el que se agrega el recurso.

-p GlobalDevicePaths=      Especifica la propiedad de extensión de la que depende el software de Availability Suite.

-p AffinityOn=True      Especifica que el recurso `SUNW.HAStoragePlus` debe realizar un switchover de afinidad para los dispositivos globales y los sistemas de archivos del cluster en -x `GlobalDevicePaths=`. Por ese motivo, si el grupo de recursos de replicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HAStoragePlus(5)`.

- 5 Agregue un recurso de nombre de host lógico al grupo de recursos de replicaciones.

```
nodeA# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-prim
```

El nombre de host lógico para el grupo de recursos de replicaciones en el cluster primario es `lhost-reprg-prim`.

**6 Habilite los recursos, administre el grupo de recursos y póngalo en línea.**

```
nodeA# clresourcegroup online -e -M -n nodeA devgrp-stor-rg
```

-e      Habilita los recursos asociados.

-M      Administra el grupo de recursos.

-n      Especifica el nodo en el que poner el grupo de recursos en línea.

**7 Compruebe que el grupo de recursos esté en línea.**

```
nodeA# clresourcegroup status devgrp-stor-rg
```

Examine el campo de estado del grupo de recursos para confirmar que el grupo de recursos de replicaciones esté en línea para el nodo nodeA.

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de replicaciones en el cluster secundario” en la página 357.](#)

## ▼ Creación de un grupo de recursos de replicaciones en el cluster secundario

**Antes de empezar** Complete el procedimiento [“Creación de un grupo de recursos de replicaciones en el cluster primario” en la página 356.](#)

**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# clresourcetype 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=devgrp \
-p AffinityOn=True \
devgrp-stor
```

|                                    |                                                                                                                                                                                                                                                                                                                                                               |
|------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>create</code>                | Crea el recurso.                                                                                                                                                                                                                                                                                                                                              |
| <code>-t</code>                    | Especifica el tipo de recurso.                                                                                                                                                                                                                                                                                                                                |
| <code>-p GlobalDevicePaths=</code> | Especifica la propiedad de extensión de la que depende el software de Availability Suite.                                                                                                                                                                                                                                                                     |
| <code>-p AffinityOn=True</code>    | Especifica que el recurso <code>SUNW.HAStoragePlus</code> debe realizar un switchover de afinidad para los dispositivos globales y los sistemas de archivos del cluster en <code>-x GlobalDevicePaths=</code> . Por ese motivo, si el grupo de recursos de replications migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta. |
| <code>devgrp-stor</code>           | El recurso de <code>HAStoragePlus</code> para el grupo de recursos de replications.                                                                                                                                                                                                                                                                           |

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HAStoragePlus(5)`.

**5 Agregue un recurso de nombre de host lógico al grupo de recursos de replications.**

```
nodeC# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-sec
```

El nombre de host lógico para el grupo de recursos de replicación en el cluster principal se denomina `lhost-reprg-sec`.

**6 Habilite los recursos, administre el grupo de recursos y póngalo en línea.**

```
nodeC# clresourcegroup online -e -M -n nodeC devgrp-stor-rg
```

|                     |                                                                   |
|---------------------|-------------------------------------------------------------------|
| <code>online</code> | Lo establece en línea.                                            |
| <code>-e</code>     | Habilita los recursos asociados.                                  |
| <code>-M</code>     | Administra el grupo de recursos.                                  |
| <code>-n</code>     | 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 replications esté en línea para el nodo `nodeC`.

**Pasos siguientes** Vaya a “[Creación de un grupo de recursos de aplicaciones NFS en el cluster primario](#)” en la [página 359](#).

## ▼ Creación de un grupo de recursos de aplicaciones NFS en el cluster 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 replications en el cluster secundario” en la página 357.](#)

- 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 grupo de dispositivos `devgrp`.**

```
nodeA# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_dependencies=devgrp-stor-rg \
nfs-rg
```

```
Pathprefix=/global/etc
```

Especifica el directorio en el que los recursos del grupo pueden guardar los archivos de administración.

```
Auto_start_on_new_cluster=False
```

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

```
RG_dependencies=devgrp-stor-rg
```

Especifica el grupo de recursos del que depende el grupo de recursos de aplicaciones. En este ejemplo, el grupo de recursos de aplicaciones depende del grupo de recursos de replicación `devgrp-stor-rg`.

Si se realiza un switchover del grupo de recursos de aplicaciones a un nodo principal nuevo, el grupo de recursos de replicación realiza un switchover automáticamente. Sin embargo, si se realiza un switchover del grupo de recursos de replicación a un nodo principal nuevo, se debe realizar un switchover manual del grupo de recursos de aplicaciones.

```
nfs-rg
```

El nombre del grupo de recursos de aplicaciones.

**5 Agregue un recurso de SUNW.HAStoragePlus al grupo de recursos de aplicaciones.**

```
nodeA# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
-p FileSystemMountPoints=/global/mountpoint \
-p AffinityOn=True \
nfs-dg-rs
```

create

Crea el recurso.

-g

Especifica el grupo de recursos en el que se agrega el recurso.

-t SUNW.HAStoragePlus

Especifica que el recurso es del tipo SUNW.HAStoragePlus.

-p FileSystemMountPoints=/global/

Especifica que el punto de montaje del sistema de archivos es global.

-p AffinityOn=True

Especifica que el recurso de aplicaciones debe efectuar un switchover de afinidad para los dispositivos globales y los sistemas de archivos del cluster definidos por -p GlobalDevicePaths=. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

nfs-dg-rs

El nombre del recurso de HAStoragePlus para la aplicación NFS.

Para obtener más información sobre estas propiedades de extensión, consulte la página 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 cluster primario es

lhost-nfsrg-prim.

**7 Active los recursos, gestione el grupo de recursos de aplicaciones y ponga en línea el grupo de recursos de aplicaciones.****a. Active el recurso HAStoragePlus para la aplicación NFS.**

```
nodeA# clresource enable nfs-rs
```

**b. Ponga en línea el grupo de recursos de aplicaciones en nodeA.**

```
nodeA# clresourcegroup online -e -M -n nodeA nfs-rg
```

online      Pone el grupo de recursos en línea.

-e            Habilita los recursos asociados.

- M Administra el grupo de recursos.
- n Especifica el nodo en el que poner el grupo de recursos en línea.
- nfs - rg El nombre del grupo de recursos.

## 8 Compruebe que el grupo de recursos de aplicaciones esté en línea.

```
nodeA# clresourcegroup status
```

Examine el campo de estado del grupo de recursos para determinar si el grupo de recursos de aplicaciones está en línea para el nodo nodeA y el nodo nodeB.

**Pasos siguientes** Vaya a [“Creación de un grupo de recursos de aplicaciones NFS en el cluster secundario” en la página 361.](#)

## ▼ Creación de un grupo de recursos de aplicaciones NFS en el cluster secundario

**Antes de empezar** Siga los pasos de [“Creación de un grupo de recursos de aplicaciones NFS en el cluster primario” en la página 359.](#)

### 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.HAStoragePlus` no se ha registrado como tipo de recurso, hágalo.

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

### 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_dependencies=devgrp-stor-rg \
nfs-rg
```

```
create
```

Crea el grupo de recursos.

```
-p
```

Especifica una propiedad del grupo de recursos.

```
Pathprefix=/global/etc
```

Especifica un directorio en el que los recursos del grupo pueden guardar los archivos de administración.

`Auto_start_on_new_cluster=False`

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

`RG_dependencias=devgrp-stor-rg`

Especifica los grupos de recursos de los que depende el grupo de recursos de aplicaciones. En este ejemplo, el grupo de recursos de aplicaciones depende del grupo de recursos de replicación.

Si se realiza un switchover del grupo de recursos de aplicaciones a un nodo principal nuevo, el grupo de recursos de replicación realiza un switchover automáticamente. Sin embargo, si se realiza un switchover del grupo de recursos de replicación a un nodo principal nuevo, se debe realizar un switchover manual del grupo de recursos de aplicaciones.

`nfs-rg`

El nombre del grupo de recursos de aplicaciones.

## 5 Agregue un recurso de SUNW.HAStoragePlus al grupo de recursos de aplicaciones.

```
nodeC# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
-p FileSystemMountPoints=/global/mountpoint \
-p AffinityOn=True \
nfs-dg-rs
```

`create`

Crea el recurso.

`-g`

Especifica el grupo de recursos en el que se agrega el recurso.

`-t SUNW.HAStoragePlus`

Especifica que el recurso es del tipo SUNW.HAStoragePlus.

`-p`

Especifica una propiedad del recurso.

`FileSystemMountPoints=/global/`

Especifica que el punto de montaje del sistema de archivos es global.

`AffinityOn=True`

Especifica que el recurso de aplicaciones debe efectuar un switchover de afinidad para los dispositivos globales y los sistemas de archivos del cluster definidos por `-x`

`GlobalDevicePaths=`. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

`nfs-dg-rs`

El nombre del recurso de HAStoragePlus para la aplicación NFS.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando [man SUNW.HAStoragePlus\(5\)](#).

**6 Agregue un recurso de nombre de host lógico al grupo de recursos de aplicaciones.**

```
nodeC# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-sec
```

El nombre de host lógico del grupo de recursos de aplicaciones del cluster 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 Asegúrese de que el grupo de recursos de aplicaciones no se ponga en línea en nodeC.**

```
nodeC# clresource disable -n nodeC nfs-rs
nodeC# clresource disable -n nodeC nfs-dg-rs
nodeC# clresource disable -n nodeC lhost-nfsrg-sec
nodeC# clresourcegroup online -n "" nfs-rg
```

El grupo de recursos permanece sin conexión tras reiniciarse debido a *Auto\_start\_on\_new\_cluster=False*.

**9 Si el volumen global se monta en el cluster primario, desmonte el volumen global del cluster secundario.**

```
nodeC# umount /global/mountpoint
```

Si el volumen está montado en un cluster secundario, se da un error de sincronización.

**Pasos siguientes** Vaya a [“Ejemplo de cómo habilitar la replicación de datos”](#) en la página 363.

## 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` y `iiadm` del software de Availability Suite. Para obtener más información sobre estos comandos, consulte la documentación de Availability Suite.

Esta sección incluye los procedimientos siguientes:

- [“Habilitación de la replicación en el cluster primario”](#) en la página 363
- [“Habilitación de la replicación en el cluster secundario”](#) en la página 365

### ▼ Habilitación de la replicación en el cluster primario

**1 Acceda al nodo nodeA como superusuario o asuma un rol que proporcione la autorización RBAC `solaris.cluster.read`.****2 Purgue todas las transacciones.**

```
nodeA# lockfs -a -f
```

### 3 Compruebe 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 cluster primario al secundario.

Este paso permite la replicación del volumen maestro en el cluster principal al volumen maestro en el cluster secundario. Además, este paso permite la replicación en el mapa de bits de reflejo remoto en `vol04`.

- Si los clusters principal y secundario no están sincronizados, ejecute este comando para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

- Si los clusters principal y secundario están sincronizados, ejecute este comando para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -E lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

### 5 Habilite la sincronización automática.

Ejecute este comando para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -a on lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Este paso habilita la sincronización automática. Si el estado activo de la sincronización automática es `on`, los conjuntos de volúmenes se vuelven a sincronizar cuando el sistema reinicie o cuando haya un error.

### 6 Compruebe que el cluster esté en modo de registro.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

## 7 Habilite la instantánea puntual.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/iiadm -e ind \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
nodeA# /usr/sbin/iiadm -w \
/dev/vx/rdisk/devgrp/vol02
```

Este paso habilita la copia del volumen maestro del cluster primario en el volumen sombra del mismo cluster. El volumen maestro, el volumen sombra y el volumen de mapa de bits de un momento determinado deben estar en el mismo grupo de dispositivos. En este ejemplo, el volumen maestro es `vol01`, el volumen shadow es `vol02` y el volumen de mapa de bits de un momento determinado es `vol03`.

## 8 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -I a \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
```

Este paso asocia la instantánea de un momento determinado con el conjunto de volúmenes de reflejo remoto. El software de Availability Suite garantiza que se tome una instantánea de un momento determinado antes de que pueda haber una replicación de reflejo remoto.

**Pasos siguientes** Vaya a [“Habilitación de la replicación en el cluster secundario”](#) en la página 365.

## ▼ Habilitación de la replicación en el cluster secundario

**Antes de empezar** Complete el procedimiento [“Habilitación de la replicación en el cluster primario”](#) en la página 363.

### 1 Acceda al nodo `nodeC` como superusuario.

### 2 Purgue todas las transacciones.

```
nodeC# lockfs -a -f
```

### 3 Habilite la duplicación por duplicación remota del cluster primario al secundario.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeC# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
```

```
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

El cluster primario detecta la presencia del cluster secundario y comienza la sincronización. Para obtener información sobre el estado de los clusters, consulte el archivo de registro del sistema `/var/adm` de Availability Suite.

#### 4 Habilite la instantánea de un momento determinado independiente.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeC# /usr/sbin/iadm -e ind \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
nodeC# /usr/sbin/iadm -w \
/dev/vx/rdisk/devgrp/vol02
```

#### 5 Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software de Availability Suite:

```
nodeC# /usr/sbin/sndradm -I a \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol02 \
/dev/vx/rdisk/devgrp/vol03
```

**Pasos siguientes** Vaya a [“Ejemplo de cómo efectuar una replicación de datos”](#) en la página 366.

## 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` y `iadm` del software de Availability Suite. Para obtener más información sobre estos comandos, consulte la [documentation](#) de Availability Suite.

Esta sección incluye los procedimientos siguientes:

- [“Replicación por duplicación remota”](#) en la página 366
- [“Instantánea de un momento determinado”](#) en la página 368
- [“Procedimiento para comprobar la configuración de la replicación”](#) en la página 369

### ▼ Replicación por duplicación remota

En este procedimiento, el volumen maestro del disco primario se replica en el volumen maestro del disco secundario. El volumen maestro es `vol01` y el volumen de mapa de bits de reflejo remoto es `vol04`.

#### 1 Obtenga acceso al nodo `nodeA` como superusuario.

**2 Compruebe que el cluster esté en modo de registro.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

**3 Purgue todas las transacciones.**

```
nodeA# lockfs -a -f
```

**4 Repita los pasos del [Paso 1](#) al [Paso 3](#) en el nodo nodeC.****5 Copie el volumen principal del nodo nodeA en el volumen principal del nodo nodeC.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -m lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**6 Espere hasta que la termine replicación y los volúmenes se sincronicen.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**7 Confirme que el cluster esté en modo de replicación.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe en el volumen principal, el software de Availability Suite actualiza el volumen secundario.

**Pasos siguientes** Vaya a [“Instantánea de un momento determinado” en la página 368.](#)

## ▼ Instantánea de un momento determinado

En este procedimiento, la instantánea de un momento determinado se ha utilizado para sincronizar el volumen sombra del cluster primario con el volumen maestro del cluster primario. El volumen maestro es vol01, el volumen de mapa de bits es vol04 y el volumen shadow es vol02.

**Antes de empezar** Realice el procedimiento [“Replicación por duplicación remota” en la página 366.](#)

- 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 **Desactive el recurso que se está ejecutando en nodeA.**

```
nodeA# clresource disable -n nodeA nfs-rs
```

- 3 **Cambie el cluster primario a modo de registro.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

- 4 **Sincronice el volumen sombra del cluster primario con el volumen maestro del cluster primario.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/iiadm -u s /dev/vx/rdisk/devgrp/vol02
nodeA# /usr/sbin/iiadm -w /dev/vx/rdisk/devgrp/vol02
```

- 5 **Sincronice el volumen sombra del cluster secundario con el volumen maestro del cluster secundario.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeC# /usr/sbin/iiadm -u s /dev/vx/rdisk/devgrp/vol02
nodeC# /usr/sbin/iiadm -w /dev/vx/rdisk/devgrp/vol02
```

- 6 **Reinicie la aplicación en el nodo nodeA.**

```
nodeA# clresource enable -n nodeA nfs-rs
```

**7 Vuelva a sincronizar el volumen secundario con el volumen principal.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

**Pasos siguientes** Vaya a [“Procedimiento para comprobar la configuración de la replicación”](#) en la página 369.

**▼ Procedimiento para comprobar la configuración de la replicación**

**Antes de empezar**

Complete el procedimiento [“Instantánea de un momento determinado”](#) en la página 368.

- 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 cluster primario esté en modo de replicación, con la sincronización automática activada.**

Utilice el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe en el volumen principal, el software de Availability Suite actualiza el volumen secundario.

- 3 Si el cluster primario no está en modo de replicación, póngalo en ese modo.**

Utilice el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

- 4 Cree un directorio en un equipo cliente.**

- a. Regístrese en un equipo cliente como superusuario.**

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

**b. Cree un directorio en el equipo cliente.**

```
client-machine# mkdir /dir
```

**5 Monte el directorio en la aplicación del cluster principal y visualice el directorio montado.****a. Monte el directorio en la aplicación del cluster principal.**

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

**b. Visualice el directorio montado.**

```
client-machine# ls /dir
```

**6 Monte el directorio en la aplicación del cluster secundario y visualice el directorio montado.****a. Desmonte el directorio de la aplicación del cluster principal.**

```
client-machine# umount /dir
```

**b. Ponga el grupo de recursos de aplicaciones fuera de línea en el cluster primario.**

```
nodeA# clresource disable -n nodeA nfs-rs
nodeA# clresource disable -n nodeA nfs-dg-rs
nodeA# clresource disable -n nodeA lhost-nfsrg-prim
nodeA# clresourcegroup online -n "" nfs-rg
```

**c. Cambie el cluster primario a modo de registro.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

**d. Asegúrese de que el directorio PathPrefix esté disponible.**

```
nodeC# mount | grep /global/etc
```

**e. Ponga en línea el grupo de recursos de aplicaciones en el cluster secundario.**

```
nodeC# clresourcegroup online -eM -n nodeC nfs-rg
```

**f. Obtenga acceso al equipo cliente como superusuario.**

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

**g. Monte el directorio creado en el [Paso 4](#) en la aplicación del cluster secundario.**

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

**h. Visualice el directorio montado.**

```
client-machine# ls /dir
```

**7 Compruebe que el directorio mostrado en el [Paso 5](#) sea el mismo del [Paso 6](#).****8 Devuelva la aplicación del cluster principal al directorio montado.****a. Ponga sin conexión el grupo de recursos de aplicaciones del cluster secundario.**

```
nodeC# clresource disable -n nodeC nfs-rs
nodeC# clresource disable -n nodeC nfs-dg-rs
nodeC# clresource disable -n nodeC lhost-nfsrg-sec
nodeC# clresourcegroup online -n "" nfs-rg
```

**b. Compruebe que el volumen global esté desmontado del cluster secundario.**

```
nodeC# umount /global/mountpoint
```

**c. Ponga en línea el grupo de recursos de aplicaciones del cluster principal.**

```
nodeA# clresourcegroup online -eM -n nodeA nfs-rg
```

**d. Ponga el cluster principal en modo de replicación.**

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe en el volumen principal, el software de Availability Suite actualiza el volumen secundario.

**Véase también** [“Ejemplo de cómo gestionar una recuperación” en la página 371](#)

## Ejemplo de cómo gestionar una recuperación

En esta sección, se describe cómo provocar una conmutación por error y el modo en que la aplicación se transfiere al cluster secundario. Después de una conmutación por error, actualice las entradas del DNS. Para obtener más información, consulte [“Directrices para la gestión de una recuperación” en la página 345](#).

Esta sección incluye los procedimientos siguientes:

- [“Provocación de un switchover” en la página 372](#)
- [“Actualización de la entrada de DNS” en la página 373](#)

## ▼ Provocación de un switchover

- 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 Cambie el cluster primario a modo de registro.

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 lhost-reprg-sec \
/dev/vx/rdisk/devgrp/vol01 \
/dev/vx/rdisk/devgrp/vol04 ip sync
```

Cuando se escribe en el volumen de datos del disco, se actualiza el volumen de mapa de bits en el mismo grupo de dispositivos. No se produce ninguna replicación.

- 3 Confirme que el cluster principal y el secundario estén en modo de registro, con la sincronización automática desactivada.

- a. En `nodeA`, confirme el modo y la configuración:

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 ->
lhost-reprg-sec:/dev/vx/rdisk/devgrp/vol01
autosync:off, max q writes:4194304,max q fbas:16384,mode:sync,ctag:
devgrp, state: logging
```

- b. En `nodeC`, confirme el modo y la configuración:

Ejecute el comando siguiente para el software de Availability Suite:

```
nodeC# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/vx/rdisk/devgrp/vol01 <-
lhost-reprg-prim:/dev/vx/rdisk/devgrp/vol01
autosync:off, max q writes:4194304,max q fbas:16384,mode:sync,ctag:
devgrp, state: logging
```

Para `nodeA` y `nodeC`, el estado debe ser `logging` y el estado activo de la sincronización automática debe ser `off`.

- 4 Confirme que el cluster secundario esté listo para tomar el control del cluster principal.

```
nodeC# fsck -y /dev/vx/rdisk/devgrp/vol01
```

**5 Realice un switchover al cluster secundario.**

```
nodeC# clresourcegroup switch -n nodeC nfs-rg
```

**Pasos siguientes** Vaya a [“Actualización de la entrada de DNS” en la página 373.](#)

**▼ Actualización de la entrada de DNS**

Si desea una ilustración de cómo asigna el DNS un cliente a un cluster, consulte la [Figura A–6.](#)

**Antes de empezar** Complete el procedimiento [“Provocación de un switchover” en la página 372.](#)

**1 Inicie el comando nsupdate.**

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

**2 Elimine en ambos clusters la asignación de DNS actual entre el nombre de host lógico del grupo de recursos de aplicaciones y la dirección IP del cluster.**

```
> 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 cluster primario en orden inverso.

*ipaddress2rev* Dirección IP del cluster 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 cluster en ambos clusters.**

Asigne el nombre de host lógico principal a la dirección IP del cluster secundario y el nombre sistema lógico secundario a la dirección IP del cluster 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
```

*ipaddress2fwd* Dirección IP del cluster secundario hacia delante.

*ipaddress1fwd* Dirección IP del cluster primario hacia delante.



# Índice

---

## A

- activación de cables de transporte, 208
- actualizar espacio de nombre global, 125
- actualizar manualmente información de DID, 164–165
- adaptadores de transporte, agregación, 202
- adaptadores de transporte, agregar, 205
- adición
  - dispositivos de quórum, 173
  - dispositivos de quórum de servidor de quórum, 177
  - dispositivos de quórum NAS de Sun, 175
  - dispositivos de quórum NAS de Sun ZFS Storage Appliance, 175
  - hosts de SNMP, 262
  - usuarios de SNMP, 263
- administración
  - clusters de zona, 16, 267
  - clusters globales, 16
  - configuración de cluster global, 235–280
  - dispositivos replicados basados en almacenamiento, 97–122
  - dispositivos replicados de EMC SRDF, 109–122
  - dispositivos replicados de Hitachi TrueCopy, 98–109
  - interconexiones de clusters y redes públicas, 199–217
  - IPMP, 199–217
  - nodos sin voto de cluster global, 16
  - sistema de archivos de cluster, 122
- administración de energía, 235
- administrar
  - cluster con herramienta de interfaz gráfica de usuario (GUI), 329–336
  - administrar (*Continuación*)
    - quórum, 169–193
- agregación
  - agregación de dirección de red a un cluster de zona, 268–269
  - cables, adaptadores y conmutadores de transporte, 202
  - nodos a un cluster de zona, 220
  - nodos a un cluster global, 220
  - roles (RBAC), 55
- agregar
  - dispositivos de quórum de disco compartido con conexión directa, 173
  - funciones personalizadas (RBAC), 58
  - grupo de dispositivos, 130, 132–133
  - grupos de dispositivos de Solaris Volume Manager, 132
  - grupos de dispositivos de ZFS, 133
  - nodos, 219–225
  - sistema de archivos de cluster, 154–158
- almacenamiento SATA, admitido como dispositivo de quórum, 172
- anular registro, grupos de dispositivos de Solaris Volume Manager, 135
- anular supervisión, rutas de disco, 163–164
- aplicación
  - parche de no reinicio, 303
  - parche de no reinicio a un cluster de zona, 295
  - parches, 295
- aplicación de parches, en nodos sin voto de cluster global, 298

- aplicación de parches de software de Oracle Solaris Cluster, 293–295
- aplicaciones de conmutación por error para replicación de datos, switchover de afinidad, 342
- aplicaciones de migración tras error para replicación de datos
  - directrices
  - grupos de recursos, 343
- aplicaciones escalables para replicación de datos, 344–345
- archivo `/etc/inet/hosts`, configuración en zonas de IP exclusiva, 225
- archivo `/etc/nsswitch.conf`, modificaciones en las zonas no globales, 224
- archivo `/etc/vfstab`, 49
  - agregar puntos de montaje, 156
  - comprobación de la configuración, 157
- archivo `/var/adm/messages`, 85
- archivo `hosts`, configuración en zonas de IP exclusiva, 225
- archivo `lofi`, desinstalación, 257
- archivo `md.tab`, 21
- archivo `nsswitch.conf`, modificaciones en las zonas no globales, 224
- archivo `ntp.conf.cluster`, 245
- archivo `vfstab`
  - agregar puntos de montaje, 156
  - comprobación de la configuración, 157
- archivos
  - `/etc/vfstab`, 49
  - `md.conf`, 130
  - `md.tab`, 21
  - `ntp.conf.cluster`, 245
  - restauración interactiva, 320
- arrancar
  - cluster de zona, 61–85
  - cluster global, 61–85
  - modo que no sea de cluster, 82
  - nodos, 71–85
  - nodos de cluster de zona, 71–85
  - nodos de cluster global, 71–85
- arrancar en un modo que no sea de cluster, 82
- Asistente para agregar roles administrativos, descripción, 55

- atributos, *Ver* propiedades
- Availability Suite, uso para replicación de datos, 337

## B

- búsqueda
  - ID de nodo para cluster de zona, 238
  - ID de nodo para cluster global, 238
  - nombres de sistemas de archivos, 312

## C

- cables de transporte
  - activación, 208
  - agregación, 202
  - agregar, 205
  - desactivación, 209
- cambiar, nombre de cluster, 236–237
- cambio
  - dirección de servidor de Oracle Solaris Cluster Manager, 332
  - nombres de host privados, 243
  - número de puerto, mediante Common Agent Container, 331
  - protocolo de MIB de eventos de SNMP, 261
- cambio de nombre de nodos
  - en un cluster de zona, 248
  - en un cluster global, 248
- cerrar
  - cluster de zona, 61–85
  - cluster global, 61–85
  - nodos, 71–85
  - nodos de cluster de zona, 71–85
  - nodos de cluster global, 71–85
- cierre, zonas no globales, 71
- claves de seguridad, regenerar, 332
- `clsetup`, agregación de una dirección de red a un cluster de zona, 268–269
- cluster
  - ámbito, 26
  - aplicación de un parche de reinicio, 300
  - autenticación de nodos, 238
  - cambiar nombre, 236–237

- cluster (*Continuación*)
  - configuración de la hora del día, 240
  - copia de seguridad, 21
  - hacer copia de seguridad, 311–319
  - restauración de archivos, 319
- cluster check
  - comando
    - cambios en, 43
- Cluster Control Panel (CCP), 20
- cluster de zona
  - administración, 235–280
  - agregación de direcciones de red, 268–269
  - apagar, 61–85
  - arrancar, 61–85
  - clon, 267
  - definición, 16
  - desplazamiento de ruta de zona, 267
  - eliminación de un sistema de archivos, 267
  - estado de componentes, 31
  - montajes directos admitidos, 270–273
  - preparación para aplicaciones, 267
  - rearrancar, 67
  - validación de configuración, 43
  - visualización de configuración, 35
- cluster global
  - administración, 235–280
  - arrancar, 61–85
  - cerrar, 61–85
  - definición, 16
  - eliminar nodos, 227
  - estado de componentes, 31
  - rearrancar, 67
  - validación de configuración, 43
  - visualización de configuración, 35
- clusters de campus
  - recuperación con replicación de datos basada en
    - almacenamiento, 93
  - replicación de datos basada en
    - almacenamiento, 90–94
- clzonecluster
  - boot, 65–67
  - descripción, 25
  - halt, 61–71
  - comando /usr/cluster/bin/cluster check,
    - comprobación de archivo vfstab, 157
  - comando boot, 65–67
  - comando cconsole, 20, 23
  - comando ccp, 19, 23
  - comando claccess, 19
  - comando cldevice, 19
  - comando cldevicegroup, 19
  - comando clinterconnect, 19
  - comando clnasdevice, 19
  - comando clnode, 265, 266–267
  - comando clnode check, 19
  - comando clquorum, 19
  - comando clreslogicalhostname, 19
  - comando clresource, 19
    - eliminación de recursos y grupos de recursos, 270
  - comando clresourcegroup, 19, 266–267
  - comando clresourcetype, 19
  - comando clressharedaddress, 19
  - comando clsnmp host, 19
  - comando clsnmpmib, 19
  - comando clsnmpuser, 19
  - comando cltelemattribute, 19
  - comando cluster check, 19
    - comprobación de archivo vfstab, 157
  - comando cluster shutdown, 61–71
  - comando clzonecluster, 19
  - comando crlogin, 23
  - comando cssh, 23
  - comando ctelnet, 23
  - comando metaset, 95–97
  - comando netcon, 20
  - comando showrev -p, 26, 27
  - comandos
    - boot, 65–67
    - cconsole, 20, 23
    - ccp, 19, 23
    - claccess, 19
    - cldevice, 19
    - cldevicegroup, 19
    - clinterconnect, 19
    - clnasdevice, 19
    - clnode check, 19
    - clquorum, 19

comandos (*Continuación*)

- clreslogicalhostname, 19
- clresource, 19
- clresourcegroup, 19
- clresourcetype, 19
- clressharedaddress, 19
- clsetup, 19
- clsnmp host, 19
- clsnmp mib, 19
- clsnmp user, 19
- cltelemetryattribute, 19
- cluster check, 19, 21, 43, 49
- cluster scs shutdown, 61–71
- clzonecluster, 19, 61–71
- clzonecluster boot, 65–67
- clzonecluster verify, 43
- crlogin, 23
- cssh, 23
- ctelnet, 23
- metaset, 95–97
- netcon, 20

Common Agent Container

- cambio de número de puerto, 331
- regenerar claves de seguridad, 332

comprobación

- configuración vfstab, 157
- estado de interconexión de cluster, 201
- puntos de montaje globales, 49

comprobar

- configuración de replicación de datos, 369–371
- puntos de montaje globales, 160
- SMF, 223

conexiones seguras con consolas del cluster, 24

configuración

- dispositivos replicados de Hitachi TrueCopy, 99–100
- replicación de datos, 337–373

configuración de cluster de zona

- validación, 43
- visualización, 35

configuración de cluster global

- validación, 43
- visualización, 35

configuración de la hora del cluster, 240

- configuración de límites de carga, en nodos, 266–267
- configuraciones de ejemplo (agrupación en clusters de campus), dos salas, replicación basada en almacenamiento, 90–94
- configurar, funciones (RBAC), 53
- conjunto de procesadores dedicado, configuración, 288
- conmutación para replicación de datos, realización, 371–373
- conmutación por error de afinidad, propiedad de extensión para replicación de datos, 342
- conmutación por error de aplicaciones para replicación de datos
  - directrices
  - gestión de recuperación, 345
  - gestión, 371–373
- conmutadores de transporte, agregación, 202
- conmutadores de transporte, agregar, 205
- conmutar, nodo primario de grupo de dispositivos, 146–147
- conmutar a nodo primario de grupo de dispositivos, 146–147
- consola de administración, 20
- consolas
  - conexión con, 23
  - conexiones seguras, 24
- control de acceso basado en funciones, *Ver* RBAC
- convención de asignación de nombre, grupos de recursos de replications, 342
- convención de denominación, dispositivos de discos sin procesar, 156
- copia de seguridad
  - cluster, 21
  - sistema de archivos, 312
  - sistemas de archivos raíz, 313
- creación de reflejos local, *Ver* replicación basada en almacenamiento
- creación de reflejos remota, *Ver* replicación basada en almacenamiento

**D**

- desactivación de cables de transporte, 209

- desinstalación
    - archivo de dispositivo, 257
    - software de Oracle Solaris Cluster, 255
  - detener
    - cluster de zona, 67
    - cluster global, 67
    - nodos, 71–85
    - nodos de cluster de zona, 71–85
    - nodos de cluster global, 71–85
  - direcciones de red, agregación a un cluster de zona, 268–269
  - direcciones IP, añadir a un servicio de nombres para zonas de direcciones IP exclusivas, 225
  - disco SCSI compartido, admitido como dispositivo de quórum, 172
  - dispositivos
    - globales, 95–168
    - reconfiguración dinámica, 96–97
  - dispositivos de discos sin procesar, convenciones de denominación, 156
  - dispositivos de quórum
    - adición, 173
      - dispositivos de quórum de servidor de quórum, 177
      - dispositivos de quórum NAS de Sun, 175
      - dispositivos de quórum NAS de Sun ZFS Storage Appliance, 175
  - agregar
    - dispositivos de quórum de disco compartido con conexión directa, 173
  - eliminación, 181
  - eliminar, 171
  - enumeración de configuración, 190
  - estado de mantenimiento, poner un dispositivo en, 187
  - estado de mantenimiento, sacar a un dispositivo, 189
  - modificación de listas de nodos, 185
  - quitar el último dispositivo de quórum, 182
  - reconfiguración dinámica de dispositivos, 171
  - reemplazar, 184
  - reparar, 191
  - dispositivos de quórum de disco compartido con conexión directa, agregar, 173
  - dispositivos de quórum de servidor de quórum
    - adición, 177
    - requisitos de instalación, 177
    - resolución de problemas para quitar dispositivos, 182
  - dispositivos de quórum NAS de Sun, adición, 175
  - dispositivos del quórum
    - cambio del tiempo de espera predeterminado, 192–193
    - y replicación basada en almacenamiento, 93
  - dispositivos replicados basados en almacenamiento, administración, 97–122
  - Domain Name System (DNS), actualizar en replicación de datos, 373
  - DR, Ver reconfiguración dinámica
  - duplicación remota, realizar, 366–368
  - duplicaciones, copia de seguridad en línea, 316
- E**
- ejemplos
    - crear un sistema de archivos del cluster, 157
    - ejecutar una comprobación de validación funcional, 47–48
    - listado de comprobaciones de validación interactivas, 47
  - ejemplos de configuración (agrupación en clusters de campus), dos salas, replicación de datos basada en almacenamiento, 90–94
  - eliminación
    - de un cluster de zona, 226
    - dispositivos de quórum, 181
    - hosts de SNMP, 262
    - nodos sin voto en un cluster global, 230
    - recursos y grupos de recursos de un cluster de zona, 270
    - usuarios de SNMP, 264
  - eliminar
    - dispositivos de quórum, 171
    - grupos de dispositivos de Solaris Volume Manager, 135
    - matrices de almacenamiento, 231
    - nodos, 225, 227
    - nodos de todos los grupos de dispositivos, 135

**EMC SRDF**

- administración, 109–122
  - configuración de dispositivos DID, 112–113
  - configuración de grupo de replicación, 110–112
  - Copia adaptable, 91
  - ejemplo de configuración, 115–122
  - modo Dominó, 91
  - prácticas recomendadas, 94
  - recuperación después de un fallo completo de la sala
    - primaria de un cluster de campus, 119–122
  - requisitos, 92
  - restricciones, 92
  - verificación de configuración, 114
- enumeración, configuración de quórum, 190
- enumerar, configuración de grupo de dispositivos, 144
- espacio de nombre
  - global, 95–97, 125
- espacio de nombres, migración, 127
- espacio de nombres de dispositivos globales, migración, 127
- estado
  - componente de cluster de zona, 31
  - componente de cluster global, 31
- estado de mantenimiento
  - nodos, 251
  - poner un dispositivo de quórum en, 187
  - sacar un dispositivo de quórum de, 189

**F**

- fence\_level, *Ver* durante replicación
- fuera de servicio, dispositivo de quórum, 187
- función
  - agregar funciones personalizadas, 58
  - configurar, 53

**G**

- global
  - espacio de nombre, 95–97
  - puntos de montaje, comprobar, 160
- globales
  - dispositivos, 95–168

globales, dispositivos (*Continuación*)

- configurar permisos, 96
  - puntos de montaje, comprobación, 49
- grupo de dispositivos de discos básicos,
  - agregar, 132–133
- grupos de dispositivos
  - agregar, 132
  - configuración de replicación de datos, 350
- disco básico
  - agregar, 132–133
  - eliminar y anular registro, 135
  - enumerar configuración, 144
  - estado de mantenimiento, 147
  - información general de administración, 123
  - modificar propiedades, 140
  - propiedad primaria, 140
- SVM
  - agregar, 130
- grupos de recursos
  - replicación de datos
    - configuración, 342
    - directrices para configurar, 341
    - rol en conmutación por error, 342
- grupos de recursos de aplicaciones
  - configurar para replicación de datos, 359–361
  - directrices, 342
- grupos de recursos de direcciones compartidas para replicación de datos, 344

**H**

- habilitación y deshabilitación de una MIB de eventos de SNMP, 260
- hacer copia de seguridad, cluster, 311–319
- herramienta de administración de GUI, 18, 329–336
  - Oracle Solaris Cluster Manager, 329
- herramienta de administración de línea de comandos, 18
- herramienta User Accounts, descripción, 59
- Hitachi TrueCopy
  - administración, 98–109
  - configuración de dispositivos DID, 101–103
  - configuración de grupo de replicación, 99–100
  - ejemplo de configuración, 104–109

Hitachi TrueCopy (*Continuación*)  
 modo Datos o modo Estado, 91  
 prácticas recomendadas, 94  
 requisitos, 92  
 restricciones, 92  
 verificación de configuración, 103–104  
 Hitachi Universal Replicator, 91  
 prácticas recomendadas, 94  
 requisitos, 92  
 restricciones, 92  
 hosts  
 adición y eliminación de SNMP, 262

## I

imprimir, rutas de disco erróneas, 164  
 información de DID, actualizar  
 manualmente, 164–165  
 información de versión, 26, 27  
 información general, quórum, 169–193  
 iniciar  
 cluster de zona, 65–67  
 cluster global, 65–67  
 nodos, 71–85  
 nodos de cluster global, 71–85  
 Oracle Solaris Cluster Manager, 333  
 inicio, zonas no globales, 71  
 inicio de sesión remoto, 23  
 instantánea, captura, 339  
 instantánea de un momento determinado  
 definición, 339  
 ejecutar, 368–369  
 instantáneas, *Ver* replicación basada en  
 almacenamiento  
 interconexiones de cluster, comprobación de  
 estado, 201  
 interconexiones de clusters  
 administración, 199–217  
 reconfiguración dinámica, 201  
 IPMP  
 administración, 214  
 estado, 34  
 grupos en zonas de direcciones IP exclusivas  
 configurar, 224

## K

/kernel/drv/, archivo md.conf, 130

## L

límites de carga  
 configuración en nodos, 265, 266–267  
 propiedad `concentrate_load`, 265  
 propiedad `preemption_mode`, 265

## M

mantenimiento, dispositivo de quórum, 187  
 mapa de bits  
 instantánea de un momento determinado, 339  
 replicación por duplicación remota, 338  
 matrices de almacenamiento, eliminar, 231  
 mensajes de error  
 archivo /var/adm/messages, 85  
 eliminar nodos, 233–234  
 MIB  
 cambio de protocolo de eventos de SNMP, 261  
 habilitación y deshabilitación de eventos de  
 SNMP, 260  
 MIB de eventos  
 cambio de protocolo de SNMP, 261  
 habilitación y deshabilitación de SNMP, 260  
 migración, espacio de nombres de dispositivos  
 globales, 127  
 modificación, listas de nodos de dispositivo de  
 quórum, 185  
 modificar  
 nodos primarios, 146–147  
 propiedad `numsecondaries`, 142  
 propiedades, 140  
 usuarios (RBAC), 59  
 montaje directo, exportación de un sistema de archivos  
 a un cluster de zona, 270–273  
 montaje en bucle de retorno, exportación de un sistema  
 de archivos a un cluster de zona, 270–273

**N**

NAS de Network Appliance, admitido como dispositivo de quórum, 172

NAS de Sun, admitido como dispositivo de quórum, 172

Network File System (NFS), configurar sistemas de archivos de aplicaciones para replicación de datos, 353–354

**nodos**

- agregar, 219–225

- aplicación de un parche de reinicio a un cluster global, 295

- arrancar, 71–85

- autenticación, 238

- búsqueda de ID, 238

- cambio de nombre en un cluster de zona, 248

- cambio de nombre en un cluster global, 248

- cerrar, 71–85

- cómo poner un nodo en estado de mantenimiento, 251

- conexión con, 23

- configuración de límites de carga, 266–267

- eliminación de nodos sin voto de un cluster global, 230

- eliminación de un cluster de zona, 226

- eliminar

  - mensajes de error, 233–234

- eliminar de grupos de dispositivos, 135

- eliminar nodos de un cluster global, 227

- primarios, 96–97, 140

- secundarios, 140

nodos con voto de cluster global, administración de sistema de archivos de cluster, 122

**nodos de cluster de zona**

- cerrar, 71–85

- especificando dirección IP y NIC, 219–225

- iniciar, 71–85

- reiniciar, 78–82

**nodos de cluster global**

- arrancar, 71–85

- cerrar, 71–85

- reiniciar, 78–82

nodos de votación de cluster global, recursos compartidos de CPU, 283

**nodos sin voto de cluster global**

- administración, 16

- administración de sistema de archivos de cluster, 122

- agregación de un nombre de host privado, 246

- aplicación d parches, 298

- cambio de nombre de host privado, 246

- cierre y reinicio, 71

- nombre de host privado, supresión, 248

- recursos compartidos de CPU, 285, 288

**nombres de host privados**

- asignar a zonas, 223

- cambio, 243

- cambio en nodos sin voto de cluster global, 246

- nodos sin voto de cluster global, 246

- supresión en nodos sin voto de cluster global, 248

número de puerto, cambio mediante Common Agent Container, 331

**O**

opciones de montaje para sistemas de archivos de cluster, requisitos, 156

Oracle Solaris Cluster Manager, 18, 329

- cambio de dirección de servidor, 332

- funciones de RBAC, configurar, 330

- iniciar, 333

**P****parches**

- aplicación a cluster y firmware, 300

- aplicación de un parche de no reinicio, 303

- aplicación de un parche de no reinicio a un cluster de zona, 295

- aplicación de un parche de reinicio a un cluster global, 295

- sugerencias, 294

perfiles, derechos de RBAC, 54–55

perfiles de derechos, RBAC, 54–55

permisos, dispositivo global, 96

planificador por reparto equitativo, configuración de recursos compartidos de CPU, 282

- prácticas recomendadas
  - EMC SRDF, 94
  - Hitachi TrueCopy, 94
  - Hitachi Universal Replicator, 94
  - replicación de datos basada en almacenamiento, 94
- PROM OpenBoot (OBP), 242
- propiedad autoboot, 223
- propiedad failback, 140
- propiedad numsecondaries, 142
- propiedad primaria de grupos de dispositivos, 140
- propiedades
  - failback, 140
  - numsecondaries, 142
  - preferenced, 140
- propiedades de extensión para replicación de datos
  - recurso de aplicación, 360, 362
  - recurso de replicación, 356, 357
- puntos de montaje
  - globales, 49
  - modificación del archivo `/etc/vfstab`, 156

## Q

- quitar
  - cables, adaptadores y conmutadores de transporte, 205
  - sistema de archivos de cluster, 158–160
  - último dispositivo de quórum, 182
- quórum
  - administración, 169–193
  - información general, 169–193

## R

- RBAC, 53–60
  - Oracle Solaris Cluster Manager, 330
  - para nodos con voto de cluster global, 54
  - para nodos sin voto, 54
  - perfiles de derechos (descripción), 54–55
  - tareas
    - agregación de roles, 55
    - agregar funciones personalizadas, 58
    - configurar, 53

- RBAC, tareas (*Continuación*)
  - modificar usuarios, 59
  - usar, 53
- realizar copia de seguridad, duplicaciones en línea, 316
- rearrancar, cluster de zona, 67
- reconfiguración dinámica, 96–97
  - dispositivos de quórum, 171
  - interconexiones de clusters, 201
  - interfaces de red pública, 216
- recuperación, cluster con replicación de datos basada en almacenamiento, 93
- recurso de nombre de host lógico, rol en recuperación de replicación de datos, 342
- recursos
  - eliminación, 270
  - visualización de tipos configurados, 30
- recursos compartidos de CPU
  - configurar, 281
  - controlar, 281
  - nodos de votación de cluster global, 283
  - nodos sin voto de cluster global, 285
  - nodos sin voto de cluster global, conjunto de procesadores dedicado, 288
- red pública
  - administración, 199–217
  - reconfiguración dinámica, 216
- reemplazar dispositivos de quórum, 184
- regenerar, claves de seguridad, 332
- reiniciar
  - cluster global, 67
  - nodos de cluster de zona, 78–82
  - nodos de cluster global, 78–82
- remoto, inicio de sesión, 23
- reparación del archivo `/var/adm/messages`
  - completo, 85
- reparar, dispositivo de quórum, 191
- replicación, *Ver* replicación de datos
- replicación, basada en almacenamiento, 90–94
- replicación asíncrona de datos, 91
- replicación de datos, 87–94
  - actualizar una entrada DNS, 373
  - asíncrona, 339
  - basada en almacenamiento, 88, 90–94
  - basada en host, 88

replicación de datos (*Continuación*)

- comprobar configuración, 369–371
  - configuración
    - grupos de dispositivos, 350
    - switchover de afinidad, 342, 356
  - configuración de ejemplo, 346
  - configurar
    - grupos de recursos de aplicaciones
      - NFS, 359–361
    - sistemas de archivos para una aplicación
      - NFS, 353–354
  - definición, 88–89
  - directrices
    - configurar grupos de recursos, 341
    - gestión de conmutación, 345
    - gestión de recuperación, 345
  - duplicación remota, 338, 366–368
  - ejemplo, 366–371
  - gestión de una recuperación, 371–373
  - grupos de recursos
    - aplicación, 342
    - aplicaciones de migración tras error, 343
    - aplicaciones escalables, 344–345
    - configuración, 342
    - convención de asignación de nombre, 342
    - crear, 356–357
    - dirección compartida, 344
  - habilitación, 363–366
  - hardware y software necesarios, 348
  - instantánea de un momento determinado, 339, 368–369
  - introducción, 338
  - sincrónica, 339
- replicación de datos asincrónica, 339
- replicación de datos basada en almacenamiento, 90–94
- definición, 88
  - prácticas recomendadas, 94
  - recuperación, 93
  - requisitos, 92
  - restricciones, 92
  - y dispositivos del quórum, 93
- replicación de datos basada en host
- definición, 88
  - ejemplo, 337–373

- replicación de datos sincrónica, 339
- replicación por duplicación remota, definición, 338
- replicación remota, *Ver* replicación basada en almacenamiento
- replicación síncrona de datos, 91
- restauración
  - archivos de cluster, 319
  - archivos de manera interactiva, 320
  - sistema de archivos raíz, 320
    - desde metadispositivo, 323
    - desde volumen, 323
- rol, agregación de roles, 55
- ruta de disco
  - anular supervisión, 163–164
  - resolver error de estado, 164–165
  - supervisar, 95–168
    - imprimir rutas de disco erróneas, 164
- ruta de zona, desplazamiento, 267
- rutas de disco compartido
  - habilitar rearranque automático, 167–168
  - inhabilitar reinicio automático, 168
  - supervisar, 160–168

**S**

- SATA, 173
- secundarios
  - establecer número, 142
  - número predeterminado, 140
- servicio de nombres, añadir asignaciones de direcciones
  - IP para zonas de direcciones IP exclusivas, 225
- servicios multiusuario, comprobar, 223
- servidor de quórum de Oracle Solaris Cluster, admitido
  - como dispositivo de quórum, 172
- servidores de quórum, *Ver* dispositivos de quórum de servidor de quórum
- shell seguro, 23, 24
- sistema de archivos
  - aplicación NFS
    - configurar para replicación de datos, 353–354
  - búsqueda de nombres, 312
  - copia de seguridad, 312
  - eliminación desde un cluster de zona, 267

- sistema de archivos (*Continuación*)
    - restauración de raíz
      - descripción, 320
      - desde metadispositivo, 323
      - desde volumen, 323
  - sistema de archivos de cluster, 95–168
    - administración, 122
    - agregar, 154–158
    - nodos con voto de cluster global, 122
    - nodos sin voto de cluster global, 122
    - quitar, 158–160
  - sistema DNS, directrices para actualizar, 345
  - sistema operativo Oracle Solaris
    - comando `svcadm`, 243
    - control CPU, 281
    - definición de cluster de zona, 15
    - definición de cluster global, 15
    - instrucciones especiales para iniciar nodos, 75–78
    - instrucciones especiales para reiniciar un nodo, 78–82
    - replicación basada en host, 89
    - tareas administrativas para un cluster global, 16
  - sistema operativo Solaris
    - Ver también* sistema operativo Oracle Solaris
  - sistemas de archivos de cluster, opciones de
    - montaje, 156
  - sistemas de archivos del cluster, comprobación de la
    - configuración, 157
  - sistemas de archivos globales, *Ver* sistemas de archivos del cluster
  - SMF, comprobar servicios en línea, 223
  - SNMP
    - adición de usuarios, 263
    - cambio de protocolo, 261
    - deshabilitación de hosts, 262
    - eliminación de usuarios, 264
    - habilitación de hosts, 262
    - habilitación y deshabilitación de una MIB de eventos, 260
  - software Oracle Solaris, SMF, 223
  - Solaris Volume Manager, nombres de dispositivos de
    - discos sin procesar, 156
  - SRDF
    - Ver* EMC SRDF
  - ssh, 24
  - Sun ZFS Storage Appliance
    - adición como dispositivo de quórum, 175
    - admitido como dispositivo de quórum, 172
  - supervisar
    - rutas de disco, 161–163
    - rutas de disco compartido, 167–168
  - switchback, directrices para realizar en replicación de
    - datos, 346
  - switchover de afinidad, configuración para replicación
    - de datos, 356
  - switchover para replicación de datos, switchover de
    - afinidad, 342
- T**
- tiempo de espera, cambio del valor predeterminado de
    - un dispositivo del quórum, 192–193
  - tipos de dispositivos de quórum, lista de tipos
    - admitidos, 172
  - tipos de dispositivos de quórum admitidos, 172
  - tolerancia a desastres, definición, 338
  - transporte, adaptadores, 205
  - transporte, cables, 205
  - transporte, conmutadores, 205
  - TrueCopy
    - Ver* Hitachi TrueCopy
- U**
- último dispositivo de quórum, quitar, 182
  - Universal Replicator, Hitachi Universal Replicator, 91
  - usar, funciones (RBAC), 53
  - `/usr/cluster/bin/clresource`, eliminación de
    - grupos de recursos, 270
  - usuarios
    - adición de SNMP, 263
    - eliminación de SNMP, 264
    - modificar propiedades, 59
  - utilidad `clsetup`, 18, 19, 25

## V

visualización de recursos configurados, 30  
volume, *Ver* replicación basada en almacenamiento

## Z

### ZFS

- agregar grupos de dispositivos, 133
- eliminación de un sistema de archivos, 270–273
- replicación, 133
- restricciones para sistemas de archivos root, 122
- ZFS Storage Appliance, *Ver* dispositivos de quórum de Sun ZFS Storage Appliance
- zonas de direcciones IP compartidas, *Ver* zonas de Oracle Solaris
- zonas de direcciones IP exclusivas, *Ver* zonas de Oracle Solaris
- zonas de Oracle Solaris
  - modificaciones en el archivo `nsswitch.conf`, 224
  - propiedad `autoboot`, 223
  - zonas de direcciones IP compartidas, 223
  - zonas de direcciones IP exclusivas
    - configurar grupos IPMP, 224
  - zonas de IP exclusiva
    - configuración de archivo `hosts`, 225