

Guía de administración del sistema de Oracle® Solaris Cluster 4.3



Referencia: E62253
Julio de 2016

Referencia: E62253

Copyright © 2000, 2016, Oracle y/o sus filiales. Todos los derechos reservados.

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

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

Si este software o la documentación relacionada se entrega al Gobierno de EE.UU. o a cualquier entidad que adquiera las licencias en nombre del Gobierno de EE.UU. entonces 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 proporcionar acceso a, o información sobre contenidos, productos o servicios de terceros. Oracle Corporation o sus filiales no son responsables y por ende desconocen cualquier tipo de garantía sobre el contenido, los productos o los servicios de terceros a menos que se indique otra cosa en un acuerdo en vigor formalizado entre Ud. y Oracle. Oracle Corporation y sus filiales no serán responsables frente a cualesquiera pérdidas, costos o daños en los que se incurra como consecuencia de su acceso o su uso de contenidos, productos o servicios de terceros a menos que se indique otra cosa en un acuerdo en vigor formalizado entre Ud. y Oracle.

Accesibilidad a la documentación

Para obtener información acerca del compromiso de Oracle con la accesibilidad, visite el sitio web del Programa de Accesibilidad de Oracle en <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=docacc>.

Acceso a Oracle Support

Los clientes de Oracle que hayan adquirido servicios de soporte disponen de acceso a soporte electrónico a través de My Oracle Support. Para obtener información, visite <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=info> o <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=trs> si tiene problemas de audición.

Contenido

Uso de esta documentación	21
 1 Introducción a la administración de Oracle Solaris Cluster	23
Visión general de la administración de Oracle Solaris Cluster	24
Trabajo con un cluster de zona	24
Restricciones de las funciones del sistema operativo Oracle Solaris	26
Herramientas de administración	26
Interfaz de explorador de Oracle Solaris Cluster Manager	26
Interfaz de línea de comandos	27
Preparación para administrar el cluster	29
Documentación de la configuración de hardware de Oracle Solaris Cluster	29
Uso de la consola de administración	29
Copia de seguridad del cluster	30
Administración del cluster	30
Inicio de sesión de manera remota en el cluster	32
Conexión segura a las consolas del cluster	32
▼ Obtención de acceso a las utilidades de configuración del cluster	33
▼ Cómo visualizar la información de versión de Oracle Solaris Cluster	34
▼ Visualización de tipos de recursos configurados, grupos de recursos y recursos	36
▼ Comprobación del estado de los componentes del cluster	38
▼ Comprobación del estado de la red pública	41
▼ Visualización de la configuración del cluster	42
▼ Validación de una configuración básica de cluster	51
▼ Comprobación de los puntos de montaje globales	57
▼ Cómo ver el contenido de los logs de comandos de Oracle Solaris Cluster	59
 2 Oracle Solaris Cluster y RBAC	63
Configuración y uso de RBAC con Oracle Solaris Cluster	63
Perfiles de derechos de RBAC de Oracle Solaris Cluster	63

Creación y asignación de un rol RBAC con un perfil de derechos de gestión de Oracle Solaris Cluster	65
▼ Creación de una función desde la línea de comandos	65
Modificación de las propiedades de RBAC de un usuario	66
▼ Modificación de las propiedades de RBAC de un usuario con la herramienta de cuentas de usuario	67
▼ Modificación de las propiedades de RBAC de un usuario desde la línea de comandos	68
3 Cierre y arranque de un cluster	69
Visión general sobre el cierre y el inicio de un cluster	69
▼ Cierre de un cluster	71
▼ Inicio de un cluster	73
▼ Reinicio de un cluster	77
Cierre e inicio de un solo nodo de un cluster	84
▼ Cierre de un nodo	85
▼ Inicio de un nodo	89
▼ Reinicio de un nodo	92
▼ Reinicio de un nodo en el modo que no es de cluster	95
Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad	98
▼ Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad	98
4 Métodos de replicación de datos	101
Comprensión de la replicación de datos	102
Métodos admitidos de replicación de datos	103
Uso de la replicación de datos basada en almacenamiento en un cluster de campus	104
Requisitos y restricciones del uso de la replicación de datos basada en almacenamiento en un cluster de campus	106
Problemas de recuperación manual en el uso de la replicación de datos basada en almacenamiento en un cluster de campus	107
Mejores prácticas para utilizar la replicación de datos basada en almacenamiento	108
5 Administración de dispositivos globales, supervisión de rutas de disco y sistemas de archivos de cluster	109
Información general sobre la administración de dispositivos globales y el espacio de nombre global	109

Permisos de dispositivos globales en Solaris Volume Manager	110
Reconfiguración dinámica con dispositivos globales	110
Configuración y administración de dispositivos replicados basados en el almacenamiento	111
Administración de dispositivos replicados de EMC Symmetrix Remote Data Facility	112
Información general sobre la administración de sistemas de archivos de clusters	125
Restricciones del sistema de archivos de cluster	125
Administración de grupos de dispositivos	125
▼ Actualización del espacio de nombre de dispositivos globales	127
▼ Cómo cambiar el tamaño de un dispositivo lofi que se utiliza para el espacio de nombres de dispositivos globales	128
Migración del espacio de nombre de dispositivos globales	129
▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de lofi	130
▼ Cómo migrar el espacio de nombres de dispositivos globales de un dispositivo lofi a una partición dedicada	131
Adición y registro de grupos de dispositivos	132
▼ Cómo agregar y registrar un grupo de dispositivos (Solaris Volume Manager)	133
▼ Adición y registro de un grupo de dispositivos (disco básico)	135
▼ Adición y registro de un grupo de dispositivos replicado (ZFS)	136
▼ Cómo configurar una agrupación de almacenamiento ZFS local sin HAStoragePlus	137
Mantenimiento de grupos de dispositivos	139
Cómo eliminar un grupo de dispositivos y anular su registro (Solaris Volume Manager)	139
▼ Eliminación de un nodo de todos los grupos de dispositivos	139
▼ Cómo eliminar un nodo de un grupo de dispositivos (Solaris Volume Manager)	141
▼ Eliminación de un nodo de un grupo de dispositivos de discos básicos	143
▼ Cambio de propiedades de los grupos de dispositivos	145
▼ Establecimiento del número de secundarios para un grupo de dispositivos	146
▼ Enumeración de la configuración de grupos de dispositivos	149
▼ Conmutación al nodo primario de un grupo de dispositivos	150
▼ Colocación de un grupo de dispositivos en estado de mantenimiento	151
Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento	153
▼ Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento	154

▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento	155
▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento	155
▼ Modificación del protocolo de protección en un solo dispositivo de almacenamiento	157
Administración de sistemas de archivos de cluster	159
▼ Adición de un sistema de archivos de cluster	159
▼ Eliminación de un sistema de archivos de cluster	162
▼ Comprobación de montajes globales en un cluster	164
Administración de la supervisión de rutas de disco	165
▼ Supervisión de una ruta de disco	166
▼ Anulación de la supervisión de una ruta de disco	167
▼ Impresión de rutas de disco erróneas	168
▼ Resolución de un error de estado de ruta de disco	169
▼ Supervisión de rutas de disco desde un archivo	169
▼ Activación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	171
▼ Desactivación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	172
6 Administración de quórum	173
Administración de dispositivos de quórum	173
Reconfiguración dinámica con dispositivos de quórum	175
Adición de un dispositivo de quórum	175
Eliminación o sustitución de un dispositivo de quórum	182
Mantenimiento de dispositivos de quórum	186
Cambio del timeout 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	194
▼ 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	200
Reconfiguración dinámica con interconexiones de clusters	201
▼ Comprobación del estado de la interconexión de cluster	202

▼ Adición de cables de transporte de clúster, adaptadores o conmutadores de transporte	203
▼ Eliminación de cables de transporte de clúster, adaptadores de transporte y conmutadores de transporte	205
▼ Cómo habilitar un cable de transporte de clúster	207
▼ Cómo deshabilitar un cable de transporte de clúster	208
▼ Determinación del número de instancia de un adaptador de transporte	210
▼ Modificación de la dirección de red privada o del intervalo de direcciones de un cluster	211
Resolución de problemas de interconexiones de cluster	213
Administración de redes públicas	215
Administración de grupos de varias rutas de red IP en un cluster	215
Reconfiguración dinámica con interfaces de red pública	217
8 Administración de nodos del cluster	219
Agregación de un nodo a un cluster o cluster de zona	219
▼ Cómo agregar un nodo a un cluster o cluster de zona existente	220
Restauración de nodos del cluster	222
▼ Cómo restaurar un nodo desde el archivo unificado	223
Eliminación de nodos de un cluster	227
▼ Eliminación de un nodo de un cluster de zona	228
▼ Eliminación de un nodo de la configuración de software del cluster	229
▼ Eliminación de la conectividad entre una matriz y un único nodo, en un cluster con conectividad superior a dos nodos	232
▼ Corrección de mensajes de error	234
9 Administración del cluster	237
Información general sobre la administración del cluster	237
▼ Cambio del nombre del cluster	238
▼ Asignación de un ID de nodo a un nombre de nodo	240
▼ Cómo trabajar con autenticación para nodos de cluster nuevos	240
▼ Restablecimiento de la hora del día en un cluster	242
▼ SPARC: Visualización de la PROM OpenBoot en un nodo	244
▼ Cómo cambiar un nombre de host privado de nodo	245
▼ 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	249
▼ Cómo poner un nodo en estado de mantenimiento	250
▼ Cómo sacar un nodo del estado de mantenimiento	252

▼ Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster	254
Resolución de problemas de desinstalación de nodos	257
Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster	258
Configuración de límites de carga	265
Cómo realizar tareas administrativas del cluster de zona	268
▼ Cómo configurar un cluster de zona desde el archivo unificado	269
▼ Cómo instalar un cluster de zona desde el archivo unificado	270
▼ Cómo agregar una dirección de red a un cluster de zona	272
▼ Cómo eliminar un cluster de zona	273
▼ Cómo eliminar un sistema de archivos de un cluster de zona	274
▼ Cómo eliminar un dispositivo de almacenamiento de un cluster de zona	277
Procedimientos de resolución de problemas para realizar pruebas	279
Ejecución de una aplicación fuera del cluster global	279
Restauración de un conjunto de discos dañado	282
10 Configuración del control del uso de la CPU	287
Introducción al control de la CPU	287
Elección de un escenario	287
Planificador por reparto equitativo	288
Configuración del control de CPU	288
▼ Cómo controlar el uso de CPU en un nodo de cluster global	288
11 Actualización de software	291
Visión general de la actualización de software de Oracle Solaris Cluster	291
Actualización de software de Oracle Solaris Cluster	292
Consejos de actualización	293
Actualización a un paquete específico	294
Actualización y aplicación de parches de un cluster de zona	295
Actualización de un servidor de quórum o servidor de instalación AI	297
Desinstalación de un paquete	298
▼ Cómo desinstalar un paquete	298
▼ Cómo desinstalar paquetes de servidor de quórum o servidor de instalación AI	299
12 Copias de seguridad y restauraciones de clusters	301
Copia de seguridad de un cluster	301
▼ Copias de seguridad en línea para reflejos (Solaris Volume Manager)	301

▼ Copias de seguridad de la configuración del cluster	303
Restauración de archivos de cluster	304
▼ Cómo restaurar el sistema de archivos raíz ZFS (/) (Solaris Volume Manager)	305
13 Uso de la interfaz de explorador de Oracle Solaris Cluster Manager	309
Visión general de Oracle Solaris Cluster Manager	309
Acceso al software de Oracle Solaris Cluster Manager	310
▼ Cómo acceder a Oracle Solaris Cluster Manager	310
Resolución de problemas de Oracle Solaris Cluster Manager	311
Configuración de soporte de accesibilidad para Oracle Solaris Cluster Manager	313
Uso de la topología para supervisar el cluster	314
▼ Cómo usar la topología para supervisar y actualizar el cluster	314
A Ejemplo de despliegue: configuración de replicación de datos basada en host entre clusters con el software Availability Suite	317
Comprensión del software Availability Suite en un cluster	318
Métodos de replicación de datos utilizados por Availability Suite	318
Replicación de reflejo remoto	318
Instantánea de un momento determinado	319
Replicación en la configuración de ejemplo	320
Directrices para la configuración de replicación de datos basada en host entre clusters	321
Configuración de grupos de recursos de replications	322
Configuración de grupos de recursos de aplicaciones	323
Directrices para la gestión de una recuperación	327
Mapa de tareas: ejemplo de configuración de una replicación de datos	329
Conexión e instalación de clusters	329
Ejemplo de configuración de grupos de dispositivos y de recursos	331
▼ Configuración de un grupo de dispositivos en el cluster primario	332
▼ Configuración de un grupo de dispositivos en el cluster secundario	334
▼ Configuración de sistemas de archivos en el cluster primario para la aplicación NFS	335
▼ Configuración del sistema de archivos en el cluster secundario para la aplicación NFS	336
▼ Creación de un grupo de recursos de replications en el cluster primario ...	337
▼ Creación de un grupo de recursos de replications en el cluster secundario	339
▼ Creación de un grupo de recursos de aplicaciones NFS en el cluster primario	341

▼ Creación de un grupo de recursos de aplicaciones NFS en el cluster secundario	344
Ejemplo de cómo activar la replicación de datos	347
▼ Activación de la replicación en el cluster primario	347
▼ Activación de la replicación en el cluster secundario	349
Ejemplo de cómo efectuar una replicación de datos	350
▼ Replicación por duplicación remota	350
▼ Instantánea de un momento determinado	352
▼ Procedimiento para comprobar la configuración de la replicación	353
Ejemplo de cómo gestionar una recuperación	356
▼ Actualización de la entrada de DNS	356
 Índice	 359

Lista de figuras

FIGURA 1	Configuración de dos salas con replicación de datos basada en almacenamiento	105
FIGURA 2	Replicación de reflejo remoto	319
FIGURA 3	Instantánea de un momento determinado	320
FIGURA 4	Replicación en la configuración de ejemplo	321
FIGURA 5	Configuración de grupos de recursos en una aplicación de migración tras error	325
FIGURA 6	Configuración de grupos de recursos en una aplicación escalable	327
FIGURA 7	Asignación del DNS de un cliente a un cluster	328
FIGURA 8	Ejemplo de configuración del cluster	330

Lista de tablas

TABLA 1	Servicios de Oracle Solaris Cluster	26
TABLA 2	Herramientas de administración de Oracle Solaris Cluster	31
TABLA 3	Lista de tareas: cerrar e iniciar un cluster	70
TABLA 4	Mapa de tareas: cierre e inicio de un nodo	85
TABLA 5	Mapa de tareas: reconfiguración dinámica con dispositivos de disco y cinta	111
TABLA 6	Mapa de tareas: Administración de dispositivo replicado basado en almacenamiento de EMC SRDF	112
TABLA 7	Mapa de tareas: administrar grupos de dispositivos	126
TABLA 8	Mapa de tareas: administrar sistemas de archivos de cluster	159
TABLA 9	Mapa de tareas: administrar la supervisión de rutas de disco	165
TABLA 10	Lista de tareas: administrar el quórum	174
TABLA 11	Mapa de tareas: reconfiguración dinámica con dispositivos de quórum	175
TABLA 12	Lista de tareas: administrar la interconexión de cluster	200
TABLA 13	Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas	201
TABLA 14	Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas	218
TABLA 15	Mapa de tareas: agregar un nodo a un cluster global o a un cluster de zona existente	220
TABLA 16	Mapa de tareas: eliminar un nodo	227
TABLA 17	Lista de tareas: administrar el cluster	237
TABLA 18	Mapa de tareas: creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster	259
TABLA 19	Otras tareas del cluster de zona	268
TABLA 20	Actualización de software de Oracle Solaris Cluster	293
TABLA 21	Mapa de tareas: hacer una copia de seguridad de los archivos del cluster	301
TABLA 22	Mapa de tareas: restaurar los archivos del cluster	304
TABLA 23	Mapa de tareas: ejemplo de configuración de una replicación de datos	329
TABLA 24	Hardware y software necesarios	330

TABLA 25	Resumen de los grupos y de los recursos en la configuración de ejemplo	332
-----------------	--	-----

Lista de ejemplos

EJEMPLO 1	Visualización de la información de versión de Oracle Solaris Cluster	34
EJEMPLO 2	Visualización de tipos de recursos, grupos de recursos y recursos configurados	37
EJEMPLO 3	Comprobación del estado de los componentes del cluster	39
EJEMPLO 4	Comprobación del estado de la red pública	42
EJEMPLO 5	Visualización de la configuración del cluster global	43
EJEMPLO 6	Visualización de la información del cluster de zona	50
EJEMPLO 7	Comprobación de la configuración del cluster global con resultado correcto en todas las comprobaciones básicas	54
EJEMPLO 8	Listado de comprobaciones de validación interactivas	55
EJEMPLO 9	Ejecución de una comprobación de validación funcional	55
EJEMPLO 10	Comprobación de la configuración del cluster global con una comprobación con resultado no satisfactorio	56
EJEMPLO 11	Comprobación de puntos de montaje globales	57
EJEMPLO 12	Visualización del contenido de los logs de comandos de Oracle Solaris Cluster	60
EJEMPLO 13	Creación de una función de operador personalizado con el comando <code>smrole</code>	66
EJEMPLO 14	Cierre de un cluster de zona	72
EJEMPLO 15	SPARC: Cierre de un cluster global	72
EJEMPLO 16	x86: Cierre de un cluster global	73
EJEMPLO 17	SPARC: Inicio de un cluster global	74
EJEMPLO 18	x86: Inicio de un cluster	75
EJEMPLO 19	Reinicio de un cluster de zona	79
EJEMPLO 20	SPARC: Reinicio de un cluster global	80
EJEMPLO 21	x86: Reinicio de un cluster	81
EJEMPLO 22	SPARC: Cierre de nodos del cluster global	87
EJEMPLO 23	x86: Cierre de nodos del cluster global	87
EJEMPLO 24	Cierre de un nodo de un cluster de zona	88
EJEMPLO 25	SPARC: Inicio de un nodo del cluster global	90
EJEMPLO 26	x86: Inicio de un nodo de cluster	91

EJEMPLO 27	SPARC: Reinicio de un nodo del cluster global	94
EJEMPLO 28	Reinicio de un nodo del cluster global	95
EJEMPLO 29	SPARC: Arranque de un nodo del cluster global en un modo que no sea de cluster	97
EJEMPLO 30	Creación de los pares de réplicas	118
EJEMPLO 31	Comprobación de la configuración de la replicación de datos	118
EJEMPLO 32	Visualización de DID correspondientes a los discos utilizados	120
EJEMPLO 33	Combinación de instancias DID	122
EJEMPLO 34	Visualización de DID combinados	122
EJEMPLO 35	Recuperación manual de datos de EMC SRDF después de una conmutación por error del sitio principal	124
EJEMPLO 36	Actualización del espacio de nombre de dispositivos globales	127
EJEMPLO 37	Agregación de un grupo de dispositivos de Solaris Volume Manager	135
EJEMPLO 38	Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)	142
EJEMPLO 39	Eliminación de un nodo de un grupo de dispositivos básicos	144
EJEMPLO 40	Modificación de las propiedades de grupos de dispositivos	146
EJEMPLO 41	Modificación del número deseado de nodos secundarios (Solaris Volume Manager)	148
EJEMPLO 42	Definición del número de nodos secundarios con el valor predeterminado	148
EJEMPLO 43	Enumeración del estado de todos los grupos de dispositivos	150
EJEMPLO 44	Enumeración de la configuración de un determinado grupo de dispositivos	150
EJEMPLO 45	Conmutación del nodo primario de un grupo de dispositivos	151
EJEMPLO 46	Colocación de un grupo de dispositivos en estado de mantenimiento	153
EJEMPLO 47	Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento	154
EJEMPLO 48	Visualización del protocolo SCSI de un solo dispositivo	155
EJEMPLO 49	Establecimiento de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento	157
EJEMPLO 50	Configuración del protocolo de protección de un solo dispositivo	159
EJEMPLO 51	Eliminación de un sistema de archivos de cluster	164
EJEMPLO 52	Supervisión de una ruta de disco en un solo nodo	166
EJEMPLO 53	Supervisión de una ruta de disco en todos los nodos	167
EJEMPLO 54	Relectura de la configuración de disco desde CCR	167
EJEMPLO 55	Anulación de la supervisión de una ruta de disco	168
EJEMPLO 56	Impresión de rutas de disco erróneas	168
EJEMPLO 57	Supervisión de rutas de disco desde un archivo	170
EJEMPLO 58	Eliminación del último dispositivo de quórum	185

EJEMPLO 59	Colocación de un dispositivo de quórum en estado de mantenimiento	188
EJEMPLO 60	Restablecimiento del número de votos de quórum (dispositivo de quórum)	190
EJEMPLO 61	Enumeración en una lista la configuración de quórum	191
EJEMPLO 62	Inicio de todos los servidores de quórum configurados	194
EJEMPLO 63	Inicio de un servidor de quórum en concreto	194
EJEMPLO 64	Detención de todos los servidores de quórum configurados	195
EJEMPLO 65	Detención de un servidor de quórum específico	195
EJEMPLO 66	Visualización de la configuración de un servidor de quórum	196
EJEMPLO 67	Visualización de la configuración de varios servidores de quórum	196
EJEMPLO 68	Visualización de la configuración de todos los servidores de quórum en ejecución	196
EJEMPLO 69	Limpieza de la información obsoleta de clusters de la configuración del servidor de quórum	198
EJEMPLO 70	Comprobación del estado de la interconexión de cluster	202
EJEMPLO 71	Comprobación de la agregación de un cable de transporte, un adaptador de transporte o un conmutador de transporte de cluster	204
EJEMPLO 72	Comprobación de la eliminación de un cable de transporte, un adaptador de transporte o un conmutador de transporte	207
EJEMPLO 73	Eliminación de un nodo de la configuración de software del cluster	231
EJEMPLO 74	Cambio del nombre del cluster	239
EJEMPLO 75	Asignación de ID del nodo al nombre del nodo	240
EJEMPLO 76	Procedimiento para evitar que un equipo nuevo se agregue al cluster global	241
EJEMPLO 77	Procedimiento para permitir que todos los equipos nuevos se agreguen al cluster global	241
EJEMPLO 78	Procedimiento para especificar que se agregue un equipo nuevo al cluster global	241
EJEMPLO 79	Configuración de la autenticación en UNIX estándar	242
EJEMPLO 80	Configuración de la autenticación en DES	242
EJEMPLO 81	Cambio de nombre de host privado	247
EJEMPLO 82	Cómo poner un nodo del cluster global en estado de mantenimiento	252
EJEMPLO 83	Eliminación de un nodo de cluster del estado de mantenimiento y restablecimiento del recuento de votos de quórum	253
EJEMPLO 84	Eliminación de un cluster de zona de un cluster global	274
EJEMPLO 85	Eliminación de un sistema de archivos local de alta disponibilidad en un cluster de zona	276
EJEMPLO 86	Eliminación de un sistema de archivos ZFS de alta disponibilidad en un cluster de zona	277
EJEMPLO 87	Eliminación de un conjunto de discos de Solaris Volume Manager de un cluster de zona	278

EJEMPLO 88	Eliminación de un dispositivo de DID de un cluster de zona	279
EJEMPLO 89	Restauración del sistema de archivos raíz ZFS (/) (Solaris Volume Manager)	306

Uso de esta documentación

La Guía de administración del sistema de Oracle® Solaris Cluster presenta procedimientos para administrar una configuración de Oracle Solaris Cluster en sistemas basados en SPARC y en x86.

- **Visión general:** describe cómo llevar a cabo una configuración de Oracle Solaris Cluster
- **Destinatarios:** técnicos, administradores de sistemas y proveedores de servicios autorizados.
- **Conocimiento requerido:** experiencia avanzada en la resolución de problemas y en el reemplazo de hardware.

Biblioteca de documentación del producto

La documentación y los recursos para este producto y los productos relacionados se encuentran disponibles en <http://www.oracle.com/pls/topic/lookup?ctx=E62278>.

Comentarios

Puede dejar sus comentarios sobre esta documentación en <http://www.oracle.com/goto/docfeedback>.

Introducción a la administración de Oracle Solaris Cluster

En este capítulo, se proporciona la siguiente información sobre la administración de clusters globales y de zona. Además, se describen los procedimientos para utilizar las herramientas de administración de Oracle Solaris Cluster:

- “Visión general de la administración de Oracle Solaris Cluster” [24]
- “Trabajo con un cluster de zona” [24]
- “Restricciones de las funciones del sistema operativo Oracle Solaris” [26]
- “Herramientas de administración” [26]
- “Preparación para administrar el cluster” [29]
- “Administración del cluster” [30]

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

Un cluster global está compuesto por uno o varios nodos del cluster global. Un cluster global puede incluir también las zonas no globales con marca `solaris` o `solaris10` que no son nodos pero están configuradas con el servicio de datos HA para Zonas.

Un cluster de zona está compuesto por una o varias zonas no globales de la marca `solaris`, `solaris10` o `labeled` configuradas con el atributo `cluster`. No se permite ningún otro tipo de marca en un cluster de zona. El cluster de zona con la marca `labeled` solo se utiliza con la función Trusted Extensions del software de Oracle Solaris. Para crear un cluster de zona, debe usar el comando `clzonecluster`, la utilidad `clsetup` o la interfaz de explorador de Oracle Solaris Cluster Manager.

Puede ejecutar servicios admitidos en el cluster de zona del mismo modo que en un cluster global, con el aislamiento que proporcionan 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 de cluster de zona en una máquina. 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 [Oracle Solaris Cluster 4.3 Concepts Guide](#).

Visión general de la administración de Oracle Solaris Cluster

El entorno de alta disponibilidad de Oracle Solaris Cluster garantiza que los usuarios finales puedan disponer de las aplicaciones importantes. 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 los siguientes manuales.

- [Capítulo 1, “Planificación de la configuración de Oracle Solaris Cluster” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#)
- [Oracle Solaris Cluster 4.3 Concepts Guide](#)

La administración de Oracle Solaris Cluster se organiza en tareas, distribuidas en la documentación siguiente.

Nota - Algunas de estas tareas se pueden realizar mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Esto se menciona en los procedimientos de tareas individuales. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

- Tareas estándar para administrar y mantener el cluster global o el de zona con regularidad o, incluso, diariamente. 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 [Guía de administración y planificación de servicios de datos de Oracle Solaris Cluster 4.3](#).
- Tareas de servicios, como agregar o reparar hardware de almacenamiento o de red. Estas tareas se describen en [Oracle Solaris Cluster 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 hacerlo mientras el resto de los nodos continúan ejecutando las operaciones del cluster. A menos que se indique lo contrario, las tareas administrativas de Oracle Solaris Cluster deben efectuarse en el nodo del cluster global. Para los procedimientos que necesitan el cierre completo del cluster, puede reducir al mínimo el impacto sobre el sistema si el tiempo de inactividad se programa 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 ejecuta el comando. Por ejemplo, utilizar el comando `cluster` en el nodo del cluster global permite recuperar toda la información del cluster global 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 del cluster global, afecta a todos los clusters de zona del cluster global. Los comandos de cluster de zona también afectan a todos los nodos del cluster de zona, incluso aunque un nodo del cluster de zona esté inactivo al ejecutarse el comando.

Los clusters de zona admiten la administración delegada de los recursos que están bajo el control del gestor 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 del cluster global es el único facultado para crear, modificar o suprimir dependencias que superen los límites de los cluster de zona.

En la lista siguiente, se muestran las tareas administrativas principales que se realizan en un cluster de zona.

- **Inicio y reinicio de un cluster de zona:** consulte [Capítulo 3, Cierre y arranque de un cluster](#). También puede iniciar y reiniciar un cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).
- **Agregación de un nodo a un cluster de zona:** consulte [Capítulo 8, Administración de nodos del cluster](#).
- **Eliminación de un nodo de un cluster de zona:** consulte [Eliminación de un nodo de un cluster de zona \[228\]](#). También puede desinstalar el software de un nodo de cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).
- **Visualización de la configuración de un cluster de zona:** consulte [Visualización de la configuración del cluster \[42\]](#). También puede ver la configuración de un cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).
- **Validación de la configuración de un cluster de zona:** consulte [Validación de una configuración básica de cluster \[51\]](#).
- **Detención de un cluster de zona:** consulte [Capítulo 3, Cierre y arranque de un cluster](#). También puede cerrar un cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Restricciones de las funciones del sistema operativo Oracle Solaris

No active ni desactive los siguientes servicios de Oracle Solaris Cluster mediante la interfaz de gestión de la utilidad de gestión de servicios (SMF).

TABLA 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
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

Puede realizar tareas administrativas para una configuración de Oracle Solaris Cluster mediante la línea de comandos o la interfaz de explorador de Oracle Solaris Cluster Manager. En la siguiente sección, se ofrece una visión general de las herramientas de línea de comandos y Oracle Solaris Cluster Manager.

Interfaz de explorador de Oracle Solaris Cluster Manager

El software de Oracle Solaris Cluster admite una interfaz de explorador, Oracle Solaris Cluster Manager, que se puede usar para llevar a cabo diversas tareas administrativas en el cluster. Consulte [Capítulo 13, Uso de la interfaz de explorador de Oracle Solaris Cluster Manager](#) para obtener más información. También puede obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager en [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

A continuación, se muestran algunas de las tareas que se pueden realizar en Oracle Solaris Cluster Manager:

- Crear y actualizar un cluster de zona
- Crear recursos y grupos de recursos
- Agregar un sistema de archivos, un host lógico o almacenamiento compartido a un cluster de zona
- Crear un servicio de datos de Oracle Database
- Gestionar nodos en un cluster global o un cluster de zona
- Agregar y gestionar servidores y dispositivos de quórum
- Agregar y gestionar dispositivos de almacenamiento NAS y gestionar discos y grupos de dispositivos
- Gestionar asociaciones de Geographic Edition

Interfaz de línea de comandos

Puede llevar a cabo la mayoría de las tareas de administración de Oracle Solaris Cluster de manera interactiva con la utilidad `clsetup`. Siempre que sea posible, los procedimientos de administración detallados en esta guía emplean la utilidad `clsetup`.

La utilidad `clsetup` permite administrar las siguientes opciones del menú principal.

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

En la siguiente lista, se incluyen otros comandos utilizados para administrar una configuración de Oracle Solaris Cluster. Consulte las páginas del comando `man` si desea obtener información más detallada.

`if_mpadm(1M)`

Conmuta las direcciones IP de un adaptador a otro en un grupo de varias rutas de red IP.

`claccess(1CL)`

Gestiona las políticas de acceso de Oracle Solaris Cluster para agregar nodos.

`cldevice(1CL)`

Gestiona los dispositivos de Oracle Solaris Cluster.

`cldevicegroup(1CL)`

Gestiona los grupos de dispositivos de Oracle Solaris Cluster.

`clinterconnect(1CL)`

Gestiona la interconexión de Oracle Solaris Cluster.

`clnasdevice(1CL)`

Gestiona el acceso a los dispositivos NAS para una configuración de Oracle Solaris Cluster.

`clnode(1CL)`

Gestiona los nodos de Oracle Solaris Cluster.

`clquorum(1CL)`

Gestiona el quórum de Oracle Solaris Cluster.

`clreslogicalhostname(1CL)`

Gestiona los recursos de Oracle Solaris Cluster para los nombres de host lógicos.

`clresource(1CL)`

Gestiona los recursos de los servicios de datos de Oracle Solaris Cluster.

`clresourcegroup(1CL)`

Gestiona los recursos de los servicios de datos de Oracle Solaris Cluster.

`clresourcetype(1CL)`

Gestiona los recursos de los servicios de datos de Oracle Solaris Cluster.

`clressharedaddress(1CL)`

Gestiona recursos de Oracle Solaris Cluster para direcciones compartidas.

`clsetup(1CL)`

Crea un cluster de zona y configura de manera interactiva una configuración de Oracle Solaris Cluster.

`clsnmphost(1CL)`

Administra los hosts SNMP de Oracle Solaris Cluster.

`clsnmpmib(1CL)`

Administra la MIB de SNMP de Oracle Solaris Cluster.

clsnmpuser(1CL)

Administra los usuarios de SNMP de Oracle Solaris Cluster.

cltelemetryattribute(1CL)

Configura la supervisión de los recursos del sistema.

cluster(1CL)

Gestiona la configuración y el estado global de la configuración de Oracle Solaris Cluster.

clzonecluster(1CL)

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 la configuración de hardware de Oracle Solaris Cluster

Documente los aspectos del hardware que sean exclusivos 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

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

Puede usar la utilidad `pconsole` para ejecutar ventanas de terminal para cada nodo de cluster más una ventana maestra que ejecute los comandos que se introducen allí para todos los nodos

al mismo tiempo. Para obtener información acerca de la instalación del software pconsole en la consola de administración, consulte [“Cómo instalar el software de pconsole en una consola de administración” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

Puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para configurar, supervisar y administrar el cluster y los componentes del cluster. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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 el uso de una consola tiene 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 una copia de seguridad del cluster de forma regular. Aunque el software de Oracle Solaris Cluster ofrece un entorno de alta disponibilidad, con copias reflejadas de los datos en los dispositivos de almacenamiento, el software de Oracle Solaris Cluster no sustituye las copias de seguridad periódicas. Una configuración de Oracle Solaris Cluster puede sobrevivir a varios errores, pero no ofrece protección contra errores de los usuarios o del programa, ni errores fatales. 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

Administración del cluster

[Tabla 2, “Herramientas de administración de Oracle Solaris Cluster”](#) proporciona un punto de partida para administrar el cluster.

TABLA 2 Herramientas de administración de Oracle Solaris Cluster

Tarea	Herramienta	Instrucciones
Iniciar sesión en el cluster de forma remota	Use la utilidad <code>pconsole</code> de Oracle Solaris de la línea de comandos para iniciar sesión en el cluster de manera remota.	“Inicio de sesión de manera remota en el cluster” [32] “Conexión segura a las consolas del cluster” [32]
Configurar el cluster de forma interactiva	Utilice el comando <code>clzonecluster</code> o la utilidad <code>clsetup</code> .	Obtención de acceso a las utilidades de configuración del cluster [33]
Mostrar la información y el número de versión de Oracle Solaris Cluster	Utilice el comando <code>clnode</code> con el subcomando y la opción <code>show - rev - v-nodo</code> .	Cómo visualizar la información de versión de Oracle Solaris Cluster [34]
Mostrar los recursos, grupos de recursos y tipos de recursos instalados	Utilice los comandos siguientes para mostrar la información sobre recursos: ■ <code>clresource</code> ■ <code>clresourcegroup</code> ■ <code>clresourcetype</code>	Visualización de tipos de recursos configurados, grupos de recursos y recursos [36]
Supervisar gráficamente los componentes del cluster	Use Oracle Solaris Cluster Manager.	Consulte la ayuda en pantalla.
Administrar gráficamente algunos componentes del cluster	Use Oracle Solaris Cluster Manager.	Consulte la ayuda en pantalla.
Comprobar el estado de los componentes del cluster	Utilice el comando <code>cluster</code> con el subcomando <code>status</code> .	Comprobación del estado de los componentes del cluster [38]
Comprobar el estado de los grupos IPMP en la red pública	En un cluster global, utilice el comando <code>clnode status</code> con la opción <code>-m</code> . Para un cluster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	Comprobación del estado de la red pública [41]
Ver la configuración del cluster	Para un cluster global, utilice el comando <code>cluster</code> con el subcomando <code>show</code> . Para un cluster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	Visualización de la configuración del cluster [42]
Ver y mostrar los dispositivos NAS configurados	Para un cluster global o un cluster de zona, utilice el comando <code>clzonecluster</code> con el subcomando <code>show</code> .	<code>clnasdevice(1CL)</code>
Comprobar los puntos de montaje globales o verificar la configuración del cluster	Para un cluster global, utilice el comando <code>cluster</code> con el subcomando <code>check</code> . En un cluster de zona, utilice el comando <code>clzonecluster verify</code> .	Validación de una configuración básica de cluster [51]

Tarea	Herramienta	Instrucciones
Examinar el contenido de los logs de comandos de Oracle Solaris Cluster	Examine el archivo <code>/var/cluster/logs/ commandlog</code> .	Cómo ver el contenido de los logs de comandos de Oracle Solaris Cluster [59]
Examinar los mensajes del sistema de Oracle Solaris Cluster	Examine el archivo <code>/var/adm/messages</code> .	“Formatos de mensaje del sistema” de <i>Resolución de problemas de administración del sistema en Oracle Solaris 11.3</i> También puede ver los mensajes del sistema de un nodo en la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte Cómo acceder a Oracle Solaris Cluster Manager [310].
Supervisar el estado de Solaris Volume Manager	Utilice el comando <code>metastat</code> .	Solaris Volume Manager Administration Guide

Inicio de sesión de manera remota en el cluster

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

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

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

Conexión segura a las consolas del cluster

Si el concentrador de terminales o el controlador del sistema admite `ssh`, puede usar la utilidad `pconsole` para conectarse a las consolas de esos sistemas. La utilidad `pconsole` es parte del paquete `terminal/pconsole` de Oracle Solaris y se instala cuando se instala el paquete. La utilidad `pconsole` crea una ventana de terminal de host para cada host remoto que se especifica en la línea de comandos. La utilidad también abre una ventana de consola central o principal

que propaga lo que introduce allí a cada una de las conexiones que abre. Consulte la página del comando `man pconsole(1)` para obtener más información.

▼ Obtención de acceso a las utilidades de configuración del 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)`

Nota - También puede llevar a cabo este procedimiento mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

1. Asuma el rol `root` en un nodo de miembro activo de 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
```

Se muestra la q.

■ 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 usar la utilidad `clsetup` interactiva o la interfaz de explorador de Oracle Solaris Cluster Manager para crear un cluster de zona o agregar un sistema de archivos o 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 [Guía de instalación del software de Oracle Solaris Cluster 4.3](#) para obtener instrucciones sobre el uso de la utilidad `clsetup` o Oracle Solaris Cluster Manager para configurar un cluster de zona.

3. En el menú, elija la configuración.

Siga las instrucciones en pantalla para finalizar una tarea. Para obtener más detalles, siga las instrucciones de “Creación y configuración de un cluster de zona” de [Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

Véase también Consulte la ayuda en pantalla y las páginas del comando `man clsetup` o `clzonecluster` para obtener más información.

▼ Cómo visualizar la información de versión de Oracle Solaris Cluster

Para realizar este procedimiento, no es necesario que inicie sesión con el rol `root`. Siga todos los pasos de este procedimiento desde un nodo del cluster global.

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, 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 versión de Oracle Solaris Cluster para todos los paquetes de Oracle Solaris Cluster.

ejemplo 1 Visualización de la información de versión de Oracle Solaris Cluster

En el siguiente ejemplo, se muestra la información de versión del cluster para los paquetes que incluyen Oracle Solaris Cluster 4.2.

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

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

```

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

```

```

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

```

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

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Nota - También puede ver los recursos y grupos de recursos mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Antes de empezar Los usuarios que no tienen el rol `root` necesitan la autorización de RBAC `solaris.cluster.read` para utilizar este subcomando.

- **Visualice los tipos de recursos, los grupos de recursos y los recursos configurados para el cluster.**

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

Siga todos los pasos de este procedimiento desde un nodo del cluster global. Si desea obtener información sobre un determinado recurso, los grupos de recursos y los tipos de recursos, utilice el subcomando show con uno de los comandos siguientes:

- resource
- resource group
- resourcetype

ejemplo 2 Visualización de tipos de recursos, grupos de recursos y recursos configurados

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

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

```
=== Registered Resource Types ===
```

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

```
=== Resource Groups and Resources ===
```

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

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

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

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

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

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Nota - También puede comprobar el estado de los componentes de cluster mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Consulte la ayuda en pantalla de Oracle Solaris Cluster Manager para obtener más información.

Oracle Solaris Cluster Manager y el comando `cluster status` también muestran el estado de un cluster de zona.

Antes de empezar Los usuarios que no tienen el rol `root` necesitan la autorización de RBAC `solaris.cluster.read` para utilizar el subcomando `status`.

● Compruebe el estado de los componentes del cluster.

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

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

ejemplo 3 Comprobación del estado de los componentes del cluster

En el siguiente ejemplo se muestra parte de la información de estado de los componentes del cluster devueltos por el comando `cluster status`.

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

--- Node Status ---

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

=== Cluster Transport Paths ===

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

=== Cluster Quorum ===

--- Quorum Votes Summary ---

Needed  Present  Possible
-----
3        3        4

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

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

--- Quorum Votes by Device ---

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

=== Cluster Device Groups ===

--- Device Group Status ---

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

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

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

=== Cluster Resource Groups ===

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

=== Cluster Resources ===

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

```
test_1      phys-schost-1    Online    Online
            phys-schost-2    Online    Online

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

=== Zone Clusters ===

--- Zone Cluster Status ---

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

▼ Comprobación del estado de la red pública

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Para comprobar el estado de los grupos de múltiples rutas de red IP, utilice el comando con el comando `clnode status`.

Antes de empezar Los usuarios que no tienen el rol `root` necesitan la autorización de RBAC `solaris.cluster.read` para utilizar este subcomando.

Nota - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para comprobar el estado de un nodo. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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

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

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

ejemplo 4 Comprobación del estado de la red pública

En el siguiente ejemplo, se muestra parte de la información de estado de los componentes del cluster devueltos por el comando `clnode status`.

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

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

▼ Visualización de la configuración del cluster

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Nota - También puede ver la configuración de un cluster mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Antes de empezar Los usuarios que no tienen el rol `root` necesitan la autorización de RBAC `solaris.cluster.read` para utilizar el subcomando `status`.

● **Visualice la configuración de un cluster global o un cluster de zona.**

```
% cluster show
```

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

Al ejecutar el comando `cluster show` desde un nodo del 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 solo de los clusters de zona. Entre las propiedades de un cluster de zona se incluyen el nombre, el tipo de IP, el inicio automático y la ruta de zona. El subcomando `show`

se ejecuta dentro de un cluster de zona y se aplica solo a ese cluster de zona. Al ejecutar el comando `clzonecluster show` desde un nodo del cluster de zona, solo 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 detalladas. Consulte la página del comando `man cluster(1CL)` para obtener detalles. Consulte la página del comando `man clzonecluster(1CL)` para obtener más información sobre `clzonecluster`.

ejemplo 5 Visualización de la configuración del cluster global

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

```
phys-schost# cluster show

=== Cluster ===

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

=== Host Access Control ===

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

=== Cluster Nodes ===

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

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

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

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

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

SNMP MIB Name:	Event
State:	Disabled
Protocol:	SNMPv2

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

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

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

Node Name:	phys-schost-2:za
Node ID:	2
Type:	cluster
Enabled:	yes
privatehostname:	clusternode2-priv
reboot_on_path_failure:	disabled
globalzoneshares:	1

```

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

```

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

```

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

```

```

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

```

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

```

SNMP MIB Name:                  Event
State:                          Disabled
Protocol:                       SNMPv2

```

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

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

```
=== Transport Cables ===
```

```

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

```

```

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

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

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

=== Transport Switches ===

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

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

=== Quorum Devices ===

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

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

=== Device Groups ===

```

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

=== Registered Resource Types ===

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

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

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

```
Resource Type:          SUNW.haderby
RT_description:          haderby server for Oracle Solaris Cluster
RT_version:              1
```

```

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

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

--- Resources for Group HA_RG ---

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

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

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

```

```

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

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

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

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

=== DID Device Instances ===

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

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

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

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

```

```

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

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

=== NAS Devices ===

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

```

ejemplo 6 Visualización de la información del cluster de zona

El siguiente ejemplo muestra las propiedades de la configuración del cluster de zona con RAC.

```

% clzonecluster show
=== Zone Clusters ===

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

--- Solaris Resources for sczone ---

Resource Name:            net
address:                  172.16.0.1
physical:                 auto

Resource Name:            net
address:                  172.16.0.2
physical:                 auto

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

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 configurados 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` utiliza el subcomando `check` para validar la configuración básica que se necesita para que el cluster global funcione correctamente. Si ninguna comprobación falla, `cluster check` vuelve a la petición de datos del shell. Si falla alguna de las comprobaciones, `cluster check` crea informes 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.

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

- Las comprobaciones interactivas requieren información del usuario que las comprobaciones no pueden determinar. La comprobación solicita al usuario la información necesaria, por ejemplo, el número de versión del firmware. Utilice la palabra clave `-k interactive` para especificar una o más comprobaciones interactivas.
- Las comprobaciones funcionales ejercen una función o un comportamiento determinados del cluster. La comprobación solicita una entrada del usuario, por ejemplo, qué nodo debe utilizar para la conmutación por error o la confirmación para iniciar o continuar la comprobación. Utilice la palabra clave `-k functional check-id` para especificar una comprobación funcional. Realice solo 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 primero 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)` desde el nodo del cluster global, se lleva a cabo un juego 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` vuelve a la petición de datos del 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 genera 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. Asuma el rol `root` en un nodo de miembro activo de un cluster global.

```
phys-schost# su
```

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

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

- a. **Vaya al separador Patches & Updates (Parches y actualizaciones) de [My Oracle Support](#).**
- b. **En Advanced Search (Búsqueda avanzada), seleccione `Solaris Cluster` como Product (Producto) y escriba `check` (comprobación) en el campo Description (Descripción).**

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

- c. **Aplique todas las actualizaciones de software que aún no estén instaladas en el cluster.**

3. Ejecute las comprobaciones de validación básicas.

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

`-v` Modo detallado.

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

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

4. Ejecute las comprobaciones de validación interactivas.

```
phys-schost# 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 no detallado.**

```
phys-schost# 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.

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

-C *check-ID* 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.**

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

-k *functional* Especifica comprobaciones de validación funcionales en ejecución.

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

- e. **Repita el paso c y el paso d para cada comprobación funcional que quede por ejecutar.**

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

6. **Si tiene configurado un cluster de zona, verifique la configuración de este cluster para saber si se puede instalar un cluster de zona.**

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

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 4.3.](#)

ejemplo 7 Comprobación de la configuración del cluster global con resultado correcto en todas las comprobaciones básicas

El ejemplo siguiente muestra la ejecución de `cluster check` en modo detallado para los nodos `phys-schost-1` y `phys-schost-2` con un resultado correcto en todas las comprobaciones.

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

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

```
cluster check: phys-schost-2: Explorer finished.
cluster check: phys-schost-2: Starting single-node checks.
cluster check: phys-schost-2: Single-node checks finished.
cluster check: Starting multi-node checks.
cluster check: Multi-node checks finished
```

ejemplo 8 Listado de comprobaciones de validación interactivas

En el siguiente ejemplo se muestran todas las comprobaciones interactivas que están disponibles para ejecutarse en el cluster. En la salida de 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 9 Ejecución de una comprobación de validación funcional

El siguiente ejemplo muestra primero la lista detallada de comprobaciones funcionales. Luego, la descripción detallada se muestra para la comprobación F6968101, lo cual indica que la comprobación podría interrumpir los servicios del cluster. El cluster se retira de la producción. Luego, se ejecuta la comprobación funcional 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
'clresourcegroup switch' on specified resource group
either to specified node or to all nodes in succession.
Version: 1.2
Revision Date: 12/10/10
```

Elimine el cluster de la producción

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

```
=====

>>> Functional Check

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

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

Select 'continue' for more details on this check.

    1) continue
    2) exit

choice: 1

=====

>>> Check Description <<<

Siga las instrucciones que aparecen en pantalla
```

ejemplo 10 Comprobación de la configuración del cluster global con una comprobación con resultado no satisfactorio

El ejemplo siguiente muestra el nodo phys-schost-2 del 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 -n phys-schost-1,phys-schost-2 -o/var/cluster/logs/
cluster_check/Dec5/

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

```
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE: Global filesystem /etc/vfstab entries are not consistent across
all Oracle SolarisCluster 4.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
```

▼ Comprobación de los puntos de montaje globales

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

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

1. Asuma el rol `root` en un nodo de miembro activo de 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 11 Comprobación de puntos de montaje globales

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

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

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

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

...
=====
= ANALYSIS DETAILS =
=====
-----
CHECK ID : 3065
SEVERITY : HIGH
FAILURE: Global filesystem /etc/vfstab entries are not consistent across
all Oracle Solaris Cluster 4.x nodes.
ANALYSIS : The global filesystem /etc/vfstab entries are not consistent across
all nodes in this cluster.
Analysis indicates:
FileSystem '/global/phys-schost-1' is on 'phys-schost-1' but missing from 'phys-schost-2'.
RECOMMEND: Ensure each node has the correct /etc/vfstab entry for the
filesystem(s) in question.
...
#
# cat /var/cluster/logs/cluster_check/Dec5/cluster_check-results.phys-schost-1.txt

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

▼ Cómo ver el contenido de los logs 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 cuando se configura el cluster, y la operación finaliza cuando se cierra el cluster. Los comandos se registran en todos los nodos activos que se han iniciado 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`
- `clsnmphost`
- `clsnmpmib`
- `clsnmpuser`
- `cltelemetryattribute`
- `cluster`
- `clzonecluster`

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

- Fecha y registro de hora
- 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 de Oracle Solaris Cluster *no* muestra un estado de salida en el archivo `commandlog`.

Por defecto, el archivo `commandlog` se archiva periódicamente una vez por semana. Para modificar las políticas 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\(1\)](#).

En cualquier momento dado, el software de Oracle Solaris Cluster conserva hasta ocho archivos `commandlog` archivados anteriormente en cada nodo de 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 12 Visualización del contenido de los logs 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
```


♦ ♦ ♦ 2 C A P Í T U L O 2

Oracle Solaris Cluster y RBAC

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

- “Configuración y uso de RBAC con Oracle Solaris Cluster” [63]
- “Perfiles de derechos de RBAC de Oracle Solaris Cluster” [63]
- “Creación y asignación de un rol RBAC con un perfil de derechos de gestión de Oracle Solaris Cluster” [65]
- “Modificación de las propiedades de RBAC de un usuario” [66]

Configuración 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 capítulo, se describen los pasos específicos para configurar y utilizar RBAC con el software de Oracle Solaris Cluster.

Tarea	Instrucciones
Más información sobre RBAC	Capítulo 1, “Sobre el uso de los derechos para controlar los usuarios y los procesos” de <i>Protección de los usuarios y los procesos en Oracle Solaris 11.3</i>
Instalar, administrar elementos y utilizar RBAC	Capítulo 3, “Asignación de derechos en Oracle Solaris” de <i>Protección de los usuarios y los procesos en Oracle Solaris 11.3</i>
Más información sobre elementos y herramientas de RBAC	Capítulo 8, “Referencia para derechos Oracle Solaris” de <i>Protección de los usuarios y los procesos en Oracle Solaris 11.3</i>

Perfiles de derechos de RBAC de Oracle Solaris Cluster

Algunos comandos y opciones de Oracle Solaris Cluster que ejecuta en la línea de comandos usan RBAC para la autorización. Los comandos y las opciones de Oracle Solaris Cluster

que requieren autorización de RBAC necesitarán uno o más de los siguientes niveles de autorización. Los perfiles de derechos de RBAC de Oracle Solaris Cluster se aplican a los nodos de 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 crear, suprimir y cambiar las propiedades de un objeto de cluster.

Para obtener más información sobre la autorización de RBAC requerida por 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 usuarios o roles, a fin de proporcionarles diferentes niveles de acceso a Oracle Solaris Cluster. Oracle proporciona los siguientes perfiles de derechos con el software de Oracle Solaris Cluster.

Perfil de derechos	Autorizaciones y atributos de seguridad	Derechos otorgados
Comandos de Oracle Solaris Cluster	Una lista de los comandos de Oracle Solaris Cluster que se ejecutan con el atributo de seguridad <code>euid=0</code> .	Ejecute comandos de Oracle Solaris Cluster seleccionados para configurar y gestionar un cluster, incluidos los siguientes subcomandos para todos los comandos de Oracle Solaris Cluster. ■ <code>list</code> ■ <code>show</code> ■ <code>status</code> <code>scha_control</code> <code>scha_resource_get</code> <code>scha_resource_setstatus</code> <code>scha_resourcegroup_get</code> <code>scha_resourcetype_get</code>
Usuario de Oracle Solaris básico	Este perfil de derechos de Oracle Solaris contiene autorizaciones de Oracle Solaris, además de la siguiente autorización:	<code>solaris.cluster.read</code> : permite realizar operaciones de lectura de listar y mostrar, entre otras, para los comandos de Oracle Solaris Cluster, además de acceder a la interfaz de explorador de Oracle Solaris Cluster Manager.
Operación del cluster	El perfil de derechos es específico del software de Oracle Solaris Cluster y cuenta con las autorizaciones siguientes:	<code>solaris.cluster.read</code> : permite realizar operaciones de lectura de listar, mostrar, exportar y ver estado, entre otras, además de acceder a la interfaz de explorador de Oracle Solaris Cluster Manager.

Perfil de derechos	Autorizaciones y atributos de seguridad	Derechos otorgados
Administrador del sistema	Este perfil de derechos de Oracle Solaris contiene las mismas autorizaciones que el perfil de administración del cluster.	<code>solaris.cluster.admin</code> : permite cambiar el estado de los objetos de 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 <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 un rol RBAC con un perfil de derechos de gestión de Oracle Solaris Cluster

Use esta tarea para crear un rol RBAC nuevo con un perfil de derechos de gestión de Oracle Solaris Cluster y para asignar usuarios a este rol nuevo.

▼ Creación de una función desde la línea de comandos

1. Seleccione un método para crear una función:

- Para los roles en el ámbito local, use el comando `roleadd` para especificar un rol local nuevo y sus atributos. Para obtener más información, consulte la página del comando [man roleadd\(1M\)](#).
- De manera alternativa, para los roles en el ámbito local, edite el archivo `user_attr` para agregar un usuario con `type=role`. Para obtener más información, consulte la página del comando [man user_attr\(4\)](#).

Use este método sólo en caso de emergencia.

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

Este comando requiere autenticación por parte del rol `root` que puede crear otros roles. Puede aplicar el comando `roleadd` a todos los servicios de nombres. Este comando se ejecuta como cliente del servidor de Solaris Management Console.

2. 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 13 Creación de una función de operador personalizado con el comando `smrole`

En la siguiente secuencia se muestra cómo se crea un rol 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 de cuentas de usuarios o con la línea de comandos. Para modificar las propiedades RBAC de un

usuario, consulte [Modificación de las propiedades de RBAC de un usuario desde la línea de comandos \[68\]](#). Elija uno de los siguientes procedimientos.

- [Modificación de las propiedades de RBAC de un usuario con la herramienta de cuentas de usuario \[67\]](#)
- [Modificación de las propiedades de RBAC de un usuario desde la línea de comandos \[68\]](#)

▼ 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 la recopilación de herramientas de usuario como usuario root o asumir un rol que tenga asignado un perfil de derechos de administrador del sistema.

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

Para ejecutar la herramienta de cuentas de usuarios, inicie Solaris Management Console, como se describe en “[Uso de sus derechos administrativos asignados](#)” de *Protección de los usuarios y los procesos en Oracle Solaris 11.3*. 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Elija el comando adecuado:**
 - Para cambiar propiedades de usuario que están asignadas a un usuario definido en el ámbito local o en un depósito LDAP, use el comando `usermod`. Para obtener más información, consulte la página del comando man [usermod\(1M\)](#).
 - Otra posibilidad es editar el archivo `user_attr`.
Use este método sólo en caso de emergencia.
 - Para gestionar roles localmente o en un servicio de nombres, como un depósito LDAP, use los comandos `roleadd` o `rolemod`. Para obtener más información, consulte las páginas del comando man [roleadd\(1M\)](#) o [rolemod\(1M\)](#).

Estos comandos requieren autenticación como rol `root` que pueda modificar archivos de usuario. Puede aplicar estos comandos a todos los servicios de nombres. Consulte [“Comandos que se utilizan para la gestión de usuarios, roles y grupos” de Gestión de las cuentas de usuario y los entornos de usuario en Oracle Solaris 11.3](#).

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

3

♦ ♦ ♦ C A P Í T U L O 3

Cierre y arranque de un cluster

En este capítulo, se proporciona información sobre las tareas de cierre e inicio de un cluster global, un cluster de zona y nodos determinados, además de los procedimientos para llevarlas a cabo.

- “Visión general sobre el cierre y el inicio de un cluster” [69]
- “Cierre e inicio de un solo nodo de un cluster” [84]
- “Reparación de un sistema de archivos /var que haya alcanzado el límite de su capacidad” [98]

Para consultar una descripción general de los procedimientos descritos en este capítulo, consulte [Reinicio de un nodo en el modo que no es de cluster](#) [95] y [Tabla 4, “Mapa de tareas: cierre e inicio de un nodo”](#).

Visión general sobre el cierre y el inicio de un cluster

El comando `cluster shutdown` de Oracle Solaris Cluster detiene los servicios del cluster global de manera organizada y cierra sin errores un cluster global 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` dentro de un cluster de zona). Para obtener más información, consulte la página del comando [man cluster\(1CL\)](#).

En los procedimientos descritos en este capítulo, `phys-schost#` refleja una petición de datos de cluster global. La petición de datos 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 evacuate` para cerrar nodos individuales. Para obtener más información, consulte [Cierre de un cluster \[71\]](#), [“Cierre e inicio de un solo nodo de un cluster” \[84\]](#) o la página del comando `man clnode(1CL)`.

También puede evacuar un nodo mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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`, y lleva a todos los nodos del cluster a la solicitud de OpenBoot PROM `ok` en los sistemas basados en SPARC, o al mensaje `Press any key to continue` del menú de GRUB en los sistemas basados en x86. Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3](#). El comando `clzonecluster halt` ejecuta el comando `zoneadm -z zone-cluster-name halt` para detener (pero no cerrar) las zonas del cluster de zona.

Nota - Si es necesario, puede iniciar un nodo en el modo que no es de cluster para que el nodo no participe como miembro del cluster. El modo que no es de cluster es útil al instalar software de cluster o para efectuar determinados procedimientos administrativos. Si desea obtener más información, consulte [Reinicio de un nodo en el modo que no es de cluster \[95\]](#).

TABLA 3 Lista de tareas: cerrar e iniciar un cluster

Tarea	Instrucciones
Detener el cluster.	Cierre de un cluster [71]
Iniciar el cluster mediante el inicio de todos los nodos. Los nodos deben contar con una conexión operativa con la interconexión de cluster para conseguir convertirse en miembros del cluster.	Inicio de un cluster [73]
Reiniciar el cluster.	Reinicio de un cluster [77]

▼ Cierre de un cluster

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



Atención - 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 una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

1. **(x86 solamente) 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. **Asuma un rol que proporcione la autorización de 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 zone-cluster-name
```

- **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 14 Cierre de un cluster de zona

En el siguiente ejemplo se cierra un cluster de zona denominado *zona_sc*.

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

ejemplo 15 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 cual permite que se muestre la petición de datos ok. La opción `-g 0` establece el período de 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 16 x86: Cierre de un cluster global

En el ejemplo siguiente se muestra la salida de la consola cuando se detiene el funcionamiento normal del cluster global y se cierran todos los nodos. En este ejemplo, el indicador ok no se visualiza en todos los nodos. La opción `-g 0` establece el período de gracia de cierre en cero y 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 [Inicio de un cluster \[73\]](#) para reiniciar un cluster global o un cluster de zona que se ha cerrado.

▼ Inicio de un cluster

Este procedimiento explica cómo iniciar 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 una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Nota - Para crear un cluster de zona, siga las instrucciones de [“Creación y configuración de un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#) o utilice la interfaz de explorador de Oracle Solaris Cluster Manager para crear el cluster de zona.

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.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3](#).

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 zone-cluster-name
```

■ **Si tiene más de un cluster de zona, puede iniciar todos los clusters de zona. Use el signo más (+) en lugar de *zone-cluster-name*.**

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

El comando `cluster status` informa sobre el estado de los nodos del cluster global.

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

Si ejecuta el comando de estado `clzonecluster status` desde un nodo del cluster global, el comando informa el estado del nodo del cluster de zona.

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

Nota - Si el sistema de archivos `/var` de un nodo alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en ese nodo. Si surge este problema, consulte [Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad \[98\]](#). Para obtener más información, consulte la página del comando `man clzonecluster(1CL)`.

ejemplo 17 SPARC: Inicio de un cluster global

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

Cuando se reinicia 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.
...
```

ejemplo 18 x86: Inicio de un cluster

El ejemplo siguiente muestra la salida de la consola cuando el nodo `phys-schost-1` se inicia en el cluster. Aparecen mensajes similares en las consolas de los otros nodos del cluster.

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

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

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

Ch B, SCSI ID: 0 SEAGATE ST336605LC 160
SCSI ID: 1 SEAGATE ST336605LC 160
```

```

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

```

```

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

```

```

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

```

```

SunOS - Intel Platform Edition Primary Boot Subsystem, vsn 2.0

```

Current Disk Partition Information

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

Please select the partition you wish to boot: * *

Solaris DCB

loading /solaris/boot.bin

SunOS Secondary Boot version 3.00

Solaris Intel Platform Edition Booting System

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

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

Press ESCape to interrupt autoboot in 2 seconds.

Initializing system

Please wait...

Warning: Resource Conflict - both devices are added

NON-ACPI device: ISY0050

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

ACPI device: ISY0050

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

<<< Current Boot Parameters >>>

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

Boot args:

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

<<< timeout in 5 seconds >>>

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

▼ Reinicio de un cluster

Para cerrar un cluster de zona, ejecute el comando `cluster shutdown` y, luego, inicie el cluster global con el comando `boot` en cada nodo. Para cerrar un cluster de zona, use el comando `clzonecluster halt`; después, ejecute el comando `clzonecluster boot` para iniciar el cluster de zona. También puede usar el comando `clzonecluster reboot`. Para obtener más información, consulte las páginas del comando `man cluster(1CL)`, `boot(1M)` y `clzonecluster(1CL)`.

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, 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. Asuma un rol que proporcione la autorización de 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 zone-cluster-name
```

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. Inicie todos los nodos.

No importa el orden en que se inicien los nodos, a menos que haga modificaciones de configuración entre las 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.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3](#).

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

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

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

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

5. Compruebe que los nodos se hayan iniciado 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 nodo alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en ese nodo. Si surge este problema, consulte [Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad \[98\]](#).

ejemplo 19 Reinicio de un cluster de zona

El ejemplo siguiente muestra cómo detener e iniciar un cluster de zona denominado *sparse-sczone*. 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-sczone'
died.
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster 'sparse-sczone'
died.
Sep  5 19:17:46 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone'
died.
phys-schost#
phys-schost# clzonecluster boot sparse-sczone
Waiting for zone boot commands to complete on all the nodes of the zone cluster "sparse-
sczone"...
```

```
phys-schost# Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 1 of cluster
'sparse-sczone' joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 2 of cluster 'sparse-sczone'
joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 3 of cluster 'sparse-sczone'
joined.
Sep  5 19:18:23 schost-4 cl_runtime: NOTICE: Membership : Node 4 of cluster 'sparse-sczone'
joined.

phys-schost#
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---

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

ejemplo 20 SPARC: Reinicio de un cluster global

En el ejemplo siguiente, se muestra la salida de la consola cuando se detiene el funcionamiento normal del cluster global, todos los nodos se cierran y muestran la petición de datos ok y se reinicia el cluster global. La opción `-g 0` establece el período de gracia en cero y 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:

```

ejemplo 21 x86: Reinicio de un cluster

El siguiente ejemplo muestra la salida de la consola cuando se detiene el funcionamiento normal del cluster, se cierran todos los nodos y se reinicia el cluster. La opción `-g 0` establece el período de gracia en cero y 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.

```

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

ATI RAGE SDRAM BIOS P/N GR-xlint.007-4.330
*
BIOS Lan-Console 2.0
Copyright (C) 1999-2001 Intel Corporation
MAC ADDR: 00 02 47 31 38 3C
AMIBIOS (C)1985-2002 American Megatrends Inc.,
Copyright 1996-2002 Intel Corporation

```

```

SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064
2 X Intel(R) Pentium(R) III CPU family      1400MHz
Testing system memory, memory size=2048MB
2048MB Extended Memory Passed
512K L2 Cache SRAM Passed
ATAPI CD-ROM SAMSUNG CD-ROM SN-124

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

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

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

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

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

SunOS - Intel Platform Edition                Primary Boot Subsystem, vsn 2.0

Current Disk Partition Information

Part#   Status   Type        Start      Length
=====
1       Active   X86 BOOT    2428       21852
2                SOLARIS     24280      71662420
3                <unused>
4                <unused>
Please select the partition you wish to boot: *      *

Solaris DCB

loading /solaris/boot.bin

SunOS Secondary Boot version 3.00

Solaris Intel Platform Edition Booting System

Autobooting from bootpath: /pci@0,0/pci8086,2545@3/pci8086,1460@1d/

```

```
pci8086,341a@7,1/sd@0,0:a
```

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

Press ESCape to interrupt autoboot in 2 seconds.

Initializing system

Please wait...

Warning: Resource Conflict - both devices are added

NON-ACPI device: ISY0050

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

ACPI device: ISY0050

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

<<< Current Boot Parameters >>>

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

Boot args:

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

<<< timeout in 5 seconds >>>

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

Size: 275683 + 22092 + 150244 Bytes

/platform/i86pc/kernel/unix loaded - 0xac000 bytes used

SunOS Release 5.9 Version Generic_112234-07 32-bit

Copyright 1983-2003 Sun Microsystems, Inc. All rights reserved.

Use is subject to license terms.

configuring IPv4 interfaces: e1000g2.

Hostname: phys-schost-1

Booting as part of a cluster

NOTICE: CMM: Node phys-schost-1 (nodeid = 1) with votecount = 1 added.

NOTICE: CMM: Node phys-schost-2 (nodeid = 2) with votecount = 1 added.

NOTICE: CMM: Quorum device 1 (/dev/did/rdisk/dls2) added; votecount = 1, bitmask of nodes with configured paths = 0x3.

NOTICE: clcomm: Adapter e1000g3 constructed

NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being constructed

NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being initiated

NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 online

NOTICE: clcomm: Adapter e1000g0 constructed

NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being constructed

NOTICE: CMM: Node phys-schost-1: attempting to join cluster.

NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being initiated

NOTICE: CMM: Quorum device /dev/did/rdisk/dls2: owner set to node 1.

NOTICE: CMM: Cluster has reached quorum.

NOTICE: CMM: Node phys-schost-1 (nodeid = 1) is up; new incarnation number = 1068496374.

NOTICE: CMM: Node phys-schost-2 (nodeid = 2) is up; new incarnation number = 1068496374.

NOTICE: CMM: Cluster members: phys-schost-1 phys-schost-2.

NOTICE: CMM: node reconfiguration #1 completed.

NOTICE: CMM: Node phys-schost-1: joined cluster.

WARNING: mod_installdrv: no major number for rsmrdt

```
ip: joining multicasts failed (18) on clprivnet0 - will use link layer
broadcasts for multicast
The system is coming up. Please wait.
checking ufs filesystems
/dev/rdisk/clt0d0s5: is clean.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 online
NIS domain name is dev.eng.mycompany.com
starting rpc services: rpcbind keyserp ybind done.
Setting netmask of e1000g2 to 192.168.255.0
Setting netmask of e1000g3 to 192.168.255.128
Setting netmask of e1000g0 to 192.168.255.128
Setting netmask of clprivnet0 to 192.168.255.0
Setting default IPv4 interface for multicast: add net 224.0/4: gateway phys-schost-1
syslog service starting.
obtaining access to all attached disks

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

phys-schost-1 console login:
```

Cierre e inicio de un solo nodo de un cluster

Puede cerrar un nodo del cluster global o un nodo del cluster de zona. En esta sección, se proporcionan 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` solo cuando vaya a cerrar todo un cluster global.

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

Para obtener más información, consulte las páginas del comando `man clnode(1CL)`, `shutdown(1M)` y `clzonecluster(1CL)`.

En los procedimientos descritos en este capítulo, `phys-schost#` refleja una petición de datos de cluster global. La petición de datos de shell interactivo de `clzonecluster` es `clzc:schost>`.

TABLA 4 Mapa de tareas: cierre e inicio de un nodo

Tarea	Herramienta	Instrucciones
Detener un nodo.	Para un nodo de cluster global, use los comandos <code>clnode evacuate</code> y <code>shutdown</code> . Para un nodo del cluster de zona, utilice el comando <code>clzonecluster halt</code> .	Cierre de un nodo [85]
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, utilice el comando <code>boot</code> o <code>b</code> . Para un nodo del cluster de zona, utilice el comando <code>clzonecluster boot</code> .	Inicio de un nodo [89]
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 los comandos <code>clnode evacuate</code> y <code>shutdown</code> , seguidos de <code>boot</code> o <code>b</code> . Para un nodo del cluster de zona, utilice el comando <code>clzonecluster reboot</code> .	Reinicio de un nodo [92]
Iniciar un nodo de manera que no participe como miembro del cluster.	Para un nodo del cluster global, ejecute los comandos <code>clnode evacuate</code> y <code>shutdown</code> , seguidos de <code>boot -x</code> en SPARC o de la edición de entradas del menú 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.	Reinicio de un nodo en el modo que no es de cluster [95]

▼ 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 cortas y largas de los nombres de comando, los comandos son idénticos.



Atención - 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.

Nota - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para evacuar un nodo del cluster global y cambiar todos los grupos de recursos y grupos de dispositivos al siguiente nodo preferido. También puede cerrar un nodo del cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.admin` en el nodo de cluster que se cerrará.**

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 zone-cluster-name
```

Al especificar un determinado nodo de un cluster de zona, solo se detiene ese nodo. Por defecto, el comando `halt` detiene 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` conmuta todos los grupos de recursos y grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. (También puede ejecutar `clnode evacuate` dentro de un nodo de un 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.**

Ejecute el comando `shutdown` en el nodo del cluster global que desea 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 22 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 y 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 -i0
Wed Mar 10 13:47:32 phys-schost-1 cl_runtime:
WARNING: CMM monitoring disabled.
phys-schost-1#
INIT: New run level: 0
The system is coming down. Please wait.
Notice: rgmd is being stopped.
Notice: rpc.pmfd is being stopped.
Notice: rpc.fed is being stopped.
umount: /global/.devices/node@1 busy
umount: /global/phys-schost-1 busy
The system is down.
syncing file systems... done
Program terminated
ok
```

ejemplo 23 x86: 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 y 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 -i0
Shutdown started.    Wed Mar 10 13:47:32 PST 2004

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

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

```
umount: /global/.devices/node@1 busy
The system is down.
syncing file systems... done
WARNING: CMM: Node being shut down.
Type any key to continue
```

ejemplo 24 Cierre de un nodo de un cluster de zona

En el ejemplo siguiente, se muestra el uso de `clzonecluster halt` para cerrar un nodo presente en un cluster de zona denominado *sparse-sczzone*. 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-sczzone	schost-1	sczone-1	Online	Running
	schost-2	sczone-2	Online	Running
	schost-3	sczone-3	Online	Running
	schost-4	sczone-4	Online	Running

```
phys-schost#
```

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

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

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

```
phys-host#
```

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

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

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

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

```
phys-schost#
```

Véase también Consulte [Inicio de un nodo \[89\]](#) para ver cómo reiniciar un nodo del cluster global que se haya cerrado.

▼ Inicio de un nodo

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

Nota - También puede iniciar un nodo de cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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 cierren o reinicien.

Nota - La configuración del quórum puede afectar el inicio de un nodo. En un cluster de dos nodos, debe tener un dispositivo de quórum configurado de manera que el recuento total de quórum correspondiente al cluster sea de tres. Es conveniente tener un recuento de quórum para cada nodo y un recuento de quórum para el dispositivo de quórum. De esta forma, si se 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 en funcionamiento. 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 dominio de E/S o control se reinicia, el dominio invitado no envía ni recibe latidos. 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.

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, inicie 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.

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

- Si tiene un cluster de zona, puede especificar un nodo para que se inicie.

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

2. Compruebe que el nodo se haya iniciado 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 solo puede iniciarse en modo de cluster si el nodo que lo aloja se inicia en modo de cluster.

Nota - Si el sistema de archivos `/var` de un nodo alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda reiniciarse en ese nodo. Si surge este problema, consulte [Reparación de un sistema de archivos /var que haya alcanzado el límite de capacidad \[98\]](#).

ejemplo 25 SPARC: Inicio de un nodo del cluster global

El ejemplo siguiente muestra la salida de la consola cuando el nodo `phys-schost-1` se inicia 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:

```

ejemplo 26 x86: Inicio de un nodo de cluster

El ejemplo siguiente muestra la salida de la consola cuando el nodo phys-schost-1 se inicia en el cluster.

```

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

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

<<< timeout in 5 seconds >>>

Select (b)oot or (i)nterpreter: Size: 276915 + 22156 + 150372 Bytes
/platform/i86pc/kernel/unix loaded - 0xac000 bytes used
SunOS Release 5.9 Version on81-feature-patch:08/30/2003 32-bit
Copyright 1983-2003 Sun Microsystems, Inc. All rights reserved.
Use is subject to license terms.
configuring IPv4 interfaces: e1000g2.
Hostname: phys-schost-1
Booting as part of a cluster
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) with votecount = 1 added.
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) with votecount = 1 added.
NOTICE: CMM: Quorum device 1 (/dev/did/rdisk/dls2) added; votecount = 1, bitmask
of nodes with configured paths = 0x3.
WARNING: CMM: Initialization for quorum device /dev/did/rdisk/dls2 failed with
error EACCES. Will retry later.
NOTICE: clcomm: Adapter e1000g3 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being constructed
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g3 - phys-schost-2:e1000g3 online
NOTICE: clcomm: Adapter e1000g0 constructed
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being constructed
NOTICE: CMM: Node phys-schost-1: attempting to join cluster.
WARNING: CMM: Reading reservation keys from quorum device /dev/did/rdisk/dls2
failed with error 2.
NOTICE: CMM: Cluster has reached quorum.
NOTICE: CMM: Node phys-schost-1 (nodeid = 1) is up; new incarnation number =
1068503958.

```

```
NOTICE: CMM: Node phys-schost-2 (nodeid = 2) is up; new incarnation number =
1068496374.
NOTICE: CMM: Cluster members: phys-schost-1 phys-schost-2.
NOTICE: CMM: node reconfiguration #3 completed.
NOTICE: CMM: Node phys-schost-1: joined cluster.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 being initiated
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 online
NOTICE: CMM: Retry of initialization for quorum device /dev/did/rdisk/dls2 was
successful.
WARNING: mod_installdr: no major number for rsmrdt
ip: joining multicasts failed (18) on clprivnet0 - will use link layer
broadcasts for multicast
The system is coming up. Please wait.
checking ufs filesystems
/dev/rdisk/clt0d0s5: is clean.
NIS domain name is dev.eng.mycompany.com
starting rpc services: rpcbind keyserp ypbind done.
Setting netmask of e1000g2 to 192.168.255.0
Setting netmask of e1000g3 to 192.168.255.128
Setting netmask of e1000g0 to 192.168.255.128
Setting netmask of clprivnet0 to 192.168.255.0
Setting default IPv4 interface for multicast: add net 224.0/4: gateway phys-schost-1
syslog service starting.
obtaining access to all attached disks

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

phys-schost-1 console login:
```

▼ Reinicio de un nodo

phys-schost# refleja una petición de datos de cluster global. Siga este procedimiento en un cluster global.

Nota - También puede reiniciar un nodo de cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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 cortas y largas de los nombres de comando, los comandos son idénticos.



Atención - Si se excede el timeout de un método para un recurso y no se lo puede detener, el nodo se reinicia solamente 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.

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 cierren o reinicien.

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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.admin` en el nodo que se cerrará.**

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. En un cluster de zona, también se pueden ejecutar los comandos `clnode evacuate` y `shutdown`.

En el caso de un cluster global, escriba los comandos siguientes en el nodo que desee cerrar. El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. El comando también conmuta todos los grupos de recursos de las zonas globales del nodo especificado a las zonas globales de otros nodos que se sitúen a continuación en el orden de preferencia.

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

- Indique el nodo del cluster de zona que desee cerrar y reiniciar.

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

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 se haya iniciado 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 27 SPARC: Reinicio de un nodo del cluster global

El ejemplo siguiente 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 e inicio, 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 28 Reinicio de un nodo del cluster global

En el ejemplo siguiente, se muestra cómo reiniciar 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#

```

▼ Reinicio de un nodo en el modo que no es de cluster

Puede iniciar un nodo del cluster global en un modo sin cluster, donde el nodo no participa en la pertenencia al cluster. El modo que no es de cluster resulta útil a la hora de instalar software del cluster o en ciertos procedimientos administrativos, como la actualización de un nodo. Un nodo de un cluster de zona no puede estar en un estado de inicio diferente del que tenga el nodo del

cluster global subyacente. Si el nodo del cluster global se inicia en el modo que no es de cluster, el nodo del cluster de zona asume automáticamente el modo que no es de cluster.

`phys-schost#` refleja una petición de datos 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 cortas y largas de los nombres de comando, los comandos son idénticos.

1. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.admin` en el cluster que se iniciará en modo no de cluster.**

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

2. **Cierre el nodo del cluster de zona o el nodo del cluster global.**

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

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

```
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 zone-cluster-name
```

También puede utilizar los comandos `clnode evacuate` y `shutdown` dentro de un 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.**

Se muestra el menú de GRUB.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3](#).

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

Se muestra la pantalla de parámetros de inicio de GRUB.

- c. **Agregue `-x` al comando para especificar que el sistema arranque en un modo que no sea de 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/kernel/$ISADIR/unix -B $ZFS-BOOTFS -x
```

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

La pantalla muestra el comando editado.

- e. **Escriba `b` para iniciar el nodo en el modo sin 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 29 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 no 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 llenar el sistema de archivos /var. Si el sistema de archivos /var de un nodo de cluster alcanza su límite de capacidad, es posible que Oracle Solaris Cluster no pueda iniciarse en ese nodo en el próximo inicio del sistema. 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 informa que el sistema de archivos /var está completo y continúa ejecutando los servicios de Oracle Solaris Cluster, siga este procedimiento para borrar el sistema de archivos completo. Consulte [“Formatos de mensaje del sistema” de Resolución de problemas de administración del sistema en Oracle Solaris 11.3](#) para obtener más información.

1. **Asuma el rol root en el nodo del cluster con todo el sistema de archivos /var.**

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

En este capítulo, se describen las tecnologías de replicación de datos que puede usar con el software de 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 de Oracle Solaris Cluster admite los siguientes tipos de replicación de datos:

- Entre clusters: use Oracle Solaris Cluster Geographic Edition para recuperar datos en caso de problema grave.
- En un cluster: úselo como sustitución de la creación de reflejo 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 gestionar grupos de dispositivos de Solaris Volume Manager, ZFS o discos raw, consulte [“Administración de grupos de dispositivos” \[125\]](#).

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 la estructura de Oracle Solaris Cluster Geographic Edition para gestionar la replicación de datos para la recuperación ante desastres, consulte [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#).

En este capítulo, se incluye la siguiente sección:

- [“Comprensión de la replicación de datos” \[102\]](#)
- [“Uso de la replicación de datos basada en almacenamiento en un cluster de campus” \[104\]](#)

Comprensión de la replicación de datos

El software de Oracle Solaris Cluster admite la replicación de datos basada en host y en almacenamiento.

- *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, o entre un cluster y un host que no se encuentra en un cluster, es Función Availability Suite de Oracle Solaris.

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 que ejecutan el sistema operativo Oracle Solaris escriban datos en un volumen compartido, por ejemplo Oracle RAC.

- Para obtener más información sobre el uso de la replicación de datos basada en host entre dos clusters, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Solaris Availability Suite](#).
- Para ver un ejemplo de la replicación basada en host que *no* usa la estructura de Geographic Edition, consulte el apéndice A, [Apéndice A, Ejemplo de despliegue: configuración de replicación de datos basada en host entre clusters con el software Availability Suite](#).
- *La replicación de datos basada en almacenamiento* usa el software del controlador de almacenamiento para mover el trabajo de la replicación de datos desde los nodos del cluster hasta el dispositivo de almacenamiento. Este software libera algo de la capacidad de procesamiento del nodo para satisfacer las solicitudes del cluster. EMC SRDF es un ejemplo de software basado en almacenamiento que puede replicar los datos en 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 de campus” \[104\]](#).
- Para obtener más información sobre el uso de 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 4.3 Geographic Edition Overview](#).

Métodos admitidos de replicación de datos

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

- **Replicación entre clusters:** para recuperación ante desastres con el software de Oracle Solaris Cluster Geographic Edition, puede usar la replicación basada en host o en almacenamiento para realizar una replicación de datos entre clusters. Por lo general, se seleccionará o la replicación basada en hosts o la replicación basada en almacenamiento, en lugar de una combinación de ambas. Puede gestionar ambos tipos de replicación con el software de Geographic Edition.

Tipo de replicación	Productos de replicación de datos
Replicación basada en host	Función Availability Suite de Oracle Solaris.
Replicación basada en almacenamiento	EMC Symmetrix Remote Data Facility (SRDF). Hitachi TrueCopy and Hitachi Universal Replicator, mediante la estructura de Oracle Solaris Cluster Geographic Edition. Oracle ZFS Storage Appliance.

Para obtener más información, consulte [“Data Replication” de Oracle Solaris Cluster 4.3 Geographic Edition Overview](#).

Si desea usar la replicación basada en host sin el software Oracle Solaris Cluster Geographic Edition, consulte las instrucciones descritas en [Apéndice A, Ejemplo de despliegue: configuración de replicación de datos basada en host entre clusters con el software Availability Suite](#).

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 dentro de un cluster de campus:** este método se utiliza como sustitución para la creación de reflejo basada en host.

Tipo de replicación	Productos de replicación de datos
Replicación basada en almacenamiento	EMC Symmetrix Remote Data Facility (SRDF). Hitachi TrueCopy and Hitachi Universal Replicator, mediante la estructura de Oracle Solaris Cluster Geographic Edition.

- **Replicación basada en aplicaciones:** Oracle Data Guard es un ejemplo de software de replicación basada en aplicaciones. Este tipo de software se utiliza solamente para la recuperación ante desastres con la estructura de Geographic Edition con el fin de replicar una base de datos Oracle o una base de datos RAC de una sola instancia.

Para obtener más información, consulte [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard](#).

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

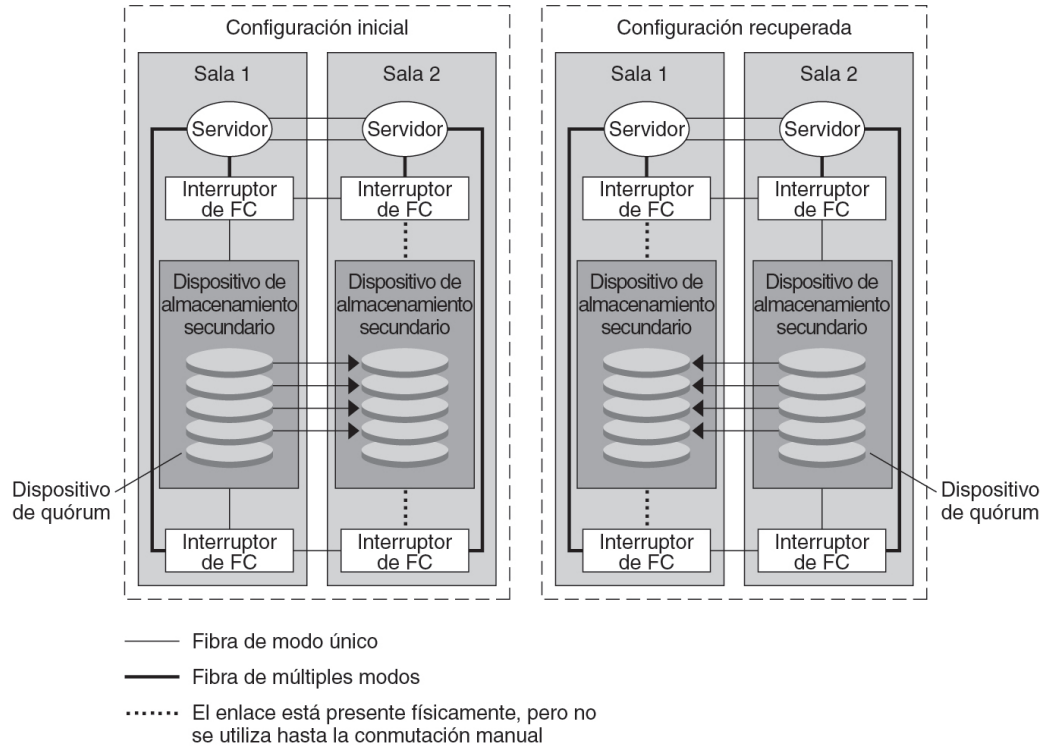
La replicación de datos basada en almacenamiento utiliza software instalado en el dispositivo de almacenamiento para gestionar la replicación dentro de un cluster, llamado cluster de campus. Dicho software es específico de su dispositivo de almacenamiento particular y si no se utiliza para la recuperación 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 utilice, puede usar el failover automático o manual con la replicación de datos basada en almacenamiento. Oracle Solaris Cluster admite el failover manual y automático de los replicadores con el software de EMC SRDF.

En esta sección, se describe la replicación de datos basada en almacenamiento como se la utiliza en un cluster de campus. En [Figura 1, “Configuración de dos salas con replicación de datos basada en almacenamiento”](#), se proporciona un ejemplo de configuración de dos salas donde los datos se replican entre dos matrices de almacenamiento. En esta configuración, la matriz de almacenamiento principal se 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 [Figura 1, “Configuración de dos salas con replicación de datos basada en almacenamiento”](#), se muestra que el dispositivo de quórum está en un volumen sin replicar. Un volumen replicado no se puede utilizar como dispositivo del quórum.

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



La replicación sincrónica basada en almacenamiento con EMC SRDF se admite en Oracle Solaris Cluster. La replicación asincrónica no se admite para 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 `data` o `status` para `Fence_level`, el dispositivo de almacenamiento principal rechaza las actualizaciones si estas últimas no se pueden copiar al dispositivo de almacenamiento remoto.

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

Para garantizar la integridad de los datos, utilice rutas múltiples y el paquete de 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.

- Si va a configurar el cluster para la conmutación por error automática, utilice la replicación sincrónica.
Para obtener instrucciones de configuración del cluster para failover automático de volúmenes replicados, consulte [“Configuración y administración de dispositivos replicados basados en el almacenamiento” \[111\]](#). Para obtener detalles sobre los requisitos para diseñar un cluster de campus, consulte [“Shared Data Storage” de Oracle Solaris Cluster Hardware Administration Manual](#).
- 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.
- 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 permiten que el usuario defina los 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:rdg-group-number
```

Esta entrada permite que el software del cluster automatice el movimiento de la aplicación entre los dos sitios SRDF síncronos. *rdg-group-number* en 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 4.3 Geographic Edition Overview](#).

- Oracle Real Application Clusters (Oracle RAC) no es compatible con SRDF 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 admiten conjuntos de discos de varios propietarios de Solaris Volume Manager para Sun 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 de campus” \[104\]](#) para obtener más información.

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

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 [Figura 1, “Configuración de dos salas con replicación de datos basada en almacenamiento”](#)), 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

Cuando use el software de EMC SRDF para la replicación de datos basada en almacenamiento, utilice 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” [109]
- “Configuración y administración de dispositivos replicados basados en el almacenamiento” [111]
- “Información general sobre la administración de sistemas de archivos de clusters” [125]
- “Administración de grupos de dispositivos” [125]
- “Administración de la configuración del protocolo SCSI para dispositivos de almacenamiento” [153]
- “Administración de sistemas de archivos de cluster” [159]
- “Administración de la supervisión de rutas de disco” [165]

Para obtener una descripción detallada de los procedimientos de este capítulo, consulte [Tabla 7, “Mapa de tareas: administrar grupos de dispositivos”](#).

Para obtener información conceptual relacionada con los dispositivos globales, los espacios de nombres globales, los grupos de dispositivos, la supervisión de las rutas de disco y el sistema de archivos del cluster, consulte [Oracle Solaris Cluster 4.3 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" para que pueda agregar, registrar y eliminar grupos de dispositivos mediante el

comando `metaset` de Solaris Volume Manager. Para obtener más información, consulte la página del comando `man metaset(1M)`.

El software de Oracle Solaris Cluster crea automáticamente un grupo de dispositivos de discos raw para cada dispositivo de cinta y disco 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 cambiar los permisos del dispositivo global `/dev/global/dsk/d3s0` y establecer un valor de 644, debe ejecutar el comando siguiente en todos los nodos del cluster:

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

Reconfiguración dinámica con dispositivos globales

Al completar operaciones de reconfiguración dinámica en dispositivos de disco y cinta de un cluster, debe tener en cuenta varios aspectos.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la función de reconfiguración dinámica de Oracle Solaris también se aplican a la reconfiguración dinámica de Oracle Solaris Cluster. La única excepción consiste en la inactividad del sistema operativo. Por lo tanto, consulte la documentación sobre la función de reconfiguración dinámica de Oracle Solaris *antes* de usar la función de reconfiguración dinámica con el software de Oracle Solaris Cluster. Debe prestar especial atención a los problemas que afectan a los dispositivos de E/S no conectados en red durante las operaciones de desconexión de reconfiguración dinámica.

- Oracle Solaris Cluster rechaza las operaciones de eliminación de tarjetas de reconfiguración dinámica en dispositivos activos del nodo principal. Las operaciones de reconfiguración dinámica pueden llevarse a cabo en dispositivos inactivos del nodo principal y en cualquier dispositivo de los nodos secundarios.
- Tras la operación de reconfiguración dinámica, el acceso a los datos del cluster no se modifica.
- Oracle Solaris Cluster rechaza las operaciones de reconfiguración dinámica que afectan la disponibilidad de los dispositivos de quórum. Si desea obtener más información, consulte [“Reconfiguración dinámica con dispositivos de quórum” \[175\]](#).



Atención - Si el nodo principal actual falla durante la operación de reconfiguración dinámica en un nodo secundario, la disponibilidad del cluster se ve afectada. 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, complete los siguientes pasos en el orden indicado.

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

Tarea	Para obtener instrucciones
1. Si en el nodo principal actual debe realizarse una operación de reconfiguración dinámica que afecta a un grupo de dispositivos activo, cambiar los nodos principal y secundario antes de llevar a cabo la operación de eliminación de reconfiguración dinámica en el dispositivo.	Conmutación al nodo primario de un grupo de dispositivos [150]
2. Efectuar la operación de eliminación de reconfiguración dinámica en el dispositivo que se va a eliminar.	Consulte la documentación que se incluye con el sistema.

Configuración y 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 admite el software EMC Symmetrix Remote Data Facility para la replicación basada en almacenamiento.

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



Atención - El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager o disco raw) debe ser idéntico al del grupo de dispositivos replicado.

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 6 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 [112]
Configurar el dispositivo DID	Configuración de dispositivos DID para la replicación con EMC SRDF [115]
Registrar el grupo replicado	Cómo agregar y registrar un grupo de dispositivos (Solaris Volume Manager) [133]
Verificar la configuración	Cómo comprobar la configuración del grupo de dispositivos globales replicados de EMC SRDF [117]
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 [122]

▼ Cómo configurar un grupo de replicación de EMC SRDF

- Antes de empezar**
- 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.



Atención - El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager o disco raw) debe ser idéntico al del grupo de dispositivos replicado.

1. **Asuma un rol que cuente con 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=device-group: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 4.3 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 replicas 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. Asegúrese de que las instancias de dispositivo DID especificadas sean réplicas de ellas mismas y que pertenezcan al grupo de replicación especificado.

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. **Asuma un rol que proporcione la autorización 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.**

```
# cldevice list -v
```

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 replication-type

Especifica el tipo de replicación. Para EMC SRDF, el tipo es **SRDF** .

-g replication-device-group

Especifica el nombre del grupo de dispositivos tal y como se muestra en el comando `syndg show`.

-d destination-instance

Especifica la instancia DID que corresponde al dispositivo RDF1.

source-instance

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 `sddidadm` para deshacer la combinación de dos dispositivos DID.

sddidadm -b device

-b device

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, realice los siguientes pasos adicionales.

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 `sddidadm -b` para eliminar la información existente. El último paso consiste en utilizar el comando `cldevice combine` 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 discos raw. Para obtener más información, consulte lo siguiente:

- [Cómo agregar y registrar un grupo de dispositivos \(Solaris Volume Manager\) \[133\]](#)
- [Adición y registro de un grupo de dispositivos \(disco básico\) \[135\]](#)
- [Adición y registro de un grupo de dispositivos replicado \(ZFS\) \[136\]](#)



Atención - El nombre del grupo de dispositivos de Oracle Solaris Cluster que ha creado (Solaris Volume Manager o disco raw) 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.**

```
# 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 de Oracle Solaris Cluster necesarios 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.

EJEMPLO 30 Creación de los pares de réplicas

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.

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

En todos los nodos del lado RDF1, ejecute los siguientes comandos

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

En todos los nodos del lado RDF2, ejecute los siguientes comandos

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

EJEMPLO 31 Comprobación de la configuración de la replicación de datos

Se ejecutan los siguientes comandos en un nodo del cluster.

```
# 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/c5t600604800018790002353594D303637d0s2 0067  RW  4315
}

```

Device Group RDF Information

...

symrdf -g dg1 establish

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

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

```

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

```

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

#

symrdf -g dg1 query

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

Target (R2) View					Source (R1) View					MODES
Standard	ST				LI	ST				
Logical	A				N	A				
Device	T	R1 Inv	R2 Inv	K	T	R1 Inv	R2 Inv		RDF Pair	
Dev	E	Tracks	Tracks	S	Dev	E	Tracks	Tracks	MDA STATE	
DEV001	0067	WD	0	0	RW	0067	RW	0	0 S.. Synchronized	
Total										
MB(s)	0.0		0.0		0.0		0.0			

Legend for MODES:

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

#

EJEMPLO 32 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`.

Ejecute estos comandos en el lado RDF1

```
# symdbg show dg1
```

Group Name: dg1

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

Device Group RDF Information

...

Obtenga los DID correspondientes

```
# cldevice list | grep c5t6006048000018790002353594D303637d0
217      pmoney1:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdisk/d217
```

```
217      pmoney2:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdsk/d217
```

Muestre los DID correspondientes

```
# cldevice show d217
```

```
=== DID Device Instances ===
```

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

```
#
```

Ejecute estos comandos en el lado RDF2

```
# symdg show dg1
```

```
Group Name:  dg1
```

```
Group Type          : RDF2      (RDFA)
```

```
...
```

```
Standard (STD) Devices (1):
```

```
{
-----
LdevName          PdevName          Sym          Cap
Dev  Att.  Sts
-----
DEV001            /dev/rdsk/c5t6006048000018799018253594D303637d0s2  0067    WD    4315
}
```

```
Device Group RDF Information
```

```
...
```

Obtenga los DID correspondientes

```
# cldevice list | grep c5t6006048000018799018253594D303637d0
```

```
108      pmoney4:/dev/rdsk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108
108      pmoney3:/dev/rdsk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108
```

Muestre los DID correspondientes

```
# cldevice show d108
```

```
=== DID Device Instances ===
```

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

```
#
```

EJEMPLO 33 Combinación de instancias DID

Desde el lado RDF2, escriba:

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

EJEMPLO 34 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 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. Este procedimiento recupera el grupo de dispositivos SRDF mediante la sincronización de los datos de la sala de secundaria original (este procedimiento utiliza phys-campus-2 para la sala secundaria) a la sala principal original (phys-campus-1). El procedimiento también 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, vea [Cómo agregar y registrar un grupo de dispositivos \(Solaris Volume Manager\) \[133\]](#).

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 de error, la sala principal de este procedimiento es phys-campus-1 y la sala secundaria es phys-campus-2.

1. **Inicie sesión en la sala primaria de un cluster de campus y asuma un rol que cuente con 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
```

Sugerencia - 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 35 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

Ningún comando de Oracle Solaris Cluster especial es necesario 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 utilizan UFS y también se pueden montar automáticamente durante el inicio. Los sistemas de archivos de cluster solo son visibles desde el nodo de un cluster global.

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` no se admite en los directorios que no están vacíos. Para obtener más información, consulte la página del comando `man unlink(1M)`.
- No se admite el comando `lockfs -d`. Como solución alternativa, utilice el comando `lockfs -n`.
- No puede volver a montar un sistema de archivos de cluster con la opción de montaje `directio` agregada en el momento del nuevo montaje.

Administración de grupos de dispositivos

A medida que cambian las necesidades del cluster, es posible que necesite agregar, eliminar o modificar los grupos de dispositivos del cluster. Oracle Solaris Cluster ofrece una interfaz interactiva, denominada `clsetup` que se puede utilizar para realizar 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.

Nota - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para colocar un grupo de dispositivos en línea o fuera de línea. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

El software de Oracle Solaris Cluster crea automáticamente un grupo de dispositivos de discos raw para cada dispositivo de cinta y disco del cluster. Sin embargo, los grupos de dispositivos del cluster permanecen fuera de línea hasta que se acceda a dichos grupos como dispositivos globales.



Atención - No ejecute `metaset -s setname -f -t` en un nodo del cluster que se inicie fuera del cluster si hay otros nodos activos como miembros del cluster y al menos uno de ellos es propietario del juego de discos.

TABLA 7 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 [127]
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 [128]
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> [130] Cómo migrar el espacio de nombres de dispositivos globales de un dispositivo <code>lofi</code> a una partición dedicada [131]
Agregar conjuntos de discos de Solaris Volume Manager y registrarlos como grupos de dispositivos mediante el comando <code>metaset</code>	Cómo agregar y registrar un grupo de dispositivos (Solaris Volume Manager) [133]
Agregar y registrar un grupo de dispositivos de discos raw con el comando <code>cldevicegroup</code>	Adición y registro de un grupo de dispositivos (disco básico) [135]
Agregar un grupo de dispositivos con nombre para ZFS mediante el comando <code>cldevicegroup</code>	Adición y registro de un grupo de dispositivos replicado (ZFS) [136]
Eliminar grupos de dispositivos de Solaris Volume Manager de la configuración con los comandos <code>metaset</code> y <code>metaclear</code>	“Cómo eliminar un grupo de dispositivos y anular su registro (Solaris Volume Manager)” [139]
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 [139]
Eliminar un nodo de un grupo de dispositivos de Solaris Volume Manager con el comando <code>metaset</code>	Cómo eliminar un nodo de un grupo de dispositivos (Solaris Volume Manager) [141]
Eliminar un nodo de un grupo de dispositivos de discos raw con el comando <code>cldevicegroup</code>	Eliminación de un nodo de un grupo de dispositivos de discos básicos [143]
Modificar las propiedades del grupo de dispositivos con <code>clsetup</code> para generar <code>cldevicegroup</code>	Cambio de propiedades de los grupos de dispositivos [145]
Mostrar las propiedades y los grupos de dispositivos con el comando <code>cldevicegroup show</code>	Enumeración de la configuración de grupos de dispositivos [149]
Cambiar el número deseado de nodos secundarios para un grupo de dispositivos con <code>clsetup</code> para generar <code>cldevicegroup</code>	Establecimiento del número de secundarios para un grupo de dispositivos [146]

Tarea	Instrucciones
Conmutar el primario de un grupo de dispositivos con el comando <code>cldevicegroup switch</code>	Conmutación al nodo primario de un grupo de dispositivos [150]
Poner un grupo de dispositivos en estado de mantenimiento con el comando <code>metaset</code>	Colocación de un grupo de dispositivos en estado de mantenimiento [151]

▼ Actualización del espacio de nombre de dispositivos globales

Al agregar un nuevo dispositivo global, actualice manualmente el espacio de nombres 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 lo tiene si el sistema de archivos `/global/.devices/node@nodeID` no está montado.

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.**
2. **En cada nodo del cluster, ejecute el comando `devfsadm`.**
Este comando puede ejecutarse simultáneamente en todos los nodos del cluster. Para obtener más información, consulte la página del comando [man `devfsadm\(1M\)`](#).
3. **Reconfigure el espacio de nombres.**

```
# cldevice populate
```

4. **Antes de intentar crear conjuntos de discos, compruebe que el comando `cldevice populate` haya finalizado su ejecución en cada uno de los nodos.**

El comando `cldevice` se llama a sí mismo de manera remota en todos los nodos, incluso al ejecutarse en un solo nodo. Para determinar si ha concluido el procesamiento del comando `cldevice populate`, ejecute el comando siguiente en todos los nodos del cluster.

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

ejemplo 36 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. **Asuma un rol que proporcione la autorización 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 [Reinicio de un nodo en el modo que no es de cluster \[95\]](#).

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
```

Asegúrese de que el dispositivo `lofi` esté desconectado

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

El comando no devuelve ninguna salida si el dispositivo está desconectado

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 de bucle de retorno (`lofi`), en lugar de crear uno para dispositivos globales en una partición dedicada.

Nota - Se admite ZFS para sistemas de archivos raíz, con una excepción importante: Si emplea una partición dedicada del disco de inicio para el sistema de archivos de los dispositivos globales, como sistema de archivos se debe usar solo 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`. Como alternativa, si en cambio se utiliza un dispositivo `lofi` para alojar el espacio de nombre de dispositivos globales, no hay límites en el uso de ZFS para sistemas de archivos raíz.

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` \[130\]](#)
- [Cómo migrar el espacio de nombres de dispositivos globales de un dispositivo `lofi` a una partición dedicada \[131\]](#)

▼ Migración del espacio de nombre de dispositivos globales desde una partición dedicada hasta un dispositivo de `lofi`

1. **Asuma el rol de usuario `root` en el nodo del cluster global cuya ubicación de espacio de nombres desea 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 [Reinicio de un nodo en el modo que no es de cluster \[95\]](#).

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 nombres de dispositivos globales.**

Esta entrada presenta una ruta de montaje que comienza con `/global/.devices/node@ID_nodo`.

6. **Desmonte la partición de dispositivos globales `/global/.devices/node@nodeID`.**

7. **Desactive y vuelva a activar los servicios SMF `globaldevices` and `scmountdev`**

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

A continuación, se crea un dispositivo `lofi` en `/.globaldevices` y se monta como 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.**

```
# 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.**

▼ **Cómo migrar el espacio de nombres de dispositivos globales de un dispositivo `lofi` a una partición dedicada**

1. **Asuma el rol de usuario `root` en el nodo del cluster global cuya ubicación de espacio de nombres desea 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 [Reinicio de un nodo en el modo que no es de cluster \[95\]](#).

3. **En un disco local del nodo, cree una nueva partición que cumpla con los requisitos siguientes:**
 - Un tamaño mínimo de 512 MB
 - Usa el sistema de archivos UFS
4. **Agregue una entrada al archivo `/etc/vfstab` para que la nueva partición se monte como sistema de archivos de dispositivos globales.**

- **Determine el ID de nodo del nodo actual.**

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

- **Cree la nueva entrada en el archivo `/etc/vfstab` con el formato siguiente:**

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

Por ejemplo, si la partición que decide usar es `/dev/did/rdisk/d5s3`, la entrada nueva que se agrega al archivo `/etc/vfstab` sería la siguiente:

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

5. **Desmonte la partición de dispositivos globales** `/global/.devices/node@nodeID`.

6. **Quite el dispositivo de `lofi` asociado con el archivo `/.globaldevices`.**

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

7. **Elimine el archivo `/.globaldevices`.**

```
# rm /.globaldevices
```

8. **Desactive y vuelva a activar los servicios SMF `globaldevices` and `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 y complete el espacio de nombres de dispositivos globales.**

- a. **Desde un nodo del cluster, complete los espacios de nombres de dispositivos globales.**

```
# cldevice populate
```

- b. **Antes de pasar a realizar ninguna otra acción en cualquiera de los nodos, compruebe que el proceso concluya en todos los nodos del 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 un disco raw.

▼ Cómo agregar y registrar 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.



Atención - El nombre del grupo de dispositivos de Oracle Solaris Cluster que se crea (Solaris Volume Manager o disco raw) debe ser idéntico al nombre del grupo de dispositivos replicados.

1. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en uno de los nodos conectado a los discos donde está 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
```

<code>-s diskset</code>	Especifica el conjunto de discos que se va a crear.
<code>-a -h nodelist</code>	Agrega la lista de nodos que pueden controlar el conjunto de discos.
<code>-M</code>	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 manera predeterminada, se obtiene un nodo secundario, independientemente del número de nodos incluidos en dicho grupo de dispositivos. Puede cambiar el número deseado de nodos secundarios con la utilidad `clsetup` una vez creado el grupo de dispositivos. Consulte [Establecimiento del número de secundarios para un grupo de dispositivos \[146\]](#) 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 DID completo, que tiene el formato `/dev/did/rdisk/dN`, al agregar una unidad a un conjunto de discos.

En el ejemplo siguiente, las entradas del dispositivo DID `/dev/did/rdisk/d3` indican que `phys-schost-1` y `phys-schost-2` comparten la unidad.

```
=== DID Device Instances ===
DID Device Name:                /dev/did/rdisk/d1
Full Device Path:                phys-schost-1:/dev/rdisk/c0t0d0
DID Device Name:                /dev/did/rdisk/d2
Full Device Path:                phys-schost-1:/dev/rdisk/c0t6d0
DID Device Name:                /dev/did/rdisk/d3
Full Device Path:                phys-schost-1:/dev/rdisk/c1t1d0
Full Device Path:                phys-schost-2:/dev/rdisk/c1t1d0
...
```

6. Agregue las unidades al conjunto de discos.

Utilice el nombre completo de la ruta de DID.

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

`-s setname` Especifique 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`) al agregar una unidad a un conjunto de discos. Ya que el nombre de dispositivo de nivel inferior es un nombre local y no único para todo el 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 37 Agregación de un grupo de dispositivos de Solaris Volume Manager

En el siguiente ejemplo, 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`, y se verifica que el grupo de dispositivos haya sido creado.

```
# 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 `cldevicegroup`.



Atención - Si crea un grupo de dispositivos en dispositivos replicados, el nombre del grupo de dispositivos que crea (Solaris Volume Manager o disco raw) debe ser idéntico al del grupo de dispositivos replicados.

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 los dispositivos `dN` y `dX`.

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

2. Cree el grupo de dispositivos de disco básico con los dispositivos deseados.

El comando siguiente crea un grupo de dispositivos globales, *raw-disk-dg*, que contiene los dispositivos *dN* y *dX*.

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

3. Compruebe el grupo de dispositivos raw creado.

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

▼ Adición y registro de un grupo de dispositivos replicado (ZFS)

Utilice este procedimiento para crear un grupo de dispositivos ZFS replicado gestionado por HAStoragePlus.

Para crear una agrupación de almacenamiento ZFS (zpool) que no use HAStoragePlus, consulte [Cómo configurar una agrupación de almacenamiento ZFS local sin HAStoragePlus \[137\]](#).

Antes de empezar

Para replicar ZFS, debe crear un grupo de dispositivos designado y mostrar 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 raw) debe ser idéntico al del grupo de dispositivos replicado.

1. Elimine los grupos de dispositivos predeterminados que se correspondan con los dispositivos de zpool.

Por ejemplo, si dispone de un zpool denominado *mypool* que contiene dos dispositivos, */dev/did/dsk/d2* y */dev/did/dsk/d13*, debe 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.

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

5. Cree un recurso hasp-rs en el grupo de recursos creado en el [Paso 4](#) y configure la propiedad globaldevicepaths en un grupo de dispositivos del tipo disco raw.

Este grupo de dispositivos se creó en el [Paso 2](#).

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

6. Establezca el valor +++ de la propiedad rg_affinities de este grupo de recursos en el grupo de recursos creado en el [Paso 4](#).

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

7. Cree un recurso HAStoragePlus (hasp-rs) para el zpool creado en el [Paso 3](#) en el grupo de recursos creado en el [Paso 4](#) o en el [Paso 6](#).

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

```
# clresource create -g oracle-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.

▼ Cómo configurar una agrupación de almacenamiento ZFS local sin HAStoragePlus

En este procedimiento, se describe cómo configurar una agrupación de almacenamiento ZFS (zpool) en un dispositivo local sin configurar un recurso HAStoragePlus.

Nota - Para configurar una zpool local que use un recurso HAStoragePlus, consulte [Adición y registro de un grupo de dispositivos replicado \(ZFS\) \[136\]](#).

1. Muestre las asignaciones DID e identifique el dispositivo local que se va a usar.

Seleccione un dispositivo que muestre solo el nodo de cluster que usará la nueva zpool. Anote el nombre de dispositivo `cNtXdY` y el nombre de dispositivo DID `/dev/did/rdisk/dN`.

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

En el siguiente ejemplo, las entradas de los dispositivos DID `/dev/did/rdisk/d1` y `/dev/did/rdisk/d2` muestran que solo `phys-schost-1` usa esas unidades. Para los ejemplos de este procedimiento, se usará el dispositivo DID `/dev/did/rdisk/d2` con el nombre de dispositivo `c0t6d0` y se lo configurará para el nodo de cluster `phys-schost-1`.

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

2. Determine el nombre de grupo de dispositivos del dispositivo DID que se selecciona para la zpool.

En la siguiente salida de ejemplo, se muestra que `dsk/d2` es el nombre del grupo de dispositivos del dispositivo DID `/dev/did/rdisk/d2s2`. A menudo, aunque no siempre, hay una relación en la que el nombre del grupo de dispositivos es parte del nombre de dispositivo DID. Este grupo de dispositivos tiene solo un nodo, `phys-schost-1`, en su lista de nodos.

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

3. Establezca la propiedad `localonly` para el dispositivo DID.

Especifique el nombre del grupo de dispositivos identificado en el [Paso 2](#). Si desea desactivar el aislamiento del dispositivo, también incluya `default_fencing=nofencing` en el comando.

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

Para obtener más información sobre las propiedades `cldevicegroup`, consulte la página del comando [man cldevicegroup\(1CL\)](#).

4. Compruebe la configuración del dispositivo.

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

5. Cree la zpool.

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

6. (Opcional) Cree un juego de datos ZFS.

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

7. Verifique la nueva zpool.

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

Mantenimiento de grupos de dispositivos

Puede llevar a cabo una serie de tareas administrativas para los grupos de dispositivos. Algunas de estas tareas también se pueden realizar mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Cómo eliminar un grupo de dispositivos y anular su registro (Solaris Volume Manager)

Los grupos de dispositivos son conjuntos de discos de Solaris Volume Manager que se registraron en 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 [Solaris Volume Manager Administration Guide](#) para conocer los pasos para eliminar un juego de discos.

▼ Eliminación de un nodo de todos los grupos de dispositivos

Use este procedimiento para eliminar un nodo de cluster de todos los grupos de dispositivos que incluyen el nodo en las listas de posibles nodos principales.

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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo que va a eliminar de todos los grupos de dispositivos como posible nodo principal.**

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 `Device group node list` para cada grupo de 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 [Cómo eliminar un nodo de un grupo de dispositivos \(Solaris Volume Manager\) \[141\]](#) para todos los grupos de dispositivos de dicho tipo.**

4. **Determine los grupos de discos de dispositivos básicos de los cuales el nodo que se va a eliminar forma parte como miembro.**

```
# cldevicegroup list -v
```

5. **Si alguno de los grupos de dispositivos que figuran en la lista del [Paso 4](#) pertenece a los tipos de grupos `Disk` o `Local_Disk`, siga los pasos descritos en [Eliminación de un nodo de un grupo de dispositivos de discos básicos \[143\]](#) para 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
```

▼ Cómo eliminar un nodo de un grupo de dispositivos (Solaris Volume Manager)

Siga este procedimiento para eliminar un nodo de cluster de la lista de potenciales nodos principales 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.



Atención - No ejecute `metaset -s setname -f -t` en un nodo del cluster que se inicie fuera del cluster si hay otros nodos activos como miembros del cluster y al menos uno de ellos es propietario del juego 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. **Asuma el rol `root` 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` Suprime del grupo de dispositivos los nodos identificados con `-h`.

`-h nodelist` 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** para cada grupo de dispositivos del cual se elimina el nodo como posible nodo principal.

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 38 Eliminación de un nodo de un grupo de dispositivos (Solaris Volume Manager)

En el ejemplo siguiente, se muestra cómo eliminar el nombre de host `phys-schost-2` de una configuración de grupos de dispositivos. Este ejemplo elimina `phys-schost-2` como posible nodo principal para el 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 el grupo de dispositivos de Solaris Volume Manager para el nodo

```
# 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 qué nodo es el primario del grupo de dispositivos

```
# 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

Asuma el rol root en el nodo que actualmente es propietario del grupo de dispositivos

Elimine el nombre de host del grupo de dispositivos

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

Verifique la eliminación del nodo

```
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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.read` y `solaris.cluster.modify` en un nodo del cluster que no sea el nodo que se eliminará.**
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. **Desactive 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)` para obtener más información sobre la propiedad `localonly`.

4. **Compruebe que haya desactivado 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 desactivada 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 raw 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 39 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).

Identifique los grupos de dispositivos conectados al nodo que se desea eliminar y determine cuáles son grupos de dispositivos de discos raw

```
phys-schost-1# cldevicegroup show -n phys-schost-2 -t rawdisk -v +
Device Group Name:                dsk/d4
Type:                             Disk
failback:                         false
Node List:                        phys-schost-2
preferenced:                      false
localonly:                        false
autogen                           true
numsecondaries:                   1
device names:                     phys-schost-2
Device Group Name:                dsk/d1
Type:                             SVM
failback:                         false
Node List:                        pbrave1, pbrave2
preferenced:                      true
localonly:                        false
autogen                           true
numsecondaries:                   1
diskset name:                     ms1
(dsk/d4) Device group node list:  phys-schost-2
(dsk/d2) Device group node list:  phys-schost-1, phys-schost-2
(dsk/d1) Device group node list:  phys-schost-1, phys-schost-2
```

Desactive el indicador localonly en todos los discos locales del nodo

```
phys-schost-1# cldevicegroup set -p localonly=false dsk/d4
```

Verifique que el indicador localonly esté desactivado phys-schost-1# **cldevicegroup show -n phys-**

```
schost-2 -t rawdisk +
```

```
(dsk/d4) Device group type:        Disk
(dsk/d8) Device group type:        Local_Disk
```

Elimine el nodo de todos los grupos de dispositivos de discos raw

```
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 principal de un grupo de dispositivos se basa en la configuración de un atributo de preferencias 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 desactiva el atributo `preferenced`, el atributo `failback` también se desactiva automáticamente. Sin embargo, si intenta activar o volver a activar el atributo `preferenced`, debe elegir entre activar o desactivar el atributo `failback`.

Si el atributo `preferenced` se activa o se vuelve a activar, se solicitará que reestablecer el orden de los nodos en la lista de preferencias de propiedades primarias.

En este procedimiento se usa la utilidad `clsetup` para configurar o anular la configuración de los atributos `preferenced` y `failback` para 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. **Asuma un rol 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 para la opción de los grupos de dispositivos y los volúmenes.**

Se muestra el menú Grupos de dispositivos.
4. **Para modificar las propiedades clave de un grupo de dispositivos, escriba el número para la opción de modificación de las propiedades clave de un grupo de dispositivos de Solaris Volume Manager.**

Se muestra el menú Cambiar las propiedades esenciales.

5. **Para modificar la propiedad de un grupo de dispositivos, escriba el número correspondiente a la opción para modificar las preferencias o las propiedades de failback.**

Siga las instrucciones para definir las opciones `preferenced` y `failback` para un grupo de dispositivos.

6. **Compruebe que se hayan modificado los atributos del grupo de dispositivos.**

Busque la información del grupo de dispositivos mostrada por el comando siguiente.

```
# cldevicegroup show -v devicegroup
```

ejemplo 40 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 el nodo principal falla. 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 [Capítulo 3, “Key Concepts for System Administrators and Application Developers”](#) de *Oracle Solaris Cluster 4.3 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 definir la propiedad `numsecondaries` para todos los tipos de grupos de dispositivos. Consulte la página del comando `man 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.

1. **Asuma un rol 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 cambiar el número deseado de nodos secundarios, escriba el número para la opción de cambio de 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
failback:              yes
Node List:              phys-schost-1, phys-schost-2, phys-schost-3
preferenced:            yes
numsecondaries:         1
diskgroup names:       dg-schost-1
```

Nota - Entre los cambios de configuración del grupo de dispositivos 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 \[127\]](#).

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 41 Modificación del número deseado de nodos secundarios (Solaris Volume Manager)

En el ejemplo siguiente se muestra el comando `cldevicegroup` que genera `clsetup` al configurar el número de nodos secundarios para un grupo de dispositivos (`dg-schost-1`). En este ejemplo se supone que el grupo de discos y el volumen ya se habían creado.

```
# cldevicegroup set -p numsecondaries=1 dg-schost-1
# cldevicegroup show -v dg-schost-1

=== Device Groups ===

Device Group Name:      dg-schost-1
Type:                  SVM
failback:              yes
Node List:              phys-schost-1, phys-schost-2
preferenced:            yes
numsecondaries:         1
diskset names:         dg-schost-1
```

ejemplo 42 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 necesita tener el rol `root` para mostrar 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.

Interfaz de explorador de Oracle Solaris Cluster Manager

Consulte [Capítulo 13, Uso de la interfaz de explorador de Oracle Solaris Cluster Manager](#) para obtener más información.

`cldevicegroup show`

Use `cldevicegroup show` para enumerar la configuración de todos los grupos de dispositivos del cluster.

`cldevicegroup show devicegroup`

Use `cldevicegroup show grupo_dispositivos` para enumerar la configuración de un solo grupo de dispositivos.

`cldevicegroup status grupo_dispositivos`

Use `cldevicegroup status grupo_dispositivos` para determinar el estado de un solo grupo de dispositivos.

`cldevicegroup status +`

Use `cldevicegroup status +` 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 43 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 44 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.

Nota - También puede colocar en línea un grupo de dispositivos inactivo mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Consulte la ayuda en pantalla de Oracle Solaris Cluster Manager para obtener más información. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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. **Asuma un rol que proporcione la autorización 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 nodename` 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 45 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 ese grupo se conecte automáticamente cada vez que se accede a uno de sus dispositivos. Conviene poner un grupo de dispositivos en estado de mantenimiento al finalizar procedimientos de reparación que requieran consentimiento para todas las actividades de E/S hasta que terminen las operaciones de reparación. Asimismo, poner un grupo de dispositivos en estado de mantenimiento ayuda a

impedir la pérdida de datos, ya que garantiza que el grupo de dispositivos no se ponga en línea en un nodo mientras el juego de discos se está reparando en otro nodo.

Para obtener instrucciones sobre cómo restaurar un conjunto de discos dañado, consulte [“Restauración de un conjunto de discos dañado” \[282\]](#).

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.

Nota - Puede colocar un grupo de dispositivos activos en línea mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Consulte la ayuda en pantalla de Oracle Solaris Cluster Manager para obtener más información. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. Ponga el grupo de dispositivos en estado de mantenimiento.

a. Si el grupo de dispositivos está activado, desactívelo.

```
# 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 juego de discos, importe manualmente ese juego de discos.

```
# metaset -C take -f -s diskset
```



Atención - 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. Libere la propiedad del juego de discos.



Atención - Antes de anular el estado de mantenimiento del grupo de dispositivos, debe liberar la propiedad del juego de discos. Si hay un error al liberar la propiedad, pueden perderse datos.

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

5. Ponga en línea el grupo de dispositivos.

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

ejemplo 46 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.

```
[Ponga el grupo de dispositivos en estado de mantenimiento.]
# cldevicegroup disable dg-schost-1
# cldevicegroup offline dg-schost-1
[Si es necesario, importe el juego de discos manualmente].
For Solaris Volume Manager:
# metaset -C take -f -s dg-schost-1
[Complete todos los procedimientos de reparación necesarios.]
[Libere la propiedad.]
Para Solaris Volume Manager:
# metaset -C release -s dg-schost-1
[Ponga en línea el grupo de dispositivos.]
# 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, anular la configuración de un dispositivo:

- [Visualización de la configuración del protocolo SCSI global predeterminado de todos los dispositivos de almacenamiento \[154\]](#)
- [Visualización del protocolo SCSI de un solo dispositivo de almacenamiento \[155\]](#)
- [Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento \[155\]](#)
- [Modificación del protocolo de protección en un solo dispositivo de almacenamiento \[157\]](#)

▼ 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. **Asuma un rol que proporcione la autorización 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 47 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
clusterid:	0x4FES2C888
installmode:	disabled
heartbeat_timeout:	10000
heartbeat_quantum:	1000
private_netaddr:	172.16.0.0
private_netmask:	255.255.111.0
max_nodes:	64
max_privatenets:	10
udp_session_timeout:	480
concentrate_load:	False
global_fencing:	prefer3
Node List:	phys-racerxx-1, phys-racerxx-2

▼ Visualización del protocolo SCSI de un solo dispositivo de almacenamiento

`phys-schost#` refleja un indicador de 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.read`.**
2. **Desde cualquier nodo, visualice la configuración del protocolo SCSI del dispositivo de almacenamiento.**

```
# cldevice show device
```

device Nombre de la ruta del dispositivo o un nombre de dispositivo.

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

ejemplo 48 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/rdsk/c4t8d0
```

```
=== DID Device Instances ===
```

DID Device Name:	/dev/did/rdsk/d3
Full Device Path:	phappy1:/dev/rdsk/c4t8d0
Full Device Path:	phappy2:/dev/rdsk/c4t8d0
Replication:	none
default_fencing:	global

▼ Modificación de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento

La protección se puede activar o desactivar globalmente para todos los dispositivos de almacenamiento conectados a un cluster. La configuración de aislamiento predeterminado 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 \[157\]](#) para modificar la configuración predeterminada de un solo dispositivo.



Atención - 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Establezca el protocolo de protección para todos los dispositivos de almacenamiento que no sean de quórum.**

```
cluster set -p global_fencing={pathcount | prefer3 | nofencing | nofencing-noscrub}
```

```
-p global_fencing
```

Establece el algoritmo de protección predeterminado global para todos los dispositivos compartidos.

```
prefer3
```

Emplea el protocolo SCSI-3 para los dispositivos que cuenten con más de dos rutas.

pathcount

Determina el protocolo de protección mediante el número de rutas de DID conectadas al dispositivo compartido. La configuración de pathcount se utiliza con los dispositivos de quórum.

nofencing

Desactiva la protección con la configuración del estado de todos los dispositivos de almacenamiento.

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.

ejemplo 49 Establecimiento de la configuración del protocolo de protección global predeterminado para todos los dispositivos de almacenamiento

En el ejemplo siguiente se establece en el protocolo SCSI-3 el protocolo de protección para todos los dispositivos de almacenamiento del 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.



Atención - 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Establezca el protocolo de protección del dispositivo de almacenamiento.**

```
cldevice set -p default_fencing ={pathcount | \  
scsi3 | global | nofencing | nofencing-noscrub} device
```

`-p default_fencing`

Modifica la propiedad `default_fencing` del dispositivo.

`pathcount`

Determina el protocolo de protección mediante el número de rutas de DID conectadas al dispositivo compartido.

`scsi3`

Usa el protocolo SCSI-3.

`global`

Usa la configuración de protección predeterminada global. La configuración global se utiliza con los dispositivos que no son de quórum.

`nofencing`

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 este 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` solo 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 50 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=no fencing 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 8 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 [159]
Eliminar un sistema de archivos de cluster	Eliminación de un sistema de archivos de cluster [162]
Comprobar los puntos de montaje globales de un cluster para verificar la coherencia entre los nodos	Comprobación de montajes globales en un cluster [164]

▼ 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.



Atención - 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:

- El privilegio del rol de usuario root se establece en un nodo del cluster.
- Software de administración de volúmenes instalado y configurado en el cluster.
- Existe un grupo de dispositivos (como un grupo de dispositivos de Solaris Volume Manager) o porción de disco de bloques donde poder crear el sistema de archivos de cluster.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para agregar un sistema de archivos de cluster a un cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Si usó Oracle Solaris Cluster Manager para instalar servicios de datos, ya hay uno o varios sistemas de archivos de cluster si había suficientes discos compartidos donde poder crear los sistemas de archivos 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. Asuma el rol root en cualquier nodo del cluster.

Sugerencia - Para crear sistemas de archivos con mayor rapidez, asuma el rol root en el nodo principal del dispositivo global para el que desea crear un sistema de archivos.

2. Cree un sistema de archivos UFS con el comando `newfs`.



Atención - Todos los datos de los discos se destruyen al crear un sistema de archivos. Compruebe que haya especificado el nombre de dispositivo de disco correcto. Si se especifica un nombre equivocado, podría borrar datos que no tuviera previsto eliminar.

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

La tabla siguiente muestra ejemplos de nombres para el argumento *dispositivo_discos_básicos*. Cada administrador de volúmenes aplica sus propias convenciones de asignación de nombres.

Administrador de volúmenes	Nombre de dispositivo de disco de ejemplo	Descripción
Solaris Volume Manager	<code>/dev/md/nfs/rdisk/d1</code>	Dispositivo de disco raw d1 dentro del juego de discos nfs
Ninguno	<code>/dev/global/rdisk/d1s3</code>	Dispositivo de disco raw d1s3

3. En cada nodo del cluster, cree un directorio de puntos de montaje para el sistema de archivos de cluster.

Todos los nodos deben tener un punto de montaje, aunque no se acceda al sistema de archivos de cluster en un nodo determinado.

Sugerencia - 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/mount-point/
```

device-group

Nombre del directorio correspondiente al nombre del grupo de dispositivos que contiene el dispositivo.

mount-point

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 de comando `man vfstab(4)` para obtener información detallada.

- 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 del cluster sea idéntica en todos los nodos.**
- d. **Asegúrese de 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 el escenario donde `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` puede iniciar y montar `/global/oracle/logs/` únicamente después de que `phys-schost-1` inicia y monta `/global/oracle/`.

5. En cualquier nodo del cluster, ejecute la utilidad de verificación de 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 se produce ningún error, no se devuelve ninguna salida.

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

6. Monte el sistema de archivos del 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 el comando `df` o `mount` para mostrar los sistemas de archivos montados. Para obtener más información, consulte la página del comando `man df(1M)` o la página del comando `man mount(1M)`.

▼ 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. Los sistemas de archivos de cluster no se desmontan al ejecutar `shutdown` para detener un único 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:

- El privilegio del rol de usuario `root` se establece 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.

Nota - También puede eliminar un sistema de archivos de cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. Asuma el rol root en cualquier nodo del 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 suprimir la entrada del sistema de archivos de cluster que se eliminará.

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) Elimine el dispositivo de disco group/metadevice/volume/plex.

Para obtener más información, consulte la documentación del administrador de volúmenes.

ejemplo 51 Eliminación de un sistema de archivos de cluster

En el ejemplo siguiente, se elimina un sistema de archivos de cluster UFS montado en el metadispositivo o el volumen `/dev/md/oracle/rdisk/d1` de Solaris Volume Manager.

```
# 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

    En cada nodo, elimine la entrada resaltada
# pfedit /etc/vfstab
#device          device          mount  FS      fsck    mount  mount
#to mount        to fsck      point  type    pass   at boot options
#
/dev/md/oracle/dsk/d1 /dev/md/oracle/rdisk/d1 /global/oracle/d1 ufs 2 yes global,logging
```

Guarde y salga

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)` permite verificar la sintaxis de las entradas de sistemas de archivos de cluster en el archivo `/etc/vfstab`. Si no hay ningún error, no se 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. Asuma el rol root en cualquier nodo del 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.

La siguiente información adicional está disponible:

Tema	Información
Información conceptual sobre daemon de supervisión de rutas de disco	Capítulo 3, “Key Concepts for System Administrators and Application Developers” de Oracle Solaris Cluster 4.3 Concepts Guide
Descripción de las opciones del comando <code>cldevice</code> y los comandos relacionados	Página del comando <code>man cldevice(1CL)</code>
Ajuste del daemon <code>scdpmd</code>	Página del comando <code>man scdpmd.conf(4)</code>
Errores registrados que el daemon <code>syslogd</code> informa	Página del comando <code>man syslogd(1M)</code>

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`. La supervisión de rutas de disco se detiene automáticamente al quitar los dispositivos de un nodo con los comandos de Oracle Solaris Cluster.

TABLA 9 Mapa de tareas: administrar la supervisión de rutas de disco

Tarea	Instrucciones
Supervisar una ruta de disco	Supervisión de una ruta de disco [166]
Anular la supervisión de una ruta de disco	Anulación de la supervisión de una ruta de disco [167]
Imprimir el estado de rutas de disco erróneas correspondientes a un nodo	Impresión de rutas de disco erróneas [168]
Supervisar rutas de disco desde un archivo	Supervisión de rutas de disco desde un archivo [169]
Activar o desactivar el re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas	Activación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas [171] Desactivación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas [172]
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 [169]

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.

`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 - También puede activar la supervisión de una ruta de disco mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma un rol que proporcione la autorización 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 52 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 del 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 53 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` es 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 54 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.

`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 - También puede desactivar la supervisión de una ruta de disco mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma un rol que proporcione la autorización 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.**
3. **Anule la supervisión de las correspondientes rutas de disco en cada uno de los nodos.**

```
# cldevice status device
```

```
# cldevice unmonitor -n node disk
```

ejemplo 55 Anulación de la supervisión de una ruta de disco

En el ejemplo siguiente se anula la supervisión de la ruta de disco `schost-2:/dev/did/rdisk/d1` y se imprimen las rutas de disco de todo el cluster, con sus estados.

```
# cldevice unmonitor -n schost2 /dev/did/rdisk/d1
# cldevice status -n schost2 /dev/did/rdisk/d1
```

Device Instance	Node	Status
-----	----	-----
/dev/did/rdisk/d1 schost-2	Unmonitored	

▼ Impresión de rutas de disco erróneas

Siga este procedimiento para imprimir las rutas de disco erróneas correspondientes a un cluster.

1. **Asuma el rol `root` en cualquier nodo del cluster.**
2. **Imprima las rutas de disco erróneas de todo el cluster.**

```
# cldevice status -s fail
```

ejemplo 56 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 solo 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.

`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. **Asuma un rol que proporcione la autorización 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 configurationfile` 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 desea supervisar y defina 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 57 Supervisión de 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 mediante un archivo XML. El archivo XML `deviceconfig` muestra que la ruta entre `phys-schost-2` y `d3` no está bajo supervisión en la actualidad.

Exporte la configuración actual del cluster

```
# cldevice export -o deviceconfig
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
...
```

```
<deviceList readonly="true">
<device name="d3" ctd="c1t8d0">
<devicePath nodeRef="phys-schost-1" monitored="true"/>
<devicePath nodeRef="phys-schost-2" monitored="false"/>
</device>
</deviceList>
</cluster>
```

Supervise la ruta estableciendo el atributo de supervisión en true

```
<?xml version="1.0"?>
<!DOCTYPE cluster SYSTEM "/usr/cluster/lib/xml/cluster.dtd">
<cluster name="brave_clus">
...
<deviceList readonly="true">
<device name="d3" ctd="c1t8d0">
<devicePath nodeRef="phys-schost-1" monitored="true"/>
<devicePath nodeRef="phys-schost-2" monitored="true"/>
</device>
</deviceList>
</cluster>
```

Lea el archivo y active la supervisión

```
# cldevice monitor -i deviceconfig
```

Compruebe que la ruta de dispositivo se esté supervisando

```
# cldevice status
```

Véase también Para obtener más información sobre cómo exportar la configuración de un cluster y cómo usar el archivo XML resultante para definir la configuración de un cluster, consulte las páginas del comando `man cluster(1CL)` y `clconfiguration(5CL)`.

▼ Activación del rearranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas

Al activar esta función, un nodo se rearranca automáticamente siempre que se cumplan las condiciones siguientes:

- Fallan todas las rutas de disco compartido supervisadas del nodo.
- Se puede acceder a uno de los discos como mínimo desde un nodo diferente del cluster.

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 re arrancar el nodo pero falla posteriormente, el nodo re arranca automáticamente otra vez.

Al activar la propiedad `reboot_on_path_failure`, los estados de las rutas de disco local no se tienen en cuenta al determinar si es necesario re arrancar un nodo. Esto afecta sólo a discos compartidos supervisados.

Nota - También puede editar la propiedad de nodo `reboot_on_path_failure` mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **En cualquiera de los nodos del cluster, asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Active 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 +
```

▼ Desactivación del re arranque automático de un nodo si fallan todas las rutas de disco compartido supervisadas

Si se desactiva esta función y fallan todas las rutas de disco compartido supervisadas de un nodo, dicho nodo *no* re arranca automáticamente.

1. **En cualquiera de los nodos del cluster, asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Desactive 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

En este capítulo, se presentan 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 4.3 Concepts Guide](#).

- [“Administración de dispositivos de quórum” \[173\]](#)
- [“Administración de servidores de quórum de Oracle Solaris Cluster” \[193\]](#)

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 utilizar el comando `clquorum` para realizar todos los procedimientos administrativos de dispositivos de quórum. Además, puede llevar a cabo algunos procedimientos mediante la utilidad interactiva `clsetup` o la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [“Cómo acceder a Oracle Solaris Cluster Manager” de Guía de administración del sistema de Oracle Solaris Cluster 4.3](#).

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 Oracle Solaris Cluster Manager. Para obtener más información, consulte las páginas del comando `man clquorum(1CL)` y `clsetup(1CL)`.

Al trabajar con dispositivos de quórum, tenga en cuenta las directrices siguientes:

- Todos los comandos de quórum se deben ejecutar en un nodo de 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 Oracle Solaris. Si se agrega un dispositivo de quórum configurado a una agrupación de almacenamiento ZFS, 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 10 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” [175]
Eliminar un dispositivo del quórum de un cluster mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	Eliminación de un dispositivo de quórum [182]
Eliminar el último dispositivo del quórum de un cluster mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	Eliminación del último dispositivo de quórum de un cluster [183]
Reemplazar un dispositivo de quórum de un cluster mediante los procedimientos de agregar y quitar	Sustitución de un dispositivo de quórum [185]
Modificar una lista de dispositivos de quórum mediante los procedimientos de agregar y quitar	Modificación de una lista de nodos de dispositivo de quórum [186]
Poner un dispositivo del quórum en estado de mantenimiento mediante la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	Colocación de un dispositivo de quórum en estado de mantenimiento [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 la utilidad <code>clsetup</code> (para generar <code>clquorum</code>)	Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento [189]

Tarea	Para obtener instrucciones
Enumerar los dispositivos del quórum y los recuentos de votos mediante el comando <code>clquorum</code>	Enumeración de una lista con la configuración de quórum [190]

Reconfiguración dinámica con dispositivos de quórum

Debe tener en cuenta algunos aspectos al completar operaciones de reconfiguración dinámica en dispositivos de quórum de un cluster.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la función de reconfiguración dinámica de Oracle Solaris también se aplican a la reconfiguración dinámica de Oracle Solaris Cluster, excepto la inactividad del sistema operativo. Por lo tanto, consulte la documentación sobre la función de reconfiguración dinámica de Oracle Solaris *antes* de usar la función de reconfiguración dinámica con el software de Oracle Solaris Cluster. Debe prestar especial atención a los problemas que afectan a los dispositivos de E/S no conectados en red durante las operaciones de desconexión de reconfiguración dinámica.
- Oracle Solaris Cluster rechaza las operaciones de eliminación de tarjetas de reconfiguración dinámica efectuadas cuando hay una interfaz configurada para un dispositivo de quórum.
- Si la operación de reconfiguración dinámica pertenece a un dispositivo activo, Oracle Solaris Cluster rechaza la operación e identifica 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 11 Mapa de tareas: reconfiguración dinámica con dispositivos de quórum

Tarea	Para obtener instrucciones
1. Activar un nuevo dispositivo de quórum para sustituir el que se va a eliminar.	“Adición de un dispositivo de quórum” [175]
2. Desactivar el dispositivo de quórum que se va a eliminar	Eliminación de un dispositivo de quórum [182]
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 muestran 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 recuentos de votos de quórum necesarios para el cluster, las configuraciones de quórum recomendadas y el aislamiento de errores, consulte [“Quorum and Quorum Devices” de Oracle Solaris Cluster 4.3 Concepts Guide](#).



Atención - 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 de Oracle Solaris Cluster admite los siguientes tipos de dispositivos de quórum:

- LUN compartidos desde:
 - Disco SCSI compartido
 - Almacenamiento SATA (Serial Attached Technology Attachment)
 - Oracle ZFS Storage Appliance
 - Dispositivos NAS admitidos.
- Servidor de quórum de Oracle Solaris Cluster

En las secciones siguientes se presentan procedimientos para agregar estos dispositivos:

- [Adición de un dispositivo de quórum de disco compartido \[177\]](#)
- [Adición de un dispositivo de quórum de servidor de quórum \[179\]](#)

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.



Atención - 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 datos del usuario o que sea miembro de un grupo de dispositivos. El protocolo que utiliza el subsistema de quórum con un disco compartido puede verse si mira el valor de `access-mode` para el disco compartido en la salida del comando `cluster show`.

Nota - También puede crear un dispositivo de servidor de quórum o un dispositivo de quórum de disco compartido mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Consulte las páginas del comando `man clsetup(1CL)` y `clquorum(1CL)` para obtener información sobre los comandos que se utilizan 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 desactivar 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 según el ID de dispositivo (DID), que es compartido por los nodos. Use el comando `cldevice show` para ver la lista de nombres de DID. Consulte la página del comando `man cldevice(1CL)` para 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. **Asuma un rol que proporcione la autorización 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 para la opción de quórum.

Aparece el menú Quórum.

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

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

5. Escriba el número para la opción de un dispositivo del quórum de disco compartido.

La utilidad `clsetup` pregunta qué dispositivo global quiere utilizar.

6. Escriba el dispositivo global que va a usar.

La utilidad `clsetup` solicita que confirme que el nuevo dispositivo de quórum debe agregarse al dispositivo global especificado.

7. Escriba *yes* para seguir agregando el nuevo dispositivo de quórum.

Si se agrega correctamente el nuevo dispositivo de quórum, la utilidad `clsetup` muestra el correspondiente mensaje.

8. Compruebe que se haya agregado el dispositivo de quórum.

```
# clquorum list -v
```

▼ **Cómo agregar un dispositivo de quórum NAS de Oracle ZFS Storage Appliance**

Asegúrese de 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.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para agregar un dispositivo NAS de Oracle ZFS Storage Appliance. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. Consulte la documentación de instalación que se envía con Oracle ZFS Storage Appliance o la ayuda en pantalla del dispositivo para obtener instrucciones sobre cómo configurar un dispositivo iSCSI.
2. Detecte el LUN de iSCSI en cada uno de los nodos y establezca la lista de acceso de iSCSI con configuración estática.

```
# iscsiadm modify discovery -s enable

# iscsiadm list discovery
Discovery:
Static: enabled
Send Targets: disabled
iSNS: disabled

# iscsiadm add static-config iqn.LUN-name,IP-address-of-NASdevice
# devfsadm -i iscsi
# cldevice refresh
```

3. Desde un nodo del cluster, configure los DID correspondientes al LUN de iSCSI.

```
# 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)` para obtener información adicional.

5. Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.

6. Ejecute el comando `clquorum` para agregar el dispositivo NAS como un dispositivo de 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 la página del comando `man clquorum(1CL)` para obtener más información.

▼ Adición de un dispositivo de quórum de servidor de quórum

Antes de empezar Antes de 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. Para obtener información sobre cómo instalar el servidor de quórum, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

`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 - También puede crear un dispositivo de servidor de quórum mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.**
2. **Asegúrese de que todos los nodos de Oracle Solaris Cluster estén en línea y puedan comunicarse con el servidor de quórum de Oracle Solaris Cluster.**

- a. **Compruebe que los conmutadores de red conectados directamente con los nodos del cluster cumplan uno de los criterios siguientes:**

- El conmutador es compatible con el protocolo RSTP.
- El conmutador tiene activado 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 ralentiza dicha comunicación se ralentiza 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 para la opción de quórum.**

Aparece el menú Quórum.

5. **Escriba el número para la opción de agregar un dispositivo del quórum.**

A continuación, escriba **yes** para confirmar que va a agregar un dispositivo de quórum.

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

6. **Escriba el número correspondiente a la opción de un dispositivo de quórum del servidor de quórum y, luego, escriba **yes** para confirmar que va a agregar un dispositivo de quórum del servidor de quórum.**

La utilidad `clsetup` solicita que se indique el nombre del nuevo dispositivo del quórum.

7. **Escriba el nombre del dispositivo de quórum que va a agregar.**

Elija el nombre que quiera para el dispositivo de quórum. Este nombre usa sólo para procesar futuros comandos administrativos.

La utilidad `clsetup` solicita que proporcione el nombre del host del servidor de quórum.

8. **Escriba el nombre de host del servidor de quórum.**

Este nombre especifica la dirección IP del equipo en que se ejecuta el servidor de quórum o el nombre de host del equipo dentro de la red.

Según la configuración IPv4 o IPv6 del host, la dirección IP del equipo debe especificarse en el archivo `/etc/hosts`, en el archivo `/etc/inet/ipnodes` o en ambos.

Nota - Todos los nodos del 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
```

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 \[182\]](#)
- [Eliminación del último dispositivo de quórum de un cluster \[183\]](#)
- [Sustitución de un dispositivo de quórum \[185\]](#)

▼ Eliminación de un dispositivo de quórum

Cuando se elimina un dispositivo de quórum, ya no participa en la votación 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 \[183\]](#).

`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 - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para eliminar un dispositivo de quórum. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma un rol que proporcione la autorización 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`.**

```
# clsetup
```

Aparece el menú principal.

4. **Escriba el número para la opción de quórum.**

5. **Escriba el número para la opción de eliminar un dispositivo del quórum.**

Responda a las preguntas que aparecen durante el proceso de eliminación.

6. **Salga de `clsetup`.**

7. **Compruebe que se haya eliminado el dispositivo de quórum.**

```
# clquorum list -v
```

Errores más frecuentes

Si se pierde la comunicación entre el cluster y el host del servidor de quórum durante la eliminación de un dispositivo de quórum de servidor de quórum, debe limpiar la información de configuración desactualizada 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” \[197\]](#).

▼ Eliminación del último dispositivo de quórum de un cluster

Este procedimiento elimina el último dispositivo de 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 \[182\]](#).

Agregar un dispositivo de quórum implica reconfigurar el nodo que afecta al dispositivo de quórum que ha fallado y genera una situación de error grave en el equipo. La opción Forzar permite quitar el dispositivo de quórum que ha fallado sin generar una situación de error grave en el equipo. El comando `clquorum` le permite eliminar el dispositivo de la configuración. Para obtener más información, consulte la página del comando `man clquorum(1CL)`. Después de quitar el dispositivo de quórum que ha fallado, puede agregar un nuevo dispositivo con el comando `clquorum add`. Consulte [“Adición de un dispositivo de quórum” \[175\]](#).

`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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.**

2. **Elimine 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 poner el nodo que se va a eliminar en estado de mantenimiento y, a continuación, eliminar el dispositivo de quórum con el comando `clquorum remove quorum`. Las opciones del menú de administración del cluster `clsetup` no están disponibles cuando el cluster está en modo de instalación. Consulte la página del comando `man Cómo poner un nodo en estado de mantenimiento [250]` and the `clsetup(1CL)` para obtener más información.

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

```
# clquorum list -v
```

4. **Según el motivo de eliminación del último dispositivo de quórum, continúe con uno de los siguientes pasos:**

- **Si va a reemplazar el dispositivo de quórum que se ha eliminado, realice los siguientes pasos:**

- a. **Agregue el nuevo dispositivo de quórum.**

Consulte [“Adición de un dispositivo de quórum” \[175\]](#) para obtener instrucciones sobre cómo agregar el nuevo dispositivo de quórum.

b. Quite el cluster del modo de instalación.

```
# cluster set -p installmode=disabled
```

- **Si va a reducir el cluster a un cluster de un solo nodo, quite el cluster del modo de instalación.**

```
# cluster set -p installmode=disabled
```

ejemplo 58 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.

Coloque el cluster en el modo de instalación

```
# cluster set -p installmode=enabled
```

Elimine el dispositivo de quórum

```
# clquorum remove d3
```

Compruebe que se haya quitado el dispositivo de quórum

```
# 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” \[175\]](#) 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 \[182\]](#) para quitar el antiguo dispositivo de quórum de la configuración.

3. Si el dispositivo de quórum es un disco que ha tenido un error, sustituya el disco.

Consulte los procedimientos de hardware del manual de su hardware para el contenedor de discos. Consulte también [Oracle Solaris Cluster Hardware Administration Manual](#).

Mantenimiento de dispositivos de quórum

Esta sección explica los procedimientos para mantener dispositivos de quórum.

- [Modificación de una lista de nodos de dispositivo de quórum \[186\]](#)
- [Colocación de un dispositivo de quórum en estado de mantenimiento \[187\]](#)
- [Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento \[189\]](#)
- [Enumeración de una lista con la configuración de quórum \[190\]](#)
- [Reparación de un dispositivo de quórum \[192\]](#)
- [“Cambio del timeout predeterminado del quórum” \[192\]](#)

▼ Modificación de una lista de nodos de dispositivo de quórum

Puede emplear la utilidad `clsetup` para agregar un nodo a la lista de nodos de un dispositivo de quórum existente o para eliminar un nodo de esa lista. 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 del quórum, el comando `clquorum` automáticamente configura las rutas de nodo a disco para todos los nodos conectados al disco. Para obtener más información, consulte la página del comando `man clquorum(1CL)`.

`phys-schost#` refleja un indicador de 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. Asuma un rol que proporcione la autorización 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 para la opción de quórum.**

Aparece el menú Quórum.

5. **Escriba el número para la opción de eliminar un dispositivo del quórum.**

Siga las instrucciones. Se preguntará el nombre del disco que se va a eliminar.

6. **Agregue o suprima las conexiones del nodo con el dispositivo del quórum.**

7. **Escriba el número para la opción de agregar un dispositivo del quórum.**

Siga las instrucciones. Se solicitará el nombre del disco que se va a usar como dispositivo de quórum.

8. **Compruebe que se haya agregado el dispositivo de quórum.**

```
# clquorum list -v
```

▼ Colocación de un dispositivo de quórum en estado de mantenimiento

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

Ponga el dispositivo de quórum en estado de mantenimiento si el dispositivo de quórum estará fuera de servicio durante un período 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 \[250\]](#).

`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 - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para desactivar un dispositivo de quórum para colocarlo en un estado de mantenimiento. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Si el cluster está en modo de instalación, haga clic en Restablecer dispositivos de quórum para salir del modo de instalación.

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.**
2. **Ponga el dispositivo de quórum en estado de mantenimiento.**

```
# clquorum disable device
```

device 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 59 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 activar el dispositivo de quórum, consulte [Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento \[189\]](#).

Para poner un nodo en estado de mantenimiento, consulte [Cómo poner un nodo en estado de mantenimiento \[250\]](#).

▼ Procedimiento para sacar un dispositivo de quórum del estado de mantenimiento

Siga este procedimiento para sacar un dispositivo de quórum del estado de mantenimiento y restablecer el recuento de votos de quórum al valor predeterminado.



Atención - Si no especifica la opción `globaldev` o la opción `node`, el recuento de quórum se restablece para todo el cluster.

Al configurar un dispositivo de quórum, el software de Oracle Solaris Cluster asigna al dispositivo de quórum un recuento 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 \[252\]](#).
- Para obtener más información sobre recuentos de votos de quórum, consulte “[About Quorum Vote Counts](#)” de *Oracle Solaris Cluster 4.3 Concepts Guide*.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para activar un dispositivo de quórum para sacarlo de un estado de mantenimiento. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

`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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify` en todos los nodos del cluster.**
2. **Restablezca el número de quórum.**

```
# clquorum enable device
```

device 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 60 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/rdsk/d20s2
Type:                      shared_disk
Access Mode:               scsi3
Hosts (enabled):           phys-schost-2, phys-schost-3
```

▼ Enumeración de una lista con la configuración de quórum

No necesita tener el rol de usuario root para enumerar la configuración de quórum. 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.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para ver la configuración de quórum. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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

```
% clquorum show +
```

ejemplo 61 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 sustituir un dispositivo de quórum que no funciona 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” \[175\]](#) para agregar un nuevo dispositivo de quórum.

Consulte [Eliminación de un dispositivo de quórum \[182\]](#) para quitar un dispositivo de disco como dispositivo de quórum.

2. Sustituya el dispositivo de disco.

Para reemplazar el dispositivo de disco, consulte los procedimientos para el contenedor de discos en la guía del hardware. Consulte también [Oracle Solaris Cluster Hardware Administration Manual](#).

3. Agregue el disco sustituido como nuevo dispositivo de quórum.

Consulte [“Adición de un dispositivo de quórum” \[175\]](#) para agregar un disco como nuevo dispositivo de quórum.

Nota - Si agregó un dispositivo de quórum adicional en el [Paso 1](#), ahora puede eliminarlo con seguridad. Consulte [Eliminación de un dispositivo de quórum \[182\]](#) para eliminar el dispositivo de quórum.

Cambio del timeout predeterminado del quórum

Existe un timeout predeterminado de 25 segundos para la finalización de operaciones de 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 4.3](#). 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 4.3.](#)

Nota - No cambie el timeout de quórum predeterminado de 25 segundos para Oracle Real Application Clusters (Oracle RAC). 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 de quórum que se utiliza no cumple con el timeout predeterminado de 25 segundos, utilice otro dispositivo de quórum.

Administración de servidores de quórum de Oracle Solaris Cluster

El servidor de quórum de Oracle Solaris Cluster ofrece un dispositivo de quórum que no es un dispositivo de almacenamiento compartido. En esta sección, se presentan los procedimientos para administrar servidores de quórum de Oracle Solaris Cluster, como:

- [“Inicio y detención del software del servidor del quórum” \[193\]](#)
- [Inicio de un servidor de quórum \[194\]](#)
- [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” \[197\]](#)

Para obtener información sobre la instalación y configuración de servidores de quórum de Oracle Solaris Cluster, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster 4.3.](#)

Inicio y detención del software del servidor del quórum

Estos procedimientos describen cómo iniciar y detener el software Oracle Solaris Cluster.

De manera predeterminada, estos procedimientos inician y detienen un solo servidor de quórum predefinido, salvo que haya personalizado el contenido del archivo de configuración del servidor de quórum, `/etc/scqsd/scqsd.conf`. El servidor de quórum predeterminado está vinculado al puerto 9000 y usa el directorio `/var/scqsd` para la información de quórum.

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

instalación del software de Oracle Solaris Cluster 4.3. Para obtener información sobre cómo cambiar el valor de timeout de quórum, consulte [“Cambio del timeout predeterminado del quórum” \[192\]](#).

▼ Inicio de un servidor de quórum

1. **Asuma el rol de usuario `root` en el host donde desea iniciar el software de Oracle Solaris Cluster.**
2. **Use el comando `clquorumserver start` para iniciar el software.**

```
# clquorumserver start quorumserver
```

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 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 62 Inicio de todos los servidores de quórum configurados

En el ejemplo siguiente se inician todos los servidores de quórum configurados.

```
# clquorumserver start +
```

ejemplo 63 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.

```
# clquorumserver start 2000
```

▼ Detención de un servidor de quórum

1. **Asuma el rol de usuario `root` en el host donde desea iniciar el software de Oracle Solaris Cluster.**
2. **Use el comando `clquorumserver stop` para detener el software.**

```
# 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 64 Detención de todos los servidores de quórum configurados

En el ejemplo siguiente se detienen todos los servidores de quórum configurados.

```
# clquorumserver stop +
```

ejemplo 65 Detención de un servidor de quórum específico

En el ejemplo siguiente se detiene el servidor de quórum que escucha en el puerto número 2000.

```
# 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. **Asuma el rol root en el host donde desea mostrar la información del servidor del quórum.**

Los usuarios que no tienen el rol root necesitan la autorización 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`.**

```
# 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 66 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 los 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.

```
# 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 67 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.

```
# clquorumserver show qs1 qs2 qs3
```

ejemplo 68 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:

```
# 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 \[182\]](#). 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.



Atención - 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 \[182\]](#).



Atención - Si el cluster sigue utilizando este servidor de quórum, efectuar este procedimiento afectará negativamente al quórum del cluster.

1. **Asuma el rol root en el host de servidor del 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.

<code>-I clusterID</code>	El 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.
<code>quorumserver</code>	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>.
<code>-y</code>	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 69 Limpieza de la información obsoleta de clusters de la configuración del servidor de quórum

En este ejemplo, se elimina la información sobre el cluster denominado `sc-cluster` del servidor de quórum que usa 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 las interconexiones de Oracle Solaris Cluster y las 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 cluster y las redes públicas, incluidos todos los objetos de gestión de red pública (PNM).

Los objetos PNM incluyen los grupos de rutas múltiples de red de protocolo de Internet (IPMP), las agregaciones de enlaces troncales y las agregaciones de enlaces de rutas múltiples de enlaces de datos (DLMP), y las VNIC que están directamente respaldadas por las agregaciones de enlaces. El método de rutas múltiples se instala automáticamente con el sistema operativo Oracle Solaris 11 y se debe activar 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” \[215\]](#). Para obtener información sobre IPMP, consulte el [Capítulo 3, “Administración de IPMP” de Administración de redes TCP/IP, IPMP y túneles IP en Oracle Solaris 11.3](#). Para obtener información sobre las agregaciones de enlaces, consulte el [Capítulo 2, “Configuración de alta disponibilidad mediante agregaciones de enlaces” de Gestión de enlaces de datos de red de Oracle Solaris 11.3](#).

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

- [“Administración de interconexiones de clusters” \[200\]](#)
- [“Administración de redes públicas” \[215\]](#)

Para obtener una descripción general de los procedimientos de este capítulo, consulte [Tabla 12, “Lista de tareas: administrar la interconexión de cluster”](#).

Consulte [Oracle Solaris Cluster 4.3 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 los procedimientos para reconfigurar las interconexiones de cluster, como adaptadores de transporte y cables de transporte de cluster. 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. Consulte la página del comando `man clsetup(1CL)` para obtener más información. Todos los comandos de interconexión del cluster se deben ejecutar en un nodo de cluster global.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para realizar algunas de estas tareas. Para obtener instrucciones de inicio de sesión, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

Si desea información sobre los procedimientos de instalación del software de cluster, consulte [Guía de instalación del software de Oracle Solaris Cluster 4.3](#). Para ver los procedimientos sobre tareas de servicio para componentes de hardware de clústeres, consulte el [Oracle Solaris Cluster 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 12 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 [33]
Comprobar el estado de la interconexión de los clusters con <code>clinterconnect status</code>	Comprobación del estado de la interconexión de cluster [202]
Agregar un cable de transporte de cluster, un adaptador de transporte o un conmutador de transporte mediante <code>clsetup</code>	Adición de cables de transporte de clúster, adaptadores o conmutadores de transporte [203]
Quitar un cable de transporte de cluster, un adaptador de transporte o un conmutador de transporte mediante <code>clsetup</code>	Eliminación de cables de transporte de clúster, adaptadores de transporte y conmutadores de transporte [205]
Habilitar un cable de transporte de cluster mediante <code>clsetup</code>	Cómo habilitar un cable de transporte de clúster [207]
Inhabilitar un cable de transporte de cluster mediante <code>clsetup</code>	Cómo deshabilitar un cable de transporte de clúster [208]
Determinar el número de instancia de un adaptador de transporte	Determinación del número de instancia de un adaptador de transporte [210]

Tarea	Instrucciones
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 [211]

Reconfiguración dinámica con interconexiones de clusters

Debe tener en cuenta algunas consideraciones al completar operaciones de reconfiguración dinámica en interconexiones de clusters.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la función de reconfiguración dinámica de Oracle Solaris también se aplican a la reconfiguración dinámica de Oracle Solaris Cluster (excepto la inactividad del sistema operativo). Por lo tanto, consulte la documentación sobre la función de reconfiguración dinámica de Oracle Solaris *antes* de usar la función de reconfiguración dinámica con el software de Oracle Solaris Cluster. Debe prestar especial atención a los problemas que afectan a los dispositivos de E/S no conectados en red durante las operaciones de desconexión de reconfiguración dinámica.
- El software de Oracle Solaris Cluster rechaza las operaciones de eliminación de tarjetas de reconfiguración dinámica realizadas en interfaces activas de interconexión privada.
- Debe eliminar por completo un adaptador activo del cluster para poder realizar la reconfiguración dinámica en una interconexión de cluster activa. Use el menú `clsetup` o los comandos correspondientes.



Atención - El software de Oracle Solaris Cluster requiere que cada nodo de cluster tenga al menos una ruta que funcione y se comuniquen con los demás nodos del cluster. No desactive una interfaz de interconexión privada que proporcione la última ruta a cualquier nodo del cluster.

Complete los siguientes procedimientos en el orden indicado al realizar operaciones de reconfiguración dinámica en interfaces de redes públicas.

TABLA 13 Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas

Tarea	Instrucciones
1. Desactivar y eliminar la interfaz de la interconexión activa	“Reconfiguración dinámica con interfaces de red pública” [217]
2. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública.	

▼ Comprobación del estado de la interconexión 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.

Para llevar a cabo este procedimiento, no hace falta iniciar sesión como rol root.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para comprobar el estado de la interconexión de cluster. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

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

▼ Adición de cables de transporte de clúster, 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 Hardware Administration Manual](#).

`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 - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para agregar cables, adaptadores de transporte y adaptadores privados al cluster. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

- 1. Asegúrese de que los cables de transporte de cluster físico estén instalados.**

Para conocer el procedimiento acerca de la instalación de un cable de transporte de cluster, consulte el [Oracle Solaris Cluster Hardware Administration Manual](#).

- 2. Asuma el rol `root` en cualquier nodo del cluster.**

- 3. Inicie la utilidad `clsetup`.**

```
# clsetup
```

Aparece el menú principal.

- 4. Escriba el número para la opción de mostrar el menú de interconexión de cluster.**

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

Siga las instrucciones y escriba la información que se solicite.

- 6. Escriba el número para la opción de agregar el adaptador de transporte a un nodo.**

Siga las instrucciones y escriba la información que se solicite.

Si tiene pensado usar alguno de los adaptadores siguientes para la interconexión de 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
nge	set nge:nge_taskq_disable=1
e1000g	set e1000g:e1000g_taskq_disable=1

7. **Escriba el número para la opción de agregar el conmutador de transporte.**
Siga las instrucciones y escriba la información que se solicite.
8. **Compruebe que se hayan agregado el cable de transporte de cluster, el adaptador de cluster o el conmutador de transporte.**

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

ejemplo 71 Comprobación de la agregación de un cable de transporte, un adaptador de transporte o un conmutador de transporte de cluster

En el ejemplo siguiente, se muestra cómo comprobar la agregación de un cable de transporte, un adaptador de transporte o un conmutador de transporte a un nodo. El ejemplo contiene valores de configuración para el tipo de transporte de la interfaz de proveedor de enlace de datos (DLPI).

```
# clinterconnect show phys-schost-1:net5,hub2
===Transport Cables===
Transport Cable:          phys-schost-1:net5@0,hub2
Endpoint1:                phys-schost-2:net4@0
Endpoint2:                hub2@2
State:                    Enabled

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

# clinterconnect show phys-schost-1:hub2

=== Transport Switches ===
```

```

Transport Switch:          hub2
Switch State:             Enabled
Switch Type:              switch
Switch Port Names:        1 2
Switch Port State(1):     Enabled
Switch Port State(2):     Enabled
    
```

Pasos siguientes Para comprobar el estado de la interconexión del cable de transporte de cluster, consulte [Comprobación del estado de la interconexión de cluster \[202\]](#).

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

Utilice el siguiente procedimiento para eliminar cables de transporte, adaptadores de transporte y conmutadores de transporte de cluster de una configuración de nodo. Cuando se desactiva un cable, sus dos extremos continúan configurados. Un adaptador no se puede quitar si sigue en uso como extremo de un cable de transporte.



Atención - 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 desactivar un cable, compruebe siempre el estado de interconexión de cluster de un nodo. Sólo se debe desactivar una conexión por cable tras haber comprobado que sea redundante. Es decir, asegúrese de que haya disponible otra conexión. Si se desactiva 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.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para eliminar cables, adaptadores de transporte y adaptadores privados del cluster. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma el rol root en cualquier nodo del cluster.**
2. **Compruebe el estado de la ruta de transporte restante del cluster.**

```
# clinterconnect status
```



Atención - Si recibe un mensaje de error del tipo `path faulted` 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 para la opción de acceder al menú de interconexión de cluster.

5. Escriba el número para 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 para la opción de quitar el cable de transporte.

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

Nota - Si va a quitar un cable, desconecte el cable entre el puerto y el dispositivo de destino.

7. Escriba el número para la opción de quitar el adaptador de transporte de un nodo.

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

Si desea extraer un adaptador físico de un nodo, consulte [Oracle Solaris Cluster Hardware Administration Manual](#) para conocer los procedimientos de servicio de hardware.

8. Escriba el número para la opción de quitar un conmutador de transporte.

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

Nota - No es posible eliminar un conmutador si alguno de los puertos se sigue usando como extremo de algún cable de transporte.

9. Compruebe que se haya quitado el cable, adaptador o conmutador.

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
```

```
# clinterconnect show node:switch
```

El cable o adaptador de transporte eliminado del nodo correspondiente no debe aparecer en la salida de este comando.

ejemplo 72 Comprobación de la eliminación de un cable de transporte, un adaptador de transporte o un conmutador de transporte

En el siguiente ejemplo se muestra cómo comprobar la eliminación de un cable de transporte, un adaptador de transporte o un conmutador de transporte.

```
# clinterconnect show phys-schost-1:net5,hub2@0
===Transport Cables===
Transport Cable:                phys-schost-1:net5,hub2@0
Endpoint1:                      phys-schost-1:net5
Endpoint2:                      hub2@0
State:                          Enabled

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

# clinterconnect show hub2
=== Transport Switches ===
Transport Switch:               hub2
State:                          Enabled
Type:                           switch
Port Names:                     1 2
Port State(1):                  Enabled
Port State(2):                  Enabled
```

▼ Cómo habilitar un cable de transporte de clúster

Esta opción se utiliza para activar un cable de transporte de cluster 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.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para activar un cable. Haga clic en Interconexiones privadas, a continuación, en Cables, luego, en el número de cable para resaltarlo y, por último, en Activar. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. Asuma el rol root en cualquier nodo del cluster.

2. Inicie la utilidad clsetup.

```
# clsetup
```

Aparece el menú principal.

3. Escriba el número para la opción de acceder al menú de interconexión de cluster.

4. Escriba el número para la opción de activar el cable de transporte.

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 activado.

```
# clinterconnect show node:adapter,adapternode
```

La salida es similar a la siguiente:

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

▼ Cómo deshabilitar un cable de transporte de clúster

Posiblemente deba desactivar un cable de transporte de cluster para cerrar de manera temporal 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 desactiva un cable, sus dos extremos continúan configurados. Un adaptador no se puede quitar si sigue en uso como extremo de un cable de transporte.



Atención - 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 desactivar un cable, compruebe siempre el estado de interconexión de cluster de un nodo. Sólo se debe desactivar una conexión por cable tras haber comprobado que sea redundante. Es decir, asegúrese de que haya disponible otra conexión. Si se desactiva 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.

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para desactivar un cable. Haga clic en Interconexiones privadas, a continuación, en Cables, luego, en el número de cable para resaltarlo y, por último, en Desactivar. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma el rol root en cualquier nodo del cluster.**
2. **Compruebe el estado de la interconexión de cluster antes de desactivar un cable.**

```
# clinterconnect status
```



Atención - 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 para la opción de acceder al menú de interconexión de cluster.**
5. **Escriba el número para la opción de desactivar el cable de transporte.**

Siga las peticiones de datos y suministre la información requerida. Se desactivarán 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 desactivado el cable.

```
# clinterconnect show node:adapter,adapternode
```

La salida es similar a la siguiente:

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

▼ Determinación del número de instancia de un adaptador de transporte

Debe determinar el número de instancia de un adaptador de transporte para asegurarse de agregar y eliminar el adaptador de transporte correcto mediante el comando `clsetup`. El nombre del adaptador es una combinación del tipo de adaptador y del número de instancia de dicho adaptador.

1. Según el número de ranura, busque el nombre del adaptador.

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```
# prtdiag
...
===== IO Cards =====
Bus  Max
IO  Port Bus      Freq Bus  Dev,
Type  ID  Side Slot MHz  Freq Func State Name Model
-----
XYZ   8   B    2    33   33   2,0  ok   xyz11c8,0-xyz11c8,d665.11c8.0.0
XYZ   8   B    3    33   33   3,0  ok   xyz11c8,0-xyz11c8,d665.11c8.0.0
...
```

2. Con la ruta del adaptador, busque el número de instancia del adaptador.

La pantalla siguiente es un ejemplo y tal vez no coincida con su hardware.

```
# grep sci /etc/path_to_inst
"/xyz@1f,4000/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. Para realizar esta tarea mediante la línea de comandos, consulte la página del comando `man cluster(1CL)`.

Antes de empezar Asegúrese de que el acceso de shell remoto (`rsh(1M)`) o shell seguro (`ssh(1)`) para el rol de usuario `root` esté activado para todos los nodos del cluster.

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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin` en el nodo del 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 conmuta todos los grupos de recursos del nodo especificado al siguiente nodo preferido.

```
# clnode evacuate node
# cluster shutdown -g0 -y
```

2. **Inicie la utilidad `clsetup` desde un nodo.**

Cuando se ejecuta en un modo que no sea de cluster, la utilidad `clsetup` muestra el menú principal para operaciones de un modo que no sea de cluster.

3. **Seleccione la opción de menú **Cambiar intervalos y asignación de direcciones de red para el transporte del cluster**.**

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

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

La utilidad `clsetup` muestra la dirección IP de red privada predeterminada, 172.16.0.0, y le pregunta si desea aceptarla.

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* (sí) y presione la tecla *Intro*.**
- **Realice lo siguiente para cambiar la dirección IP de red privada predeterminada:**

- a. **Escriba *no* como respuesta a la pregunta de la utilidad `clsetup` sobre si desea aceptar la dirección predeterminada; a continuación, pulse la tecla *Intro*.**

La utilidad `clsetup` solicita la nueva dirección IP de red privada.

- b. **Escriba la dirección IP nueva y pulse la tecla *Intro*.**

La utilidad `clsetup` muestra la máscara de red predeterminada; a continuación, pregunta si desea aceptar la máscara de red predeterminada.

6. Cambie 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*.**

- **Realice lo siguiente para cambiar el rango de direcciones IP:**

- a. **Escriba *no* como respuesta a la pregunta de la utilidad `clsetup` sobre si desea aceptar el rango de direcciones predeterminado; a continuación, pulse la tecla *Intro*.**

Si rechaza la máscara de red predeterminada, la utilidad `clsetup` solicita el número de nodos, redes privadas y 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 `clsetup` sobre si desea continuar con la actualización.**
- 8. **Cuando haya finalizado, salga de la utilidad `clsetup`.**
- 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.
- 10. **Compruebe que el nodo haya arrancado sin errores y esté en línea.**

`# cluster status -t node`

Resolución de problemas de interconexiones de cluster

En esta sección, se presentan los procedimientos de resolución de problemas para desactivar y, a continuación, activar una interconexión de cluster, como adaptadores y cables de transporte de cluster.

No utilice los comandos `ipadm` para administrar adaptadores de transporte de cluster. Si un adaptador de transporte se desactivó mediante el comando `ipadm disable-if`, debe utilizar los comandos `clinterconnect` para desactivar la ruta de transporte y, luego, activarla.

Este procedimiento requiere que tenga instalado el software Oracle Solaris Cluster. Estos comandos se deben ejecutar desde un nodo del cluster global.

▼ Cómo activar una interconexión de cluster

Nota - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para activar una interconexión de cluster. Haga clic en Interconexiones privadas, a continuación, en Cables, luego, en el número de cable para resaltarlo y, por último, en Activar. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. Compruebe el estado de la interconexión de cluster.

```
% clinterconnect status
```

```
=== Cluster Transport Paths===
Endpoint1      Endpoint2      Status
-----
pnode1:net1    pnode2:net1    waiting
pnode1:net5    pnode2:net5    Path online
```

2. Desactive la ruta de interconexión de cluster.

a. Compruebe la ruta de interconexión de cluster.

```
% clinterconnect show | egrep -ie "cable.*pnode1"
Transport Cable: pnode1:net5,switch2@1
Transport Cable: pnode1:net1,switch1@1
```

b. Desactive la ruta de interconexión de cluster.

```
% clinterconnect disable pnode1:net1,switch1@1
```

3. Active la ruta de interconexión de cluster.

```
% clinterconnect enable pnode1:net1,switch1@1
```

4. Verifique que la interconexión de cluster esté activada.

```
% clinterconnect status
```

```
=== Cluster Transport Paths===
Endpoint1      Endpoint2      Status
-----
pnode1:net1    pnode2:net1    Path online
pnode1:net5    pnode2:net5    Path online
```

Administración de redes públicas

El software de Oracle Solaris Cluster admite que el software de Oracle Solaris implemente IPMP, agregaciones de enlaces y VNIC para redes públicas. La administración básica de redes públicas es igual en los entornos con clusters y sin clusters.

El método de rutas múltiples se instala automáticamente al instalar el sistema operativo Oracle Solaris 11 y se debe activar para utilizarlo. La administración de varias rutas se trata en la documentación correspondiente del sistema operativo Oracle Solaris. No obstante, consulte las directrices siguientes antes de administrar IPMP, agregaciones de enlaces y VNIC en un entorno de Oracle Solaris Cluster.

Para obtener información sobre IPMP, consulte el [Capítulo 3, “Administración de IPMP” de Administración de redes TCP/IP, IPMP y túneles IP en Oracle Solaris 11.3](#). Para obtener información sobre las agregaciones de enlaces, consulte el [Capítulo 2, “Configuración de alta disponibilidad mediante agregaciones de enlaces” de Gestión de enlaces de datos de red de Oracle Solaris 11.3](#).

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 directrices siguientes.

- Al configurar un recurso de servicio escalable (SCALABLE=TRUE en el archivo de registro de tipo de recurso para el tipo de recurso) que utilice el recurso de red SUNW.SharedAddress, PNM se puede configurar para supervisar el estado de todos los grupos IPMP en los nodos del cluster, además del que SUNW.SharedAddress está configurado para usar. Esta configuración permite que el servicio se reinicie y realice un failover si falla alguno de los grupos IPMP en los nodos del cluster, a fin de maximizar la disponibilidad del servicio para los clientes de red que están ubicados en las mismas subredes que los nodos del cluster. Por ejemplo:

```
# echo ssm_monitor_all > /etc/cluster/pnm/pnm.conf
```

Reinicie el nodo.

- La variable `local-mac-address?` debe presentar un valor `true` para los adaptadores de Ethernet.
- Puede utilizar grupos IPMP basados en sondeos o 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 comprometer la disponibilidad.

Si utiliza el almacenamiento iSCSI como un dispositivo de quórum, asegúrese de que el dispositivo IPMP basado en sondeos esté configurado correctamente. Si la red iSCSI

es una red privada que contiene únicamente los nodos del cluster y el dispositivo de almacenamiento iSCSI, y no hay otros hosts presentes en la red iSCSI, el mecanismo de IPMP basado en sondeos puede fallar cuando dejan de funcionar todos los nodos, excepto uno. El problema ocurre porque no hay otros hosts en la red iSCSI para el sondeo por parte de IPMP, de modo que IPMP considera que este problema es un error de red cuando solamente un nodo permanece en el cluster. IPMP desactiva el adaptador de red iSCSI y el nodo restante pierde acceso al almacenamiento iSCSI y, por lo tanto, al dispositivo de quórum. Para resolver este problema, puede agregar un enrutador a la red iSCSI de modo que otros hosts fuera del cluster respondan a los sondeos e impidan que IPMP desactive el adaptador de red. De manera alternativa, puede configurar IPMP con failover basado en enlaces en lugar de failover basado en sondeos.

- A menos que existan una o más interfaces de red pública IPv6 que no sean de enlace local en la configuración de red pública, la utilidad `scinstall` configura automáticamente un grupo IPMP de varios adaptadores para cada conjunto de adaptadores de red pública en el cluster que utiliza la misma subred. Estos grupos se basan en enlaces con sondeos transitivos. Se pueden agregar direcciones de prueba si se requiere la detección de fallos basada en sondeos.
- 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, para un grupo de rutas múltiples denominado `sc_ipmp0`, el nombre `netiflist` 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 anule la configuración (desasocie) o apague un adaptador de un grupo de múltiples rutas de red IP sin primero conmutar las direcciones IP del adaptador que se eliminará a un adaptador alternativo del grupo, mediante el comando `if_mpadm(1M)`.
- No desasocie ni elimine una interfaz de red del grupo IPMP donde está asociada la dirección IP de alta disponibilidad de Oracle Solaris Cluster. Esta dirección IP puede pertenecer al recurso de nombre de host lógico o al recurso de dirección compartida. Sin embargo, si desasocia la interfaz activa mediante el comando `ifconfig`, Oracle Solaris Cluster ahora reconoce este evento. Realiza el failover del grupo de recursos a otros nodos en buen estado si el grupo IPMP ya no puede utilizarse en el proceso. Oracle Solaris Cluster también puede reiniciar el grupo de recursos en el mismo nodo, si el grupo IPMP es válido pero falta una dirección IP de alta disponibilidad. El grupo IPMP se vuelve inutilizable por varios motivos: pérdida de conectividad IPv4, pérdida de conectividad IPv6 o ambas. 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 múltiples rutas de red IP de un cluster, use el comando `ipmpstat -g`. Para obtener más información, consulte [Capítulo 3, “Administración de IPMP” de Administración de redes TCP/IP, IPMP y túneles IP en Oracle Solaris 11.3](#).

Si desea información sobre los procedimientos de instalación del software de cluster, consulte [Guía de instalación del software de Oracle Solaris Cluster 4.3](#). Para conocer los procedimientos sobre las tareas de mantenimiento de los componentes de hardware de las redes públicas, consulte [Oracle Solaris Cluster Hardware Administration Manual](#).

Reconfiguración dinámica con interfaces de red pública

Debe tener en cuenta diversos aspectos al llevar a cabo operaciones de reconfiguración dinámica (DR) en las interfaces de red pública de un cluster.

- Todos los requisitos, los procedimientos y las restricciones documentados sobre la función de reconfiguración dinámica de Oracle Solaris también se aplican a la reconfiguración dinámica de Oracle Solaris Cluster (excepto la inactividad del sistema operativo). Por lo tanto, consulte la documentación sobre la función de reconfiguración dinámica de Oracle Solaris *antes* de usar la función de reconfiguración dinámica con el software de Oracle Solaris Cluster. Debe prestar especial atención a los problemas que afectan a los dispositivos de E/S no conectados en red durante las operaciones de desconexión de reconfiguración dinámica.
- Las operaciones de eliminación de tarjetas de reconfiguración dinámica únicamente pueden realizarse correctamente si las interfaces de red pública no están activas. Antes de quitar una interfaz de red pública, conmute las direcciones IP del adaptador que se va a quitar a otro adaptador del grupo de múltiples rutas mediante el comando `if_mpadm`. Para obtener más información, consulte la página del comando `man if_mpadm(1M)`.
- Si intenta eliminar una tarjeta de interfaz de red pública sin haberla desactivado correctamente como una interfaz de red activa, Oracle Solaris Cluster rechaza la operación e identifica la interfaz que hubiera sido afectada por la operación.



Atención - En el caso de grupos de rutas múltiples con dos adaptadores, si el adaptador de red restante falla durante la operación de eliminación de reconfiguración dinámica en el adaptador de red desactivado, la disponibilidad se verá afectada. El adaptador restante no podrá realizar un failover mientras dure la operación de reconfiguración dinámica.

Complete los siguientes procedimientos en el orden indicado al realizar operaciones de reconfiguración dinámica en interfaces de redes públicas.

TABLA 14 Mapa de tareas: reconfiguración dinámica con interfaces de redes públicas

Tarea	Instrucciones
1. Conmutar las direcciones IP del adaptador que se va a quitar a otro adaptador del grupo de múltiples rutas mediante el comando <code>if_mpadm</code>	Página del comando <code>man if_mpadm(1M)</code> . “Cómo mover una interfaz de un grupo IPMP a otro grupo IPMP” de Administración de redes TCP/IP, IPMP y túneles IP en Oracle Solaris 11.3
2. Eliminar el adaptador del grupo de rutas múltiples mediante el comando <code>ipadm</code>	Página del comando <code>man ipadm(1M)</code> “Cómo eliminar una interfaz de un grupo IPMP” de Administración de redes TCP/IP, IPMP y túneles IP en Oracle Solaris 11.3
3. Efectuar la operación de reconfiguración dinámica en la interfaz de red pública	

Administración de nodos del cluster

Este capítulo proporciona instrucciones sobre cómo agregar o eliminar un nodo de un cluster:

- “Agregación de un nodo a un cluster o cluster de zona” [219]
- “Restauración de nodos del cluster” [222]
- “Eliminación de nodos de un cluster” [227]

Para obtener información sobre las tareas de mantenimiento del cluster, consulte [Capítulo 9, Administración del cluster](#).

Agregación de un nodo a un cluster o cluster de zona

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 en un nodo del cluster global que aloje el cluster de zona, siempre y cuando el nodo del cluster global no aloje ya un nodo de ese cluster de zona específico.

Nota - El nodo que agrega debe ejecutar la misma versión del software de Oracle Solaris Cluster que el cluster al que se une.

La especificación de una dirección IP y una 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 tenía una dirección IP o una NIC especificadas, no es necesario especificar esa información para el nuevo nodo de cluster de zona.

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

En la tabla siguiente, se muestra una lista con las tareas que se deben realizar para agregar un nodo a un cluster existente. Efectúe las tareas en el orden en que se muestran.

TABLA 15 Mapa de tareas: agregar un nodo a un cluster global o a un cluster de zona existente

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 Hardware Administration Manual
Agregar almacenamiento compartido	Siga las instrucciones de Oracle Solaris Cluster Hardware Administration Manual para agregar almacenamiento compartido de forma manual. También puede utilizar Oracle Solaris Cluster Manager para agregar un dispositivo de almacenamiento compartido a un cluster de zona. Navegue por Oracle Solaris Cluster Manager hasta la página del cluster de zona y haga clic en el separador Recursos de Solaris. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte Cómo acceder a Oracle Solaris Cluster Manager [310].
Agregar el nodo a la lista de nodos autorizados.	<code>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 4.3
Agregar el nuevo nodo a un cluster existente.	Cómo agregar un nodo a un cluster o cluster de zona existente [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 4.3 Geographic Edition System Administration Guide

▼ Cómo agregar un nodo a un cluster o cluster de zona existente

Antes de agregar un host de Oracle Solaris o una máquina virtual a un cluster global o a un cluster de zona que ya existe, asegúrese de que el nodo disponga de todos los elementos de hardware necesarios correctamente instalados y configurados, incluida una conexión física operativa con la interconexión privada del cluster.

Para obtener información sobre la instalación de hardware, consulte [Oracle Solaris Cluster Hardware Administration Manual](#) o la documentación de hardware proporcionada con el servidor.

Este procedimiento activa 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.

La petición de datos `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 cortas y largas de los nombres de comando, los comandos son idénticos.

1. **En el caso de un miembro de cluster global actual, asuma el rol `root` del miembro de cluster actual. Lleve a cabo estos pasos desde un nodo de un cluster global.**
2. **Compruebe que haya finalizado correctamente todas las tareas de configuración e instalación de hardware que se muestran como requisitos previos en el mapa de tareas de [Tabla 15, “Mapa de tareas: agregar un nodo a un cluster global o a un cluster de zona existente”](#).**

3. **Instalar y configurar el software del nuevo nodo del cluster.**

Use la utilidad `scinstall` para completar la instalación y la configuración del nodo nuevo, como se describe en [Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

4. **Use la utilidad `scinstall` en el nodo nuevo para configurar ese nodo en el cluster.**
5. **Para agregar manualmente un nodo a un cluster de zona, debe especificar el host de Oracle Solaris y el nombre del nodo virtual.**

También debe especificar un recurso de red que se utilizará para la comunicación con la red pública en cada nodo. En el ejemplo siguiente, `sczone` es el nombre de zona y `sc_ipmp0` es el nombre de grupo IPMP.

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

Para obtener instrucciones sobre cómo configurar el nodo, consulte [“Creación y configuración de un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

6. **Si el nuevo nodo del cluster de zona es de tipo de marca `solaris10` y no tiene el software de Oracle Solaris Cluster instalado en el cluster de zona, proporcione la ruta de acceso a la imagen de DVD e instale el software.**

```
# clzonecluster install-cluster -d dvd-image zone-cluster-name
```

7. **Después de configurar el nodo, reinicie el nodo en modo de cluster e instale el cluster de zona en el nodo.**

```
# clzonecluster install zone-cluster-name
```

8. **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.**

Pulse la tecla Intro.

Siga las peticiones de datos 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.

9. **Cierre la utilidad `clsetup`.**

Véase también [Página del comando `man clsetup\(1CL\)`](#).

Para obtener una lista completa de las tareas necesarias para agregar un nodo de cluster, consulte [Tabla 15, “Mapa de tareas: agregar un nodo a un cluster global o a un cluster de zona existente”](#), “Mapa de tareas: agregación de un nodo de cluster”.

Para agregar un nodo a un grupo de recursos ya existente, consulte [Guía de administración y planificación de servicios de datos de Oracle Solaris Cluster 4.3](#).

Restauración de nodos del cluster

Puede usar los archivos unificados para restaurar un nodo del cluster de modo que sea exactamente igual al archivo. Antes de restaurar el nodo, debe crear un archivo de *recuperación* en los nodos del cluster. Únicamente se puede usar un archivo de *recuperación*; no se puede usar un archivo de *clonación* para restaurar un nodo del cluster. Consulte el [Paso 1](#) a continuación para obtener instrucciones sobre cómo crear el archivo de recuperación.

Este procedimiento le solicita el nombre del cluster, los nombres de los nodos y sus direcciones MAC, y la ruta a los archivos unificados. Para cada archivo que especifique, la utilidad `scinstall` verifica que el nombre del nodo de origen del archivo sea igual al del nodo que se desea restaurar. Para obtener instrucciones sobre la restauración de los nodos de un cluster desde un archivo unificado, consulte [Cómo restaurar un nodo desde el archivo unificado \[223\]](#).

▼ Cómo restaurar un nodo desde el archivo unificado

En este procedimiento, se utiliza el formato interactivo de la utilidad `scinstall` en el servidor de Automated Installer. Debe haber configurado el servidor AI e instalado los paquetes `ha-cluster/system/install` desde los repositorios de Oracle Solaris Cluster. El nombre de nodo del archivo debe ser el mismo que el nodo que va a restaurar.

Siga estas directrices para usar la utilidad `scinstall` interactiva en este procedimiento:

- La utilidad `scinstall` interactiva completa automáticamente el texto que está escribiendo. Por lo tanto, no pulse la tecla Intro más de una vez si la siguiente pantalla de menú no aparece inmediatamente.
- A menos que se indique lo contrario, puede pulsar Control-D para volver al inicio de una serie de preguntas relacionadas o al menú principal.
- Las respuestas por defecto o aquellas proporcionadas en sesiones anteriores se muestran entre corchetes ([]) al final de una pregunta. Pulse la tecla de Intro para introducir la respuesta entre corchetes sin escribirla.

1. Asuma el rol `root` en un nodo del cluster global y cree un archivo de recuperación.

```
phys-schost# archiveadm create -r archive-location
```

Al crear un archivo, excluya los conjuntos de datos ZFS que se encuentran en el almacenamiento compartido. Si tiene previsto restaurar los datos en el almacenamiento compartido, utilice el método tradicional.

Para obtener más información sobre el uso del comando `archiveadm`, consulte la página del comando `man archiveadm(1M)`.

2. Inicie sesión en el servidor de Automated Installer y asuma el rol `root`.

3. Inicie la utilidad `scinstall`.

```
phys-schost# scinstall
```

4. Escriba el número de opción para restaurar un cluster.

```
*** Main Menu ***
```

```
Please select from one of the following (*) options:
```

- ```
* 1) Install, restore, or replicate a cluster from this Automated Installer server
* 2) Securely install, restore, or replicate a cluster from this Automated Installer server
* 3) Print release information for this Automated Installer install server

* ?) Help with menu options
```

\* q) Quit

Option: 2

Seleccione la opción 1 para restaurar un nodo del cluster mediante una instalación de servidor AI no segura. Seleccione la opción 2 para restaurar un nodo del cluster mediante una instalación de servidor AI segura.

Se muestra el menú Custom Automated Installer o el menú Custom Secure Automated Installer.

**5. Escriba el número de opción para restaurar nodos del cluster desde los archivos unificados.**

Se muestra la pantalla Nombre de cluster.

**6. Escriba el nombre del cluster que contiene los nodos que desea restaurar.**

Se muestra la pantalla Nodos del cluster.

**7. Escriba los nombres de los nodos del cluster que desea restaurar desde los archivos unificados.**

Escriba un nombre de nodo por línea. Cuando haya terminado, pulse Ctrl+D y confirme la lista; para ello, escriba yes y pulse la tecla Intro. Si desea restaurar todos los nodos del cluster, especifique todos los nodos.

Si la utilidad `scinstall` no puede encontrar la dirección MAC de los nodos, escriba cada dirección cuando se le solicite.

**8. Escriba la ruta completa al archivo de recuperación.**

El archivo utilizado para restaurar un nodo *debe* ser un archivo de recuperación. El archivo de almacenamiento que utiliza para restaurar un nodo determinado debe crearse en el mismo nodo. Repita esta operación para cada nodo de cluster que desea restaurar.

**9. Para cada nodo, confirme las opciones que ha elegido, de modo que la utilidad `scinstall` realice la configuración necesaria para instalar los nodos del cluster desde este servidor AI.**

La utilidad también imprime las instrucciones para agregar las macros DHCP en el servidor DHCP y agrega o borra las claves de seguridad para nodos SPARC (si eligió la instalación segura). Siga estas instrucciones.

**10. (Opcional) Para personalizar el dispositivo de destino, actualice el manifiesto de AI para cada nodo.**

El manifiesto de AI se encuentra en el siguiente directorio:

```
/var/cluster/logs/install/autoscinstall.d/ \
cluster-name/node-name/node-name_aimanifest.xml
```

**a. Para personalizar el dispositivo de destino, actualice el elemento target del archivo de manifiesto.**

Actualice el elemento `target` del archivo de manifiesto según la manera en la que desee usar los criterios admitidos para localizar el dispositivo de destino para la instalación. Por ejemplo, puede especificar el subelemento `disk_name`.

---

**Nota** - `scinstall` supone que el disco de inicio existente en el archivo de manifiesto es el dispositivo de destino. Para personalizar el dispositivo de destino, actualice el elemento `target` del archivo de manifiesto. Para obtener más información, consulte [Parte III, “Instalación con un servidor de instalación,” de \*Instalación de sistemas Oracle Solaris 11.3\*](#) y la página del comando `man ai_manifest(4)`.

---

**b. Ejecute el comando `installadm` para cada nodo.**

```
installadm update-manifest -n cluster-name-{sparc|i386} \
-f /var/cluster/logs/install/autoscinstall.d/cluster-name/node-name/node-
name_aimanifest.xml \
-m node-name_manifest
```

Tenga en cuenta que SPARC y i386 es la arquitectura del nodo del cluster.

**11. Si utiliza una consola de administración para el cluster, abra una pantalla de la consola para cada nodo del cluster.**

- **Si el software de `pconsole` se instala y se configura en la consola de administración, use la utilidad `pconsole` para mostrar las pantallas individuales de la consola.**

Como rol `root`, utilice el siguiente comando para iniciar la utilidad `pconsole`:

```
adminconsole# pconsole host[:port] [...] &
```

La utilidad `pconsole` abre, además, una ventana maestra desde la que puede enviar entradas a todas las ventanas individuales de la consola al mismo tiempo.

- **Si no usa la utilidad `pconsole`, conecte con la consola de cada nodo por separado.**

**12. Cierre e inicie los nodos para comenzar la instalación mediante AI.**

El software de Oracle Solaris se instala con la configuración por defecto.

---

**Nota** - No puede usar este método si desea personalizar la instalación de Oracle Solaris. Si selecciona la instalación interactiva de Oracle Solaris, Automated Installer se omite, y el software de Oracle Solaris Cluster no se instala ni se configura.

Para personalizar Oracle Solaris durante la instalación, siga las instrucciones descritas en [“Cómo instalar el software de Oracle Solaris” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#) y, luego, instale y configure el cluster siguiendo las instrucciones descritas en [“Cómo instalar paquetes de software de Oracle Solaris Cluster” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

---

■ **SPARC:**

a. **Cierre todos los nodos.**

```
phys-schost# cluster shutdown -g 0 -y
```

b. **Inicie el nodo con el siguiente comando.**

```
ok boot net:dhcp - install
```

---

**Nota** - Deje un espacio a cada lado del guión (-) en el comando.

---

■ **x86**

a. **Reinicie el nodo.**

```
reboot -p
```

b. **Durante el inicio de PXE, presione Ctrl+N.**

Aparece el menú de GRUB.

c. **Inmediatamente, seleccione la entrada de instalación automatizada.**

---

**Nota** - Si no selecciona la entrada de instalación automatizada en 20 segundos, la instalación se ejecuta con el método de instalador de texto interactivo por defecto, que no instala ni configura el software Oracle Solaris Cluster.

---

Cada nodo se reiniciará automáticamente para unirse al cluster una vez finalizada la instalación. El nodo se restaura al mismo estado que tenía cuando se creó el archivo. La salida de la instalación de Oracle Solaris Cluster se registra en el archivo `/var/cluster/logs/install/sc_ai_config.log` en cada nodo.

**13. Desde un nodo, compruebe que todos los nodos se hayan unido al cluster.**

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

La salida presenta un aspecto similar al siguiente.

```
=== Cluster Nodes ===

--- Node Status ---

Node Name Status

phys-schost-1 Online
phys-schost-2 Online
phys-schost-3 Online
```

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

## Eliminación de nodos de un cluster

En esta sección, se proporcionan instrucciones para eliminar un nodo de un cluster global o un cluster de zona. También es posible eliminar un cluster de zona específico de un cluster global. En la tabla siguiente, se muestran las tareas que se deben realizar para eliminar un nodo de un cluster existente. Efectúe las tareas en el orden en que se muestran.



**Atención** - Si se elimina un nodo para una configuración de RAC mediante este procedimiento solamente, el nodo podría experimentar un error grave al reiniciarse. 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 Guía del servicio de datos de Oracle para Oracle Real Application Clusters](#). Tras completar este proceso, elimine un nodo de una configuración de RAC y siga los pasos adecuados que se muestran a continuación.

**TABLA 16** Mapa de tareas: eliminar un nodo

| Tarea                                                                                                                                                                                                                                                                                                                                            | Instrucciones                                                                                           |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------|
| Mueva todos los grupos de recursos y grupos de dispositivos fuera del nodo que se va a eliminar. Si hay un cluster de zona, inicie sesión en el cluster de zona y evacúe el nodo de cluster de zona que se encuentra en el nodo físico que se está desinstalando. Luego, elimine el nodo del cluster de zona antes de desactivar el nodo físico. | <code>clnode evacuate node</code><br><a href="#">Eliminación de un nodo de un cluster de zona [228]</a> |
| Si el nodo físico afectado ya ha fallado, simplemente elimine el nodo del cluster.                                                                                                                                                                                                                                                               |                                                                                                         |
| Verifique que se pueda eliminar el nodo comprobando los hosts permitidos.                                                                                                                                                                                                                                                                        | <code>claccess show</code><br><br><code>claccess allow -h node-to-remove</code>                         |
| Si, con el comando <code>claccess show</code> , no se muestra el nodo, no se puede eliminar. Otorgue al nodo acceso a la configuración del cluster.                                                                                                                                                                                              |                                                                                                         |
| Elimine el nodo de todos los grupos de dispositivos.                                                                                                                                                                                                                                                                                             | <a href="#">Cómo eliminar un nodo de un grupo de dispositivos (Solaris Volume Manager) [141]</a>        |

| Tarea                                                                                                                                                                                            | Instrucciones                                                                                      |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------|
| Elimine todos los dispositivos de quórum conectados al nodo que se desea eliminar.                                                                                                               | <a href="#">Eliminación de un dispositivo de quórum [182]</a>                                      |
| <b>Este paso es opcional en caso de que se elimine un nodo de un cluster compuesto por dos nodos.</b>                                                                                            | <a href="#">Eliminación del último dispositivo de quórum de un cluster [183]</a>                   |
| Aunque hay que quitar el dispositivo de quórum antes de eliminar el dispositivo de almacenamiento en el paso siguiente, inmediatamente después se puede volver agregar el dispositivo de quórum. |                                                                                                    |
| Ponga el nodo que se desea eliminar en el modo que no es de cluster.                                                                                                                             | <a href="#">Cómo poner un nodo en estado de mantenimiento [250]</a>                                |
| Elimine un nodo de la configuración de software del cluster.                                                                                                                                     | <a href="#">Eliminación de un nodo de la configuración de software del cluster [229]</a>           |
| (Opcional) Desinstale el software Oracle Solaris Cluster de un nodo de cluster.                                                                                                                  | <a href="#">Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster [254]</a> |

## ▼ 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 decide volver a agregar el nodo al cluster de zona, siga las instrucciones que se indican en [Tabla 15, “Mapa de tareas: agregar un nodo a un cluster global o a un cluster de zona existente”](#). La mayoría de estos pasos se realizan desde el nodo del cluster global.

**Nota** - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para cerrar un nodo del cluster de zona, pero no para eliminarlo. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

1. **Asuma el rol root 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 zone-cluster-name
```

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

3. **Elimine el nodo de todos los grupos de recursos en el cluster de zona.**

```
phys-schost# clresourcegroup remove-node -n zone-hostname -Z zone-cluster-name rg-name
```

Si utiliza el procedimiento descrito en la nota del Paso 2, los grupos de recursos deberían eliminarse automáticamente, por lo tanto, es posible que pueda omitir este paso.

**4. Desinstale el nodo del cluster de zona.**

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

**5. Elimine el nodo del cluster de zona de la configuración.**

Use los comandos siguientes:

```
phys-schost# clzonecluster configure zone-cluster-name
clzc:sczone> remove node physical-host=node
clzc:sczone> exit
```

---

**Nota** - Si el nodo del cluster de zona que desea eliminar reside en un sistema al que no puede accederse o que no es capaz de unirse al cluster, elimine el nodo con el shell interactivo de clzonecluster:

```
clzc:sczone> remove -F node physical-host=node
```

Si utiliza este método para eliminar el último nodo del cluster de zona, se le solicitará que suprima el cluster de zona por completo. Si elige no hacerlo, no se eliminará el último nodo. Esta supresión tiene el mismo efecto que `clzonecluster delete -F zone-cluster-name`.

---

**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

Realice este procedimiento para eliminar un nodo del cluster global.

phys-schost# refleja una petición de datos 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 cortas y largas de los nombres de comando, 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. Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo que desea eliminar.**

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

**3. Inicie el nodo del cluster global que desea eliminar en el modo que no es de cluster.**

Para un nodo de un cluster de zona, siga las instrucciones indicadas en [Eliminación de un nodo de un cluster de zona \[228\]](#) antes de realizar este paso.

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

```
ok boot -x
```

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

```
shutdown -g -y -i0
```

```
Press any key to continue
```

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

Para obtener más información sobre el inicio basado en GRUB, consulte “[Inicio de un sistema](#)” de *Inicio y cierre de sistemas Oracle Solaris 11.3*.

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

- c. **Agregue -x al comando para especificar que el sistema arranque en un modo que no sea de 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/kernel/#ISADIR/unix -B $ZFS-BOOTFS -x
```

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

La pantalla muestra el comando editado.

- 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, este 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.

**4. Elimine el nodo del cluster.**

- a. **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`.

**b. Ejecute el siguiente comando desde el nodo que desea eliminar:**

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

---

**Nota** - Si el nodo que se eliminará no está disponible o ya no se puede iniciar, ejecute el siguiente comando 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`.

---

Si desea eliminar el último nodo del cluster, este debe establecerse en el modo que no es de cluster, sin nodos activos 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 desea desinstalar el software de Oracle Solaris Cluster del nodo eliminado, continúe con [Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster \[254\]](#).

También puede eliminar el nodo del cluster y desinstalar el software de Oracle Solaris Cluster al mismo tiempo. Cambie a un directorio que no contenga ningún archivo de Oracle Solaris Cluster y escriba `scinstall -r`.

- Si no tiene previsto desinstalar el software Oracle Solaris Cluster del nodo eliminado, puede eliminar físicamente el nodo del cluster mediante la eliminación de las conexiones de hardware.

Consulte [Oracle Solaris Cluster Hardware Administration Manual](#) para obtener instrucciones.

**ejemplo 73** Eliminación de un nodo de la configuración de software del cluster

En este ejemplo, se muestra cómo eliminar un nodo (`phys-schost-2`) de 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`).

*Eliminar el nodo del cluster:*

```
phys-schost-2# clnode remove
```

```
phys-schost-1# clnode clear -F phys-schost-2
```

*Verificar la eliminación del nodo:*

```
phys-schost-1# clnode status
```

```
-- Cluster Nodes --
 Node name Status
 -
Cluster node: phys-schost-1 Online
```

**Véase también** Para desinstalar el software de Oracle Solaris Cluster del nodo eliminado, consulte [Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster \[254\]](#).

Para obtener información sobre los procedimientos de hardware, consulte [Oracle Solaris Cluster Hardware Administration Manual](#).

Para obtener una lista completa de las tareas necesarias para eliminar un nodo de cluster, consulte la [Tabla 16, “Mapa de tareas: eliminar un nodo”](#).

Para agregar un nodo a un cluster existente, consulte [Cómo agregar un nodo a un cluster o cluster de zona existente \[220\]](#).

## ▼ Eliminación de la conectividad entre una matriz y un único nodo, en un cluster con conectividad superior a dos nodos

Siga este procedimiento para desconectar una matriz de almacenamiento de un nodo único de 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.**



---

**Atención (SPARC solamente)** - Si el cluster ejecuta el software de Oracle RAC, cierre la instancia de base de datos de Oracle RAC que está en ejecución en el nodo antes de mover los grupos 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 en orden de preferencia. El comando también conmuta todos los grupos de recursos del nodo especificado al siguiente nodo preferido.

**4. Establezca los grupos de dispositivos en estado de mantenimiento.**

Para el procedimiento de establecer un grupo de dispositivos en el estado de mantenimiento, consulte [Cómo poner un nodo en estado de mantenimiento \[250\]](#).

**5. Elimine el nodo de los grupos de dispositivos.**

Si usa un disco raw, utilice el comando `cldevicegroup(1CL)` para eliminar los grupos de dispositivos.

**6. Para cada grupo de recursos que contiene un recurso HASStoragePlus, elimine el nodo de la lista de nodos del grupo de recursos.**

```
phys-schost# clresourcegroup remove-node -n node + | resourcegroup
```

Consulte [Guía de administración y planificación de servicios de datos de Oracle Solaris Cluster 4.3](#) si desea obtener 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 desea eliminar es la última que está conectada al nodo, desconecte el cable de fibra óptica que comunica el nodo y el hub o switch que está conectado a esta matriz de almacenamiento.**

En caso contrario, omita este paso.

**8. Si desea quitar el adaptador de host del nodo que planea desconectar, apague el nodo.**

Si desea quitar el adaptador de host del nodo que planea 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 sin iniciarlo.**

11. **Si se ha instalado el software de Oracle RAC, elimine el paquete de software de Oracle RAC del nodo que va a desconectar.**

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



---

**Atención (SPARC solamente)** - Si no elimina el software de Oracle RAC del nodo que desconectó, el nodo generará un aviso grave al volver a introducirlo en el cluster, lo que podría provocar una pérdida de disponibilidad de 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 de GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse `Int ro`.

13. **En el nodo, actualice el espacio de nombre del dispositivo mediante la actualización de las entradas `/devices` y `/dev`.**

```
phys-schost# devfsadm -C
phys-schost# 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 \[252\]](#).

## ▼ 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 solo en un cluster global.

```
phys-schost# boot
```

2. **¿Se ha vuelto a unir correctamente el nodo al cluster?**

- Si no fue posible, continúe con [Paso 2b](#).

- 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 \[139\]](#).
  - b. **Tras eliminar el nodo de todos los grupos de dispositivos, vuelva a [Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster \[254\]](#) 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 desinstalar el software de Oracle Solaris Cluster de un nodo de cluster \[254\]](#) 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” [237]
- “Cómo realizar tareas administrativas del cluster de zona” [268]
- “Procedimientos de resolución de problemas para realizar pruebas” [279]

Para obtener información sobre cómo agregar o eliminar un nodo del cluster, consulte [Capítulo 8, Administración de nodos del cluster](#).

### 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 es necesario que todos los nodos del cluster de zona estén activados y en funcionamiento; Oracle Solaris Cluster reproduce cualquier cambio de configuración cuando el nodo que está fuera del cluster se vuelve a unir al cluster.

---

**Nota** - De manera predeterminada, la gestión de energía está desactivada para que no interfiera con el cluster. Si activa 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 shell interactivo de `clzonecluster` es `clzc:schost>`.

**TABLA 17** Lista de tareas: administrar el cluster

| Tarea                                    | Instrucciones                                                   |
|------------------------------------------|-----------------------------------------------------------------|
| Agregar o eliminar un nodo de un cluster | <a href="#">Capítulo 8, Administración de nodos del cluster</a> |
| Cambiar el nombre del cluster            | <a href="#">Cambio del nombre del cluster [238]</a>             |

| Tarea                                                                                                                                                                          | Instrucciones                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Enumerar los ID de nodos y sus nombres de nodo correspondientes                                                                                                                | <a href="#">Asignación de un ID de nodo a un nombre de nodo [240]</a>                                                                                                                                                                                                                                                                                                                                                                                                                         |
| Permitir o denegar que se agreguen nodos nuevos al cluster                                                                                                                     | <a href="#">Cómo trabajar con autenticación para nodos de cluster nuevos [240]</a>                                                                                                                                                                                                                                                                                                                                                                                                            |
| Cambiar el tiempo para un cluster utilizando el NTP                                                                                                                            | <a href="#">Restablecimiento de la hora del día en un cluster [242]</a>                                                                                                                                                                                                                                                                                                                                                                                                                       |
| Cerrar un nodo con el indicador ok de la PROM OpenBoot en un sistema basado en SPARC o con el mensaje Press any key to continue de un menú de GRUB en un sistema basado en x86 | <a href="#">Visualización de la PROM OpenBoot en un nodo [244]</a>                                                                                                                                                                                                                                                                                                                                                                                                                            |
| Agregar o cambiar el nombre de host privado                                                                                                                                    | <a href="#">Cómo cambiar un nombre de host privado de nodo [245]</a><br><br>También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para agregar un nombre de host lógico a un cluster global o un cluster de zona. Haga clic en Tareas y, luego, en Nombre de host lógico para iniciar el asistente. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte <a href="#">Cómo acceder a Oracle Solaris Cluster Manager [310]</a> . |
| Poner un nodo de cluster en estado de mantenimiento                                                                                                                            | <a href="#">Cómo poner un nodo en estado de mantenimiento [250]</a>                                                                                                                                                                                                                                                                                                                                                                                                                           |
| Cambiar el nombre de un nodo                                                                                                                                                   | <a href="#">Cómo cambiar el nombre de un nodo [248]</a>                                                                                                                                                                                                                                                                                                                                                                                                                                       |
| Sacar un nodo del cluster fuera del estado de mantenimiento                                                                                                                    | <a href="#">Cómo sacar un nodo del estado de mantenimiento [252]</a>                                                                                                                                                                                                                                                                                                                                                                                                                          |
| Desinstalar el software del cluster de un nodo del cluster                                                                                                                     | <a href="#">Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster [254]</a>                                                                                                                                                                                                                                                                                                                                                                                            |
| Agregar y administrar una MIB de eventos de SNMP                                                                                                                               | <a href="#">Cómo activar una MIB de eventos de SNMP [259]</a><br><br><a href="#">Cómo agregar un usuario de SNMP a un nodo [263]</a>                                                                                                                                                                                                                                                                                                                                                          |
| Configurar los límites de carga de un nodo                                                                                                                                     | <a href="#">Cómo configurar límites de carga en un nodo [266]</a>                                                                                                                                                                                                                                                                                                                                                                                                                             |
| Mover un cluster de zona, preparar un cluster de zona para aplicaciones, eliminar un cluster de zona                                                                           | <a href="#">“Cómo realizar tareas administrativas del cluster de zona” [268]</a>                                                                                                                                                                                                                                                                                                                                                                                                              |

## ▼ Cambio del nombre del cluster

Si es necesario, el nombre del cluster puede cambiarse tras la instalación inicial.



**Atención** - No realice este procedimiento si el cluster está en una asociación de Oracle Solaris Cluster Geographic Edition. En cambio, siga los procedimientos descritos en [“Renaming a Cluster That Is in a Partnership” de Oracle Solaris Cluster 4.3 Geographic Edition System Administration 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. Asuma el rol root en cualquier 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 para 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, suprima la etiqueta de Oracle Solaris Cluster existente y reinicie el cluster.**

Para suprimir la instancia de la etiqueta de servicio de Oracle Solaris Cluster, complete 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 y, 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 74** 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`.

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

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

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

Durante la instalación de Oracle Solaris Cluster, 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.

No es necesario que asuma el rol de usuario root para enumerar la información de configuración para un cluster global o de zona. 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` 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. **(Opcional) Muestre 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 75** 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
```

## ▼ Cómo trabajar con autenticación para nodos de cluster nuevos

Oracle Solaris Cluster permite determinar si se pueden agregar nodos nuevos al cluster global y el tipo de autenticación que se debe 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. Consulte las páginas del comando `man keyserv(1M)` y `publickey(4)` 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. **Asuma el rol root en cualquier 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 para la opción para nuevos nodos.**  
Aparece el menú Nuevos nodos.
4. **Seleccione desde el menú y siga las instrucciones que aparecen en la pantalla.**

**ejemplo 76** 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 77** 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 78** 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 79** 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 80** 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 NTP para mantener la sincronización de hora 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 [Oracle Solaris Cluster 4.3 Concepts Guide](#) y *Guía del usuario del protocolo de hora de red* en <http://download.oracle.com/docs/cd/E19065-01/servers.10k/>.



**Atención** - 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)`, `svcadm(1M)` o `cron(1M)`. La página del comando `man ntpd(1M)` se entrega con el paquete de `service/network/ntp` Oracle Solaris 11.

`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. Asuma el rol `root` en cualquier nodo del cluster global.**

**2. Cierre el cluster global.**

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

**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.**

Se muestra el menú de GRUB.

Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3](#).

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

Se muestra la pantalla de parámetros de inicio de GRUB.

- c. Agregue -x al comando para especificar que el sistema arranque en un modo que no sea de 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/kernel/$ISADIR/unix _B $ZFS-BOOTFS -x
```

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

La pantalla muestra el comando editado.

- e. Escriba b para iniciar el nodo en el modo sin 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 de ese nodo; para ello, ejecute el comando `rddate(1M)`.**

```
phys-schost# rdate hostname
```

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 OpenBoot™ PROM.

`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 de cluster con el comando `clnode evacuate` y, a continuación, el comando `shutdown`.**

El comando `clnode evacuate` conmuta todos los grupos de dispositivos del nodo especificado al siguiente nodo por orden de preferencia. Este comando también conmuta todos los grupos de recursos del nodo especificado del cluster global al nodo que se sitúa a continuación en el orden de preferencia.

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



**Atención** - 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

Siga este procedimiento para cambiar el 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 predeterminado usa el formato `clusternodeid-priv`, por ejemplo: `clusternode3-priv`. Cambie un nombre de host privado sólo si el nombre ya se utiliza en el dominio.



**Atención** - 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. Desactive, 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 desactiva.

- 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 información sobre el uso del comando `clresource`, consulte la página del comando `man clresource(1CL)` y la [Guía de administración y planificación de servicios de datos de Oracle Solaris Cluster 4.3](#).

**2. Si el archivo de configuración NTP hace referencia al nombre de host privado que está cambiando, desactive el daemon de NTP en cada nodo del cluster.**

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

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

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

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

---

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

---

También puede ejecutar el comando `clnode` en lugar de la utilidad `clsetup` para cambiar el nombre de host privado. En el siguiente ejemplo, el nombre de nodo del cluster es `phys-schost-1`. Después de ejecutar el comando `clnode` que aparece a continuación, vaya al [Paso 6](#).

```
phys-schost# clnode set -p privatehostname=New-private-nodename phys-schost-1
```

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

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

Responda las preguntas cuando se lo solicite. Se le solicitará el nombre del nodo cuyo nombre de host privado desea cambiar (`clusternodenodeid-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. Si cambió un nombre de host privado en su configuración NTP o archivo de inclusión, actualice el archivo NTP en cada nodo.**

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

**a. Use la herramienta de edición que prefiera.**

Si efectúa este paso durante la instalación, recuerde eliminar los nombres de los nodos que estén configurados. En general, el archivo ntp.conf.sc es el mismo en todos los nodos del 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 svc:network/ntp:default
```

**8. Active todos los recursos de servicio de datos y otras aplicaciones que se hayan desactivado en el [Paso 1](#).**

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

Para obtener información sobre el uso del comando clresource, consulte la página del comando man `clresource(1CL)` y la [Guía de administración y planificación de servicios de datos de Oracle Solaris Cluster 4.3](#).

**ejemplo 81 Cambio de nombre de host privado**

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

*Disable all applications and data services as necessary*

```
phys-schost-1# svcadm disable ntp
phys-schost-1# clnode show | grep node
...
private hostname: clusternode1-priv
private hostname: clusternode2-priv
private hostname: clusternode3-priv
...
phys-schost-1# clsetup
phys-schost-1# nscd -i hosts
```

```
phys-schost-1# pfedit /etc/inet/ntp.conf.sc
...
peer clusternode1-priv
peer clusternode4-priv
peer clusternode3-priv
phys-schost-1# ping clusternode4-priv
phys-schost-1# svcadm enable ntp
Enable all applications and data services disabled at the beginning of the procedure
```

## ▼ 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, asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **(Opcional) Si cambia el nombre de un nodo en un cluster de Oracle Solaris Cluster Geographic Edition que está en una asociación, determine si debe cambiar el grupo de protección.**

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 [Capítulo 5, “Administering Cluster Partnerships” de Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide](#).

3. **Cambie los nombres de host de Oracle Solaris.**

Complete los pasos indicados en [“Cómo cambiar la identidad de un sistema” de Gestión del rendimiento, los procesos y la información del sistema en Oracle Solaris 11.3](#), 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 new-node old-node
```

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;reboot
```

9. **Verifique que el nodo muestre el nuevo nombre.**

```
clnode status -v
```

10. **Actualice la configuración de Geographic Edition para utilizar el nuevo nombre del nodo del cluster.**

La información de configuración utilizada por los grupos de protección y el producto de replicación de datos puede especificar el nombre del nodo.

11. **Puede elegir cambiar la propiedad `hostnameList` de los recursos de nombre de host lógico.**

Consulte [Cómo cambiar los nombres de host lógicos utilizados por los recursos de nombre de host lógico de Oracle Solaris Cluster existentes \[249\]](#) 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 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 \[248\]](#). Este paso es opcional.

1. **En el cluster global, asuma un rol que proporcione la autorización 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.**

```
clresourcegroup offline apache-rg
```

- b. **Desactive los recursos de nombre de host lógico de Apache.**

```
clresource disable apache-lh-res
```

- c. **Proporcione la nueva lista de nombres de host.**

```
clresource 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. **Active los recursos de nombre de host lógico de Apache.**

```
clresource enable apache-lh-res
```

- f. **Ponga en línea los grupos de recursos de Apache.**

```
clresourcegroup 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.**

```
clresource status apache-rs
```

## ▼ Cómo poner un nodo en estado de mantenimiento

Coloque un nodo del cluster global en estado de mantenimiento si el nodo estará fuera de servicio durante un período prolongado. 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, debe cerrar el nodo con los comandos `clnode evacuate` y `shutdown`. Para obtener más información, consulte las páginas del comando `man clnode(1CL)` y `cluster(1CL)`.

---

**Nota** - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para evacuar un nodo y cambiar todos los grupos de recursos y grupos de dispositivos al siguiente nodo preferido. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

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.

---

**Nota** - El comando `shutdown` de Oracle Solaris cierra un solo nodo, mientras que el comando `cluster shutdown` cierra todo el cluster.

---

Utilice el comando `clquorum disable` desde otro nodo que aún es miembro del cluster para poner un nodo del 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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en el nodo del cluster global que está poniendo en estado de mantenimiento.**

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

El comando `clnode evacuate` conmuta todos los grupos de recursos y grupos de dispositivos 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 -i0
```

4. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en otro nodo del cluster y ponga el nodo que cerró en el [Paso 3](#) en estado de mantenimiento.**

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

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 82**    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 \[252\]](#).

▼ **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 recuento de quórum predeterminado es  $N-1$ , donde  $N$  es el número de nodos con recuento 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.



**Atención** - Si no especifica la opción `globaldev` o la opción `node`, el recuento 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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en cualquier nodo del cluster global que no sea el 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 en la configuración del cluster hay dos nodos, vaya al [Paso 4](#).
  - Si hay más de dos nodos en la configuración del cluster, vaya al [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
```

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 quitó del estado de mantenimiento deber tener el estado `online` y mostrar el recuento de votos adecuado para los votos de quórum `Present` y `Possible`.

**ejemplo 83** 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. La salida de `cluster status` muestra los `Node votes` para que `phys-schost-1` sea 1 y el estado sea `online`. El `Quorum Summary` debe mostrar también un incremento en los recuentos de votos.

```
phys-schost-2# clquorum reset
```

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

ok **boot**

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

Cuando aparezca el menú GRUB, seleccione la entrada de Oracle Solaris que corresponda y pulse Intro.

```
phys-schost-1# clquorum status
```

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

| Needed | Present | Possible |
|--------|---------|----------|
| 4      | 6       | 6        |

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

| Node Name     | Present | Possible | Status |
|---------------|---------|----------|--------|
| phys-schost-2 | 1       | 1        | Online |
| phys-schost-3 | 1       | 1        | Online |

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

| Device Name          | Present | Possible | Status |
|----------------------|---------|----------|--------|
| /dev/did/rdisk/d3s2  | 1       | 1        | Online |
| /dev/did/rdisk/d17s2 | 0       | 1        | Online |
| /dev/did/rdisk/d31s2 | 1       | 1        | Online |

## ▼ Cómo desinstalar el software de Oracle Solaris Cluster de un nodo de cluster

Realice este procedimiento para anular la configuración del 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 de Oracle Solaris Cluster desde un nodo que todavía no se unió al cluster o que aún esté en modo de instalación. En cambio, vaya a [“Cómo anular la configuración del software Oracle Solaris Cluster para solucionar problemas de instalación” de Guía de instalación del software de Oracle Solaris Cluster 4.3.](#)

---

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 [Tabla 16, “Mapa de tareas: eliminar un nodo”](#).

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

---

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

---

**2. Asuma el rol root en el nodo que desee desinstalar.**

**3. 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 -i0 ok boot -x
```

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

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

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

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

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

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

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

```
ok boot -x
```

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

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

Para obtener más información sobre el inicio basado en GRUB, consulte [“Inicio de un sistema” de Inicio y cierre de sistemas Oracle Solaris 11.3.](#)

- b. **En la pantalla de parámetros de inicio, seleccione la entrada `kernel` y escriba `e` para editarla.**

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

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

La pantalla muestra el comando editado.

- e. **Escriba `b` para iniciar el nodo en el modo sin 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 el modo sin cluster, siga estos pasos para volver a agregar la opción `-x` al comando del parámetro de inicio del núcleo.

---

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

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

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

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

`-r`

Elimina la información de configuración del cluster y desinstala la estructura de Oracle Solaris Cluster y el software de servicios de datos del nodo de cluster. Puede reinstalar el nodo o eliminarlo del cluster.

-b *bootenvironmentname*

Especifica el nombre de un nuevo entorno de inicio, que es donde iniciará una vez completado el proceso de desinstalación. No es obligatorio especificar un nombre. Si no especifica un nombre para el entorno de inicio, se genera uno automáticamente.

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

8. **Si tiene previsto volver a instalar el software de Oracle Solaris Cluster en este nodo una vez finalizada la desinstalación, reinicie el nodo para iniciarlo el nuevo entorno de inicio.**
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.**

---

**Sugerencia** - Para obtener más información acerca de la migración de un espacio de nombre de dispositivos globales a un dispositivo `lofi`, consulte [“Migración del espacio de nombre de dispositivos globales” \[129\]](#).

---

## 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 desinstalar el software de Oracle Solaris Cluster de un nodo de cluster \[254\]](#) y repita el procedimiento. Asegúrese de haber realizado correctamente el [Paso 4](#) del procedimiento antes de volver a ejecutar el comando `clnode remove`.

## Lista no eliminada de grupos de dispositivos

Los siguientes mensajes de error indican que el nodo que ha eliminado todavía está en la lista con un grupo de dispositivos.

```
Verifying that no device services still reference this node ... failed
clnode: This node is still configured to host device service "service".
clnode: This node is still configured to host device service "service2".
clnode: This node is still configured to host device service "service3".
clnode: This node is still configured to host device service "dg1".
```

```
clnode: It is not safe to uninstall with these outstanding errors.
clnode: Refer to the documentation for complete uninstall instructions.
clnode: Uninstall failed.
```

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

En esta sección se describe cómo crear, configurar y gestionar la base de información de administración (MIB) de eventos del Protocolo simple de administración de red (SNMP). En esta sección, también se describe cómo activar, desactivar y cambiar la MIB de eventos de SNMP de Oracle Solaris Cluster.

El software de Oracle Solaris Cluster actualmente admite una MIB: la MIB de eventos. El software del administrador de SNMP intercepta los eventos del cluster en tiempo real. Si se activa, el administrador de SNMP envía automáticamente notificaciones de captura a todos los sistemas definidos en el comando `clsnmphost`. Debido a que los clusters generan una gran cantidad de notificaciones, únicamente los eventos con una gravedad de `min_severity` o superior se envían como notificaciones de captura. De manera predeterminada, el valor `min_severity` se establece en `NOTICE`. El valor `log_number` especifica el número de eventos que se registrarán en la tabla de MIB antes de eliminar las entradas más antiguas. La MIB mantiene una tabla de sólo lectura con los eventos más actuales para los que se ha enviado una captura. El número de eventos está limitado por el valor de `log_number`. 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.

La interfaz SNMP de evento de cluster usa el adaptador SNMP del contenedor de agente común (cacao) como infraestructura de agente SNMP. De manera predeterminada, el número de puerto para el módulo SNMP es 11161, y el número de puerto predeterminado para las capturas SNMP es 11162. Estos números de puerto se pueden cambiar mediante el comando `cacaoadm`. Para obtener más información, consulte la página del comando `man cacaoadm(1M)`.

La creación, configuración y gestión de una MIB de eventos de SNMP de Oracle Solaris Cluster puede implicar las tareas siguientes.

**TABLA 18** Mapa de tareas: creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster

| Tarea                                                                                            | Instrucciones                                                                                        |
|--------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|
| Activar una MIB de eventos de SNMP.                                                              | <a href="#">Cómo activar una MIB de eventos de SNMP [259]</a>                                        |
| Desactivar una MIB de eventos de SNMP.                                                           | <a href="#">Cómo desactivar una MIB de eventos de SNMP [260]</a>                                     |
| Cambiar una MIB de eventos de SNMP.                                                              | <a href="#">Cómo cambiar una MIB de eventos de SNMP [260]</a>                                        |
| Agregar un host de SNMP a la lista de hosts que recibirá notificaciones de captura para las MIB. | <a href="#">Cómo activar un host de SNMP para que reciba capturas de SNMP en un nodo [262]</a>       |
| Eliminar un host de SNMP.                                                                        | <a href="#">Cómo desactivar un host de SNMP para que no reciba capturas de SNMP en un nodo [263]</a> |
| Agregar un usuario de SNMP.                                                                      | <a href="#">Cómo agregar un usuario de SNMP a un nodo [263]</a>                                      |
| Eliminar un usuario de SNMP.                                                                     | <a href="#">Cómo eliminar un usuario de SNMP de un nodo [265]</a>                                    |

## ▼ Cómo activar una MIB de eventos de SNMP

En este procedimiento, se muestra cómo activar 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Active la MIB de eventos de SNMP.**

```
phys-schost-1# clsnmpmib enable [-n node] MIB
```

|                        |                                                                                                                                                                                                                 |
|------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>[-n node]</code> | Especifica el <i>node</i> en el que se ubica la MIB de eventos que desea activar. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual. |
| <i>MIB</i>             | Especifica el nombre de la MIB que desea activar. En este caso, el nombre de la MIB debe ser <i>event</i> .                                                                                                     |

## ▼ Cómo desactivar una MIB de eventos de SNMP

En este procedimiento, se muestra cómo desactivar 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Desactive la MIB de eventos de SNMP.**

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

|                      |                                                                                                                                                                                                                    |
|----------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>-n node</code> | Especifica el <i>node</i> en el que se ubica la MIB de eventos que desea desactivar. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, de forma predeterminada se utiliza el nodo actual. |
| <i>MIB</i>           | Especifica el tipo de MIB que desea desactivar. En este caso, debe especificar el tipo <i>event</i> .                                                                                                              |

## ▼ Cómo cambiar una MIB de eventos de SNMP

En este procedimiento, se muestra cómo cambiar el protocolo, el valor de gravedad mínimo y el registro de eventos 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**

## 2. Cambie el protocolo, el valor de gravedad mínimo y el registro de eventos de la MIB de eventos de SNMP.

```
phys-schost-1# clsnmpmib set -n node
-p version=SNMPv3 \
-p min_severity=WARNING \
-p log_number=100 MIB
```

-n *node*

Especifica el *node* 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. Especifique *value* de la siguiente manera:

- version=SNMPv2
- version=snmpv2
- version=2
- version=SNMPv3
- version=snmpv3
- version=3

-p *min\_severity=value*

Especifica el valor de gravedad mínimo que se debe utilizar con las MIB. Especifique *value* de la siguiente manera:

- min\_severity=NOTICE
- min\_severity=WARNING
- min\_severity=ERROR
- min\_severity=CRITICAL
- min\_severity=FATAL

-p *log\_number=number*

Especifica el número de eventos que se registrarán en la tabla de MIB antes de eliminar las entradas más antiguas. El valor predeterminado es 100 . Los valores deben estar entre 100 y 500. Especifique *value* de la siguiente manera: *log\_number=100*.

*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.

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

## ▼ Cómo activar un host de SNMP para que reciba capturas de SNMP en un nodo

En este procedimiento, se muestra cómo agregar un host SNMP de un nodo a la lista de hosts que recibirán 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. **Asuma un rol que proporcione la autorización 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 SNMPcommunity`

Especifica el nombre de la comunidad de SNMP que se usa junto con el nombre del host. El host es un sistema en la red que se puede configurar para recibir las capturas.

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 *node* del cluster del host SNMP que tiene acceso a las MIB de SNMP en el cluster. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, el valor predeterminado es el nodo donde se ejecuta el comando.

*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. Puede tratarse de un host fuera del cluster o de un nodo del cluster que intenta obtener capturas SNMP.

## ▼ Cómo desactivar un host de SNMP para que no reciba capturas de SNMP en un nodo

En este procedimiento, se muestra cómo eliminar un host SNMP de un nodo de la lista de hosts que recibirán 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. **Asuma un rol que proporcione la autorización 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 cluster de cuya configuración se eliminó el host SNMP. Puede especificar un nombre o un ID de nodo. Si no especifica esta opción, el valor predeterminado es el nodo donde se ejecuta el comando.

`host`

Especifica el nombre, la dirección IP o la dirección IPv6 del host que se ha eliminado de la configuración. Puede tratarse de un host fuera del cluster o de un nodo del cluster que intenta obtener capturas SNMP.

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 usuarios 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Agregue el usuario de SNMP.**

```
phys-schost-1# clsnmpuser create -n node -a authentication -f password user
```

*-n node*

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 authentication*

Especifica el protocolo de autenticación que se utiliza para autorizar al usuario. El valor del protocolo de autenticación puede ser SHA o MD5.

*-f password*

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:

*user:password*

Las contraseñas no pueden tener espacios ni los siguientes caracteres:

- `;` (punto y coma)
- `:` (dos puntos)
- `\` (barra diagonal inversa)
- `\n` (línea nueva)

*user*

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 usuarios 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. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.modify`.**
2. **Elimine el usuario de SNMP.**

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

`-n node` 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.

`user` Especifica el nombre del usuario de SNMP que desea eliminar.

## Configuración de límites de carga

Puede configurar de carga para activar la distribución automática de la carga del grupo de recursos en los nodos. 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 de explorador de Oracle Solaris Cluster Manager para configurar los límites de carga. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#). En el siguiente procedimiento se explica cómo configurar límites de carga mediante la línea de comandos.

---

## ▼ Cómo configurar límites de carga en un nodo

---

**Nota** - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para crear y configurar un límite de carga en un nodo del cluster global o un nodo del cluster de zona, o para editar o suprimir un límite de carga de nodo existente. Haga clic en Nodos o Clusters de zona y, luego, en el nombre del nodo para acceder a su página. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

1. **Asuma un rol que proporcione la autorización 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.**

En el siguiente comando de ejemplo, el nombre del cluster de zona es `zc1`. La propiedad de ejemplo se llama `mem_load` y tiene un límite variable de 11 y un límite de carga estricto 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.

```
clnode create-loadlimit -p limitname=mem_load -Z zc1 \
-p softlimit=11 -p hardlimit=20 node1 node2 node3
```

### 3. Asigne valores de factor de carga a cada grupo de recursos.

En el siguiente ejemplo de comando, 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.

```
clresourcegroup set -p load_factors=mem_load@50,factor2@1 rg1 rg2
```

También puede realizar este paso durante la creación del grupo de recursos con el comando `clresourcegroup create`. Consulte la página del comando man [clresourcegroup\(1CL\)](#) para obtener más información.

### 4. Si lo desea, realice una o varias tareas de configuración opcionales adicionales.

#### ■ Redistribuya la carga existente.

```
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.

#### ■ Asigne 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.

#### ■ Establezca la propiedad `Preemption_mode`.

```
clresourcegroup set -p Preemption_mode=No_cost rg1
```

Consulte la página del comando man [clresourcegroup\(1CL\)](#) para obtener más información sobre las opciones `HAS_COST`, `NO_COST` y `NEVER`.

#### ■ Establezca el indicador `Concentrate_load`.

```
cluster set -p Concentrate_load=TRUE
```

#### ■ Especifique 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
```

5. **Compruebe el estado de todos los nodos del cluster global y del cluster de zona del cluster.**

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

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

## Cómo realizar tareas administrativas del cluster de zona

Puede realizar otras tareas administrativas en un cluster de zona, como mover la ruta de zona, preparar un cluster de zona para ejecutar aplicaciones y clonar un cluster de zona. Todos estos comandos deben ejecutarse desde el nodo del cluster global.

Puede crear un nuevo cluster de zona o agregar un sistema de archivos o dispositivo de almacenamiento a un cluster de zona existente mediante la utilidad `clsetup` para iniciar el asistente de configuración del cluster de zona. Las zonas en un cluster de zona se configuran cuando ejecuta `clzonecluster install -c` para configurar los perfiles. Consulte [“Creación y configuración de un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#) para obtener instrucciones sobre cómo usar la utilidad `clsetup` o la opción `-c config_profile`.

También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para crear un cluster de zona, o agregarle a este un sistema de archivos o un dispositivo de almacenamiento. También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para editar la propiedad Resource Security del cluster de zona. Haga clic en Clusters de zona, luego, en el nombre del cluster de zona para ir a su página y, a continuación, en el separador Recursos de Solaris para administrar los componentes del cluster de zona, o bien, en Propiedades para administrar las propiedades del cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

**Nota** - Los comandos de Oracle Solaris Cluster que ejecuta sólo desde el nodo del cluster global no son válidos para usarlos con los clusters de zona. Consulte la página del comando `man` de Oracle Solaris Cluster correspondiente para obtener información sobre el uso válido de un comando en clusters de zona.

**TABLA 19** Otras tareas del cluster de zona

| Tarea                                                  | Instrucciones                                                  |
|--------------------------------------------------------|----------------------------------------------------------------|
| Mover la ruta de zona a una nueva ruta de zona         | <code>clzonecluster move -f zonepath zone-cluster-name</code>  |
| Preparar el cluster de zona para ejecutar aplicaciones | <code>clzonecluster ready -n nodename zone-cluster-name</code> |

| Tarea                                                                             | Instrucciones                                                                                                                                                                                                                                                                                                                                                                                                                               |
|-----------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Restaurar los nodos desde los archivos unificados                                 | <a href="#">Cómo restaurar un nodo desde el archivo unificado [223]</a>                                                                                                                                                                                                                                                                                                                                                                     |
| Configurar o instalar un cluster de zona desde el archivo unificado               | <a href="#">Cómo configurar un cluster de zona desde el archivo unificado [269]</a><br><a href="#">Cómo instalar un cluster de zona desde el archivo unificado [270]</a><br><br>Use un comando:<br><br><pre>clzonecluster clone -Z target-zone-cluster-name [-m copymethod] source-zone-cluster-name</pre> Detenga el cluster de zona de origen antes de usar el subcomando clone. El cluster de zona de destino ya debe estar configurado. |
| Agregar una dirección de red a un cluster de zona                                 | <a href="#">Cómo agregar una dirección de red a un cluster de zona [272]</a>                                                                                                                                                                                                                                                                                                                                                                |
| Agregar un nodo a un cluster de zona                                              | <a href="#">Cómo agregar un nodo a un cluster o cluster de zona existente [220]</a>                                                                                                                                                                                                                                                                                                                                                         |
| Eliminar un nodo de un cluster de zona                                            | <a href="#">Eliminación de un nodo de un cluster de zona [228]</a>                                                                                                                                                                                                                                                                                                                                                                          |
| Eliminar un cluster de zona                                                       | <a href="#">Cómo eliminar un cluster de zona [273]</a>                                                                                                                                                                                                                                                                                                                                                                                      |
| Eliminar un sistema de archivos de un cluster de zona                             | <a href="#">Cómo eliminar un sistema de archivos de un cluster de zona [274]</a>                                                                                                                                                                                                                                                                                                                                                            |
| Eliminar un dispositivo de almacenamiento de un cluster de zona                   | <a href="#">Cómo eliminar un dispositivo de almacenamiento de un cluster de zona [277]</a>                                                                                                                                                                                                                                                                                                                                                  |
| Restaurar los nodos del cluster de zona desde el archivo unificado                | <a href="#">Cómo restaurar un nodo desde el archivo unificado [223]</a>                                                                                                                                                                                                                                                                                                                                                                     |
| Solucionar problema de desinstalación de nodo                                     | <a href="#">“Resolución de problemas de desinstalación de nodos” [257]</a>                                                                                                                                                                                                                                                                                                                                                                  |
| Crear, configurar y gestionar la MIB de eventos de SNMP de Oracle Solaris Cluster | <a href="#">“Creación, configuración y gestión de la MIB de eventos de SNMP de Oracle Solaris Cluster” [258]</a>                                                                                                                                                                                                                                                                                                                            |

## ▼ Cómo configurar un cluster de zona desde el archivo unificado

Utilice el comando `clzonecluster` para iniciar una utilidad interactiva para configurar un cluster de zona con marca `solaris10` o `labeled` desde el archivo unificado. La utilidad `clzonecluster configure` le permite especificar un archivo de *recuperación* o un archivo de *clonación*.

Si prefiere utilizar la línea de comandos en lugar de la utilidad interactiva para configurar un cluster de zona desde un archivo, utilice el comando `clzonecluster configure -f command-file`. Consulte la página del comando `man clzonecluster(1CL)` para obtener más información.

---

**Nota** - Si el cluster de zona que desea instalar ya ha sido configurado con otros métodos admitidos, no es necesario que configure el cluster de zona desde un archivo unificado.

---

**1. Cree un archivo de clonación o recuperación.**

```
phys-schost# archiveadm create -r archive-location
```

Utilice el comando `create` para crear un archivo de clonación o la opción `-r` para crear un archivo de recuperación. Para obtener más información sobre el uso del comando `archiveadm`, consulte la página del comando `man archiveadm(1M)`.

**2. Asuma el rol `root` en un nodo del cluster global que alojará el cluster de zona.**

**3. Configure el cluster de zona a partir del archivo recuperado o clonado en el archivo unificado.**

```
phys-schost-1# clzonecluster configure zone-cluster-name
```

El comando `clzonecluster configure zone-cluster-name` inicia la utilidad interactiva, donde puede especificar `create -a archive [other-options-such-as-"-x"]`. El archivo puede ser un archivo de clonación o un archivo de recuperación.

---

**Nota** - Los miembros del cluster de zona deben agregarse a la configuración antes de que se pueda crear un cluster de zona.

---

El subcomando `configure` utiliza el comando `zonecfg` para configurar una zona en cada máquina especificada. El subcomando `configure` permite especificar las propiedades que se aplican a cada nodo del cluster de zona. Estas propiedades tienen el mismo significado establecido por el comando `zonecfg` para las zonas individuales. El subcomando `configure` admite la configuración de propiedades que son desconocidas para el comando `zonecfg`. El subcomando `configure` inicia un shell interactivo si usted no especifica la opción `-f`. La opción `-f` toma un archivo de comandos como argumento. El subcomando `configure` utiliza este archivo para crear o modificar los clusters de zona de manera no interactiva.

## ▼ Cómo instalar un cluster de zona desde el archivo unificado

Puede instalar un cluster de zona desde el archivo unificado. La utilidad `clzonecluster install` permite especificar la ruta absoluta del archivo o el archivo de imagen de Oracle Solaris 10 que debe utilizarse para la instalación. Consulte la página del comando `man solaris10(5)` para obtener detalles sobre los tipos de archivos admitidos. Se debe poder acceder a la ruta absoluta del archivo desde todos los nodos físicos del cluster donde se instalará el cluster de zona. La instalación desde el archivo unificado puede utilizar un archivo de *recuperación* o un archivo de *clonación*.

Si prefiere utilizar la línea de comandos en lugar de la utilidad interactiva para instalar un cluster de zona desde un archivo, utilice el comando `clzonecluster create -a archive -z archived-zone`. Consulte la página del comando man [clzonecluster\(1CL\)](#) para obtener más información.

#### 1. Cree un archivo de clonación o recuperación.

```
phys-schost# archiveadm create -r archive-location
```

Utilice el comando `create` para crear un archivo de clonación o la opción `-r` para crear un archivo de recuperación. Para obtener más información sobre el uso del comando `archiveadm`, consulte la página del comando man [archiveadm\(1M\)](#).

#### 2. Asuma el rol `root` en un nodo del cluster global que alojará el cluster de zona.

#### 3. Instale el cluster de zona a partir del archivo recuperado o clonado del archivo unificado.

```
phys-schost-1# clzonecluster install -a absolute_path_to_archive zone-cluster-name
```

Se debe poder acceder a la ruta absoluta del archivo desde todos los nodos físicos del cluster donde se instalará el cluster de zona. Si tiene una ubicación HTTPS del archivo unificado, utilice `-x cert|ca-cert|key=file` para especificar el certificado SSL, el certificado de la autoridad de certificación (CA) y los archivos de claves.

Los archivos unificados no contienen los recursos del nodo del cluster de zona. Los recursos del nodo se especifican recursos cuando se configura el cluster. Cuando configura un cluster de zona desde una zona global utilizando los archivos unificados, debe definir la ruta de zona.

Si el archivo unificado contiene varias zonas, use `zone-cluster-name` para especificar el nombre de zona del origen de la instalación. Consulte la página del comando man [clzonecluster\(1CL\)](#) para obtener más información.

---

**Nota** - Si el origen que utilizó para crear el archivo unificado no contiene los paquetes de Oracle Solaris Cluster, debe ejecutar `pkg install ha-cluster-packages` (y sustituir el nombre del paquete específico, como `ha-cluster-minimal` o `ha-cluster-framework-full`). Deberá iniciar la zona, ejecutar el comando `zlogin` y, a continuación, ejecutar el comando `pkg install`. Esta acción instala los mismos paquetes en el cluster de zona de destino que en el cluster global.

---

#### 4. Inicie el nuevo cluster de zona.

```
phys-schost-1# clzonecluster boot zone-cluster-name
```

## ▼ 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.

---

**Nota** - También puede agregar una dirección de red a un cluster de zona mediante la interfaz de explorador de Oracle Solaris Cluster Manager. Haga clic en Clusters de zona, luego, en el nombre del cluster de zona para ir a su página y, por último, en el separador Recursos de Solaris para administrar los componentes del cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

1. **Asuma el rol `root` en un nodo del cluster global que aloja 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 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.**  
Aparecen 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.

---

**Nota** - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para suprimir un cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

1. **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. **Suprima todos los grupos de recursos y sus recursos del cluster de zona.**

```
phys-schost# clresourcegroup delete -F -Z zone-cluster-name +
```

---

**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 zone-cluster-name
```

**4. Desinstale el cluster de zona.**

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

**5. Anule la configuración del cluster de zona.**

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

**ejemplo 84** 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 independientes StorageTek QFS
- Sistema de archivos compartidos StorageTek QFS, cuando se utiliza para admitir Oracle RAC
- Oracle Solaris ZFS (exportado como 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 StorageTek QFS
- Sistema de archivos compartidos StorageTek QFS, solo cuando se utiliza para admitir Oracle RAC
- Sistema de archivos de cluster de UFS

Puede configurar un recurso HASToragePlus o ScalMountPoint para gestionar el montaje del sistema de archivos. Para obtener instrucciones sobre cómo agregar un sistema de archivos a un

cluster de zona, consulte [“Agregación de sistemas de archivos a un cluster de zona” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

Un recurso HAStoragePlus no supervisa un sistema de archivos ZFS si el sistema de archivos tiene la propiedad mountpoint establecida en none o legacy, o la propiedad canmount establecida en off. Para el resto de los sistemas de archivos ZFS, el recurso HAStoragePlus comprueba si el sistema de archivos está montado. Si el sistema de archivos está montado, el recurso HAStoragePlus sondea la accesibilidad del sistema de archivos mediante operaciones de lectura y escritura, según el valor de la propiedad IOOption denominada ReadOnly/ReadWrite.

Si el sistema de archivos ZFS no está montado o falla el sondeo del sistema de archivos, se produce un error en el supervisor de fallos y el recurso se establece en Faulted. RGM intentará reiniciarlo, según las propiedades retry\_count y retry\_interval del recurso. Esta acción vuelve a montar el sistema de archivos si los valores de las propiedades mountpoint y canmount antes descritas no están implicados. Si el supervisor de recursos sigue fallando y supera retry\_count dentro de retry\_interval, RGM realiza el failover del recurso a otro nodo.

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.

---

**Nota** - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para eliminar un sistema de archivos de un cluster de zona. Haga clic en Clusters de zona, luego, en el nombre del cluster de zona para ir a su página y, por último, en el separador Recursos de Solaris para administrar los componentes del cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

1. **Asuma el rol root en un nodo del cluster global que aloja 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 HAStoragePlus y SUNW.ScalMountPoint, configurados para el sistema de archivos del cluster de zona que va a eliminar.**

```
phys-schost# clresource delete -F -Z zone-cluster-name 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 desactivado 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)`.

---

**Sugerencia** - 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 zone-cluster-name
```

4. **Elimine el sistema de archivos de la configuración del cluster de zona.**

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

```
clzc:zone-cluster-name> remove fs dir=filesystemdirectory
```

```
clzc:zone-cluster-name> 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 zone-cluster-name
```

**ejemplo 85** Eliminación de un sistema de archivos local 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 del tipo `HASStoragePlus`.

```
phys-schost# clzonecluster show -v sczone
...
Resource Name: fs
dir: /local/ufs-1
special: /dev/md/dsk/d0
raw: /dev/md/dsk/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 86** 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, configurado en el cluster de zona sczone en el 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 Solaris Volume Manager y dispositivos DID, de un cluster de zona. Siga este procedimiento para eliminar un dispositivo de almacenamiento desde un cluster de zona.

---

**Nota** - También puede utilizar la interfaz de explorador de Oracle Solaris Cluster Manager para eliminar un dispositivo de almacenamiento de un cluster de zona. Haga clic en Clusters de zona, luego, en el nombre del cluster de zona para ir a su página y, por último, en el separador Recursos de Solaris para administrar los componentes del cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager \[310\]](#).

---

1. **Asuma el rol root en un nodo del cluster global que aloja 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 zone-cluster dev_zone_resources
```

**3. Determine la entrada de coincidencia de los dispositivos que se van a eliminar.**

```
phys-schost# clzonecluster show -v zone-cluster
...
Resource Name: device
match: <device_match>
...
```

**4. Elimine los dispositivos de la configuración del cluster de zona.**

```
phys-schost# clzonecluster configure zone-cluster
clzc:zone-cluster-name> remove device match=devices-match
clzc:zone-cluster-name> commit
clzc:zone-cluster-name> end
```

**5. Rearranque el cluster de zona.**

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

**6. Compruebe que se hayan eliminado los dispositivos.**

```
phys-schost# clzonecluster show -v zone-cluster
```

**ejemplo 87** Eliminación de un conjunto de discos de Solaris Volume Manager de un cluster de zona

En este ejemplo, se muestra cómo eliminar un conjunto de discos de Solaris Volume Manager 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 88** Eliminación de un dispositivo de DID de un cluster de zona

En este ejemplo se muestra cómo eliminar dispositivos 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#
phys-schost# clzonecluster show -v sczone

```

## Procedimientos de resolución de problemas para realizar pruebas

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 obtener un metaconjunto de Solaris Volume Manager desde nodos iniciados en un modo no 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 nuevo dispositivo de quórum que no sea parte del metaconjunto que se obtendrá más tarde en un modo no de cluster.

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

- b. Quite el dispositivo de quórum anterior.

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

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

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
```



---

**Atención** - 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 target-node
```

3. Desconecte todos los grupos de recursos o los grupos de recursos que contengan recursos `HASStorage` o `HASStoragePlus`, y que contengan dispositivos o sistemas de archivos afectados por el metaconjunto que desea obtener más adelante en un modo no de cluster.

```
phys-schost# clresourcegroup offline resource-group
```

4. Desactive todos los recursos de los grupos de recursos que desconectó.

```
phys-schost# clresource disable resource
```

5. Anule la gestión de grupos de recursos.

```
phys-schost# clresourcegroup unmanage resource-group
```

6. Desconecte el grupo o los grupos de dispositivos correspondientes.

```
phys-schost# cldevicegroup offline device-group
```

**7. Desactive el grupo o los grupos de dispositivos.**

```
phys-schost# cldevicegroup disable device-group
```

**8. Inicie el nodo pasivo en modo que no sea de cluster.**

```
phys-schost# shutdown -g0 -i0 -y
ok> boot -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 device-group
```

**17. Inicie el grupo o los grupos de recursos.**

```
phys-schost# clresourcegroup online -eM resource-group
```

## 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 metaconjuntos de Solaris Volume Manager y metaconjuntos de Solaris Volume Manager de múltiples propietarios.

### ▼ 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 hacer copias de seguridad de las réplicas con regularidad o utilizar Oracle Explorer (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 los comandos `prtvtoc` y `metastat`) y, a continuación, volver a crear el conjunto de discos y sus componentes. Consulte [Recreación de la configuración del software de Solaris Volume Manager \[283\]](#).

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

```
/usr/sbin/prtvtoc /dev/global/rdisk/disk-name > /etc/lvm/disk-name.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 set-name >> /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 Oracle Explorer (SUNWexplo) se ejecuta en un nodo que es propietario del conjunto, recuperará la información de `prtvtoc` y `metaset -p`.

---

### ▼ Cómo depurar 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 set-name -P
```

Cuando se ejecuta este comando, se elimina la información del conjunto de discos de las réplicas de base de datos y del repositorio de Oracle Solaris Cluster. Las opciones `-P` y `-C` permiten depurar 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 set-name
```

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

---

**2. Si desea eliminar solo la información del conjunto de discos de las réplicas de base de datos, utilice el siguiente comando.**

```
/usr/sbin/metaset -s set-name -C purge
```

Por lo general, debería utilizar la opción `-P`, en lugar de la opción `-C`. Si se utiliza la opción `-C` se pueden generar problemas al volver a crear el conjunto de discos, porque el software de 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 se ha guardado la configuración actual de Solaris Volume Manager y todos sus componentes, y que se ha depurado el conjunto de discos dañado.

---

**Nota** - Utilice mediadores solo en clusters de dos nodos.

---

**1. Cree un conjunto de discos nuevo.**

```
/usr/sbin/metaset -s set-name -a -h node1 node2
```

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 set-name -aM -h node1 node2
```

**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 set-name -a -m node1 node2
```

**3. Vuelva a agregar los mismos discos al conjunto de discos de este mismo host.**

```
/usr/sbin/metaset -s set-name -a /dev/did/rdisk/disk-name /dev/did/rdisk/disk-name
```

**4. Si ha depurado el conjunto de discos y lo crea nuevamente, la tabla de contenido del volumen (VTOC) debe permanecer en los discos, de modo que puede 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/disk-name.vtoc`. Por ejemplo:

```
/usr/sbin/fmthard -s /etc/lvm/d4.vtoc /dev/global/rdsk/d4s2
```

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

Este comando 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 set-name -n -a metadvice
```

**6. Cree cada metadispositivo a partir de una configuración guardada.**

```
/usr/sbin/metainit -s set-name -a metadvice
```

**7. Si ya existe un sistema de archivos en el metadispositivo, ejecute el comando `fsck`.**

```
/usr/sbin/fsck -n /dev/md/set-name/rdsk/metadvice
```

Si el comando `fsck` sólo muestra algunos errores, como el recuento de superbloqueos, probablemente el dispositivo se haya reconstruido de manera correcta. A continuación, puede ejecutar el comando `fsck` sin la opción `-n`. Si surgen varios errores, compruebe si el metadispositivo se ha reconstruido correctamente. En caso afirmativo, revise los errores del comando `fsck` para determinar si se puede recuperar el sistema de archivos. Si no se puede recuperar, deberá restablecer los datos a partir de una copia de seguridad.

8. **Concatene el resto de los 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” [287]
- “Configuración del control de CPU” [288]

### Introducción al control de la CPU

El software de 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 las zonas, los proyectos, las agrupaciones de recursos, los juegos de procesadores y las clases de programación, consulte [Introducción a 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.

---

**Nota** - También puede usar la interfaz de explorador de Oracle Solaris Cluster Manager para ver la configuración de un cluster de zona. Para obtener instrucciones de inicio de sesión en Oracle Solaris Cluster Manager, consulte [Cómo acceder a Oracle Solaris Cluster Manager](#) [310].

---

### Elección de un escenario

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`.

## Planificador por reparto equitativo

El primer paso del procedimiento para asignar recursos compartidos de CPU a grupos de recursos es configurar el 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 CPU

En esta sección, se describe cómo controlar el uso de CPU en un nodo de cluster global.

### ▼ Cómo controlar el uso de CPU en un nodo de cluster global

Realice este procedimiento para asignar recursos compartidos de CPU a un grupo de recursos que se ejecutará en un nodo del cluster global.

Si se asignan recursos compartidos de CPU a un grupo de recursos, el software de Oracle Solaris Cluster realiza las tareas siguientes al iniciar un recurso del grupo en un nodo del cluster global:

- Aumenta el número de recursos compartidos de CPU asignados al nodo (`zone.cpu-shares`) con el número especificado de recursos compartidos de CPU, si esto todavía no se ha hecho.
- Crea un proyecto denominado `SCSLM_resource-group` en el nodo, si esto 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.cpu-shares`).
- Inicia el recurso en el proyecto `SCSLM_resource-group`.

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
```

Con la combinación de los comandos `priocntl` y `dispadmin`, se asegura de que FSS se convierta inmediatamente en el planificador predeterminado y permanezca en ese rol después del reinicio. Para obtener más información sobre la configuración de 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 que usen el control del CPU, configure el número de recursos compartidos para los nodos del cluster global y el número mínimo de CPU disponibles en el conjunto de procesadores predeterminado.**

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

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

```
-p defaultpsetmin=integer
```

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. 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.

**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
```

`-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 `zone.cpu-shares`.

*resource-group*

Especifica el nombre del grupo de recursos.

En este procedimiento, no se configura la propiedad `RG_SLM_PSET_TYPE`. En el nodo, 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`

*resource-group*           Especifica el nombre del grupo de recursos.

---

**Nota** - No elimine ni modifique el proyecto `SCSLM_resource-group`. Puede agregar más control de recursos de forma manual al proyecto; por ejemplo, puede configurar la propiedad `project.max-lwps`. Para obtener más información, consulte la página del comando `man projmod(1M)`.

---

## Actualización de software

---

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

- “Visión general de la actualización de software de Oracle Solaris Cluster” [291]
- “Actualización de software de Oracle Solaris Cluster” [292]
- “Desinstalación de un paquete” [298]

### Visión general de la actualización de software de Oracle Solaris Cluster

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

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

- **Cambiar de versión:** cambie la versión del cluster por la última versión de Oracle Solaris Cluster (con pequeños o grandes cambios) y actualice todos los paquetes del sistema operativo Oracle Solaris. Un ejemplo de una nueva versión sería una mejora de Oracle Solaris Cluster 4.0 a 5.0. Un ejemplo de una pequeña actualización sería una mejora de Oracle Solaris Cluster 4.1 a 4.2.

Ejecute la utilidad `scinstall` o el comando `scinstall -u update` para crear un nuevo entorno de inicio (una instancia iniciable de una imagen), monte el entorno de inicio en un punto de montaje que no esté en uso, actualice los bits y active el nuevo entorno de inicio. La creación del entorno de clon no consume inicialmente espacio adicional y se produce instantáneamente. Después de realizar esta actualización debe reiniciar el cluster. La actualización también actualiza el sistema operativo Oracle Solaris a la última versión compatible. Para obtener instrucciones detalladas, consulte [Oracle Solaris Cluster 4.3 Upgrade Guide](#).

Si tiene zonas de conmutación por error del tipo de marca `solaris`, siga las instrucciones de “How to Upgrade a solaris Branded Failover Zone” de [Oracle Solaris Cluster 4.3 Upgrade Guide](#).

Si tiene una zona de marca `solaris10` en un cluster de zona, siga las instrucciones de cambio de versión descritas en [“Upgrading a solaris10 Branded Zone in a Zone Cluster” de Oracle Solaris Cluster 4.3 Upgrade Guide](#).

Para actualizar un cluster de zona sin marca, siga los procedimientos descritos en [“Actualización a un paquete específico” \[294\]](#) para actualizar el cluster global subyacente. Cuando el cluster global se actualiza, sus clusters de zona también se actualizan de forma automática.

---

**Nota** - La aplicación de una SRU principal de Oracle Solaris Cluster no proporciona el mismo resultado que la actualización del software a una versión posterior de Oracle Solaris Cluster.

---

- **Actualizar:** actualice paquetes de Oracle Solaris Cluster específicos a diferentes niveles de SRU. Puede utilizar uno de los comandos `pkg` para actualizar los paquetes de IPS (Image Packaging System) en una SRU (Support Repository Update).

Las SRU se publican generalmente con frecuencia y contienen paquetes actualizados y correcciones para defectos. El depósito contiene todos los paquetes de IPS y los paquetes actualizados. Si ejecuta el comando `pkg update`, se actualiza el sistema operativo Oracle Solaris y el software de Oracle Solaris Cluster a versiones compatibles. Después de realizar esta actualización, es posible que necesite reiniciar el cluster. Para obtener instrucciones, consulte [Cómo actualizar un paquete específico \[294\]](#).

Debe ser un usuario registrado de My Oracle Support para ver y descargar las actualizaciones de software necesarias para el producto Oracle Solaris Cluster. Si no tiene una cuenta de My Oracle Support, póngase en contacto con el representante del servicio o el responsable de ventas de Oracle, o regístrese en línea, en <http://support.oracle.com>. Para obtener información sobre las actualizaciones de firmware, consulte la documentación del hardware.

---

**Nota** - Lea la actualización de software README antes de aplicar o eliminar cualquier actualización.

---

La información para usar la utilidad de gestión de paquetes de Oracle Solaris, `pkg`, se proporciona en el [Capítulo 3, “Instalación y actualización de paquetes de software” de Agregación y actualización de software en Oracle Solaris 11.3](#).

## Actualización de software de Oracle Solaris Cluster

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

**TABLA 20** Actualización de software de Oracle Solaris Cluster

| Tarea                                                                  | Instrucciones                                                                                                                                                                                                                                                                                                                                                     |
|------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Actualizar todo el cluster a una nueva versión o pequeña actualización | <a href="#">“How to Upgrade the Software (Standard Upgrade)” de Oracle Solaris Cluster 4.3 Upgrade Guide</a>                                                                                                                                                                                                                                                      |
| Revisar consejos para realizar actualizaciones correctamente           | <a href="#">“Consejos de actualización” [293]</a>                                                                                                                                                                                                                                                                                                                 |
| Actualizar un paquete específico                                       | <a href="#">Cómo actualizar un paquete específico [294]</a>                                                                                                                                                                                                                                                                                                       |
| Actualizar un servidor de quórum o servidor de instalación AI          | <a href="#">Cómo actualizar un servidor de quórum o servidor de instalación AI [298]</a>                                                                                                                                                                                                                                                                          |
| Actualizar un cluster de zona con marca solaris                        | <a href="#">Cómo actualizar un cluster de zona con marca solaris [295]</a>                                                                                                                                                                                                                                                                                        |
| Aplicar un parche en un cluster de zona con marca solaris10            | <a href="#">Cómo aplicar un parche a un cluster de zona con marca solaris10 (clzonecluster) [296]</a><br><a href="#">Cómo aplicar un parche a una zona con marca solaris10 en un cluster de zona (patchadd) [297]</a><br><a href="#">Cómo aplicar un parche en un entorno de inicio alternativo para una zona en un cluster de zona con marca solaris10 [297]</a> |
| Eliminar paquetes de Oracle Solaris Cluster                            | <a href="#">Cómo desinstalar un paquete [298]</a><br><a href="#">Cómo desinstalar paquetes de servidor de quórum o servidor de instalación AI [299]</a>                                                                                                                                                                                                           |

## Consejos de actualización

Utilice los siguientes consejos para administrar más eficazmente las actualizaciones de Oracle Solaris Cluster:

- Lea el archivo README de SRU antes de realizar la actualización.
- Compruebe los requisitos de actualización de los dispositivos de almacenamiento.
- Aplique todas las actualizaciones antes de ejecutar el cluster en un entorno de producción.
- Compruebe los niveles de hardware del firmware e instale todas las actualizaciones de firmware que puedan ser necesarias. Consulte la documentación del hardware para obtener información sobre las actualizaciones del firmware.
- Todos los nodos que actúan como miembros del cluster deben tener las mismas actualizaciones.
- Mantenga actualizados los subsistemas del cluster. Entre otros, estas actualizaciones incluyen la administración de volúmenes, el firmware de dispositivos de almacenamiento y el transporte del cluster.
- Pruebe la conmutación por error tras haber aplicado actualizaciones importantes. Prepárese para retirar la actualización si disminuye el rendimiento del cluster o si no funciona correctamente.

## Actualización a un paquete específico

Los paquetes IPS se presentaron con el sistema operativo Oracle Solaris 11. El indicador de recurso de gestión de fallo (FMRI) describe cada paquete IPS, y se puede utilizar el comando `pkg` para realizar la actualización de SRU. Como alternativa, puede utilizar el comando `scinstall -u` para realizar una actualización de SRU.

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

### ▼ Cómo actualizar un paquete específico

Realice este procedimiento en cada nodo del cluster global que desee actualizar. Los cluster de zona sin marca en el nodo de cluster también recibirán esta actualización de manera automática.

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin`.**

2. **Actualice el paquete.**

Por ejemplo, para actualizar un paquete de un editor específico, especifique el nombre del editor en `pkg-fmri`.

```
pkg update pkg-fmri
```



---

**Atención** - Si utiliza el comando `pkg update` sin ningún `pkg-fmri` especificado, se actualizan todos los paquetes instalados que tienen actualizaciones disponibles.

---

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

3. **Si actualiza un agente de servicio de datos (`ha-cluster/data-service/*` o el agente de servicio de datos genérico de `ha-cluster/ha-service/gds`), ejecute los siguientes comandos.**

```
pkg change-facet facet.version-lock.pkg-name=false
pkg update pkg-name
```

Por ejemplo:

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

```
pkg update ha-cluster/data-service/weblogic
```

---

**Nota** - Si desea congelar un agente para evitar que se actualice, ejecute los siguientes comandos.

```
pkg change-facet facet.version-lock.pkg-name=false
pkg freeze pkg-name
```

---

Para obtener más información sobre la inmovilización de un agente específico, consulte [“Control de la instalación de componentes opcionales” de Agregación y actualización de software en Oracle Solaris 11.3](#).

#### 4. Verifique que el paquete se haya actualizado.

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

## Actualización y aplicación de parches de un cluster de zona

Para actualizar un cluster de zona con marca solaris, aplique una SRU mediante el comando `scinstall-u update`.

Para aplicar un parche a un cluster de zona con marca solaris10, use el comando `clzonecluster install-cluster -p` o `patchadd`, o aplique un parche en un entorno de inicio alternativo.

En esta sección, se incluyen los siguientes procedimientos:

- [Cómo actualizar un cluster de zona con marca solaris \[295\]](#)
- [Cómo aplicar un parche a un cluster de zona con marca solaris10 \(clzonecluster\) \[296\]](#)
- [Cómo aplicar un parche a una zona con marca solaris10 en un cluster de zona \(patchadd\) \[297\]](#)
- [Cómo aplicar un parche en un entorno de inicio alternativo para una zona en un cluster de zona con marca solaris10 \[297\]](#)

### ▼ Cómo actualizar un cluster de zona con marca solaris

Realice este procedimiento para actualizar un cluster de zona con marca solaris mediante el comando `scinstall-u update` para aplicar una SRU.

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin` en un nodo del cluster global.**

2. **Desde un nodo del cluster global, actualice el nodo completo.**

```
phys-schost# scinstall -u update [-b be-name]
```

Repita este paso en cada nodo del cluster.

3. **Reinicie el cluster.**

```
phys-schost# clzonecluster reboot
```

## ▼ **Cómo aplicar un parche a un cluster de zona con marca `solaris10` (`clzonecluster`)**

Realice este procedimiento desde un nodo del cluster global.

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin` en un nodo del cluster global.**

2. **Inicie el cluster de zona en modo de ejecución fuera de línea.**

- **Si el cluster de zona no se inicia, ejecute el siguiente comando.**

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

- **Si el cluster de zona se inicia, ejecute el siguiente comando.**

```
phys-schost# clzonecluster reboot -o zone-cluster
```

3. **Desde un nodo del cluster global, aplique un parche en todo el cluster de zona con marca `solaris10`.**

```
phys-schost# clzonecluster install-cluster -p patch-spec [options] zone-cluster
```

Para obtener más información sobre el subcomando `install-cluster`, consulte la página del comando `man clzonecluster(1CL)`.

---

**Nota** - Si el comando `scinstall -u update` falla y se devuelve un mensaje que indica que el comando `pkg` está desactualizado, siga las instrucciones del mensaje.

---

4. **Rearranque el cluster de zona.**

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

## ▼ Cómo aplicar un parche a una zona con marca solaris10 en un cluster de zona (patchadd)

Realice los siguientes pasos en cada nodo de cluster de zona configurado.

1. **Ponga el cluster de zona en un estado de ejecución fuera de línea.**

```
phys-schost# clzonecluster reboot -o zone-cluster
```

2. **Inicie sesión en el nodo del cluster de zona.**

```
phys-schost# zlogin zone-cluster
```

3. **En la línea de comandos, escriba patchadd.**

```
zchost# patchadd
```

4. **Instale los parches que corresponden a la nueva versión.**

## ▼ Cómo aplicar un parche en un entorno de inicio alternativo para una zona en un cluster de zona con marca solaris10

Use este procedimiento para aplicar parches a un entorno de inicio alternativo para la zona del nodo de cluster de zona. Si es necesario, este método brinda la opción de retirar el parche del entorno de inicio.

1. **Conviértase en administrador de la zona del nodo de cluster de zona en el que desee aplicar el parche.**

2. **Cree y active un nuevo entorno de inicio y aplíquelo parches, y, luego, use el comando shutdown para reiniciar la zona.**

Siga las instrucciones de estas tareas que se incluyen en “Cómo crear y activar varios entornos de inicio en una zona con marca solaris10”, en la sección [“Acerca del inicio de una zona con marca solaris10” de Creación y uso de zonas de Oracle Solaris 10](#).

3. **Repita los pasos anteriores para las zonas del resto de los nodos de cluster de zona en los que necesite aplicar parches.**

## Actualización de un servidor de quórum o servidor de instalación AI

Utilice el siguiente procedimiento para actualizar los paquetes para el servidor de quórum o el servidor de instalación Automated Installer (AI) de Oracle Solaris 11. Para obtener más

información sobre los servidores de quórum, consulte [“Instalación y configuración del software Oracle Solaris Cluster Quorum Server” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#). Para obtener más información sobre el uso de AI, consulte [“Cómo instalar y configurar el software de Oracle Solaris y Oracle Solaris Cluster \(repositorios IPS\)” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

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

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin`.**
2. **Actualice paquetes de servidor de quórum o de servidor de instalación AI.**

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

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



---

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

---

## Desinstalación de un paquete

Puede eliminar un solo paquete o varios paquetes.

## ▼ Cómo desinstalar un paquete

1. **Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin`.**
2. **Desinstale uno o varios paquetes existentes.**

```
pkg uninstall pkg-fmri [pkg-fmri ...]
```

El comando `pkg uninstall` fallará si hay otros paquetes instalados que dependen del `pkg-fmri` que desinstala. Para desinstalar `pkg-fmri`, debe ejecutar el comando `pkg uninstall` con todos los elementos dependientes de `pkg-fmri`. Para obtener información adicional sobre la desinstalación de paquetes, consulte [Agregación y actualización de software en Oracle Solaris 11.3](#) y la página del comando `man pkg(1)`.

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

1. Asuma un rol que proporcione la autorización RBAC `solaris.cluster.admin`.
2. Desinstale paquetes de servidor de quórum o de servidor de instalación AI.

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



---

**Atención** - Este comando desinstala todos los paquetes `ha-cluster` instalados en el sistema.

---



## Copias de seguridad y restauraciones de clusters

---

Este capítulo contiene las secciones siguientes.

- “Copia de seguridad de un cluster” [301]
- “Restauración de archivos de cluster” [304]
- “Restauración de nodos del cluster” [222]

### Copia de seguridad de un cluster

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

**TABLA 21** Mapa de tareas: hacer una copia de seguridad de los archivos del cluster

| Tarea                                                                                        | Instrucciones                                                                             |
|----------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------|
| Realizar una copia de seguridad en línea para sistemas de archivos duplicados o plexados     | <a href="#">Copias de seguridad en línea para reflejos (Solaris Volume Manager) [301]</a> |
| Hacer una copia de seguridad de la configuración del cluster                                 | <a href="#">Copias de seguridad de la configuración del cluster [303]</a>                 |
| Realizar una copia de seguridad de la configuración de partición del disco de almacenamiento | Consulte la documentación del disco de almacenamiento                                     |

### ▼ Copias de seguridad en línea para reflejos (Solaris Volume Manager)

Se puede realizar la copia de seguridad de un volumen reflejado de Solaris Volume Manager sin desmontarlo ni desconectar todo el reflejo. 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. **Asuma un rol equivalente en el nodo del cluster del que esté haciendo la copia de seguridad.**
2. **Use el comando `metaset` para determinar qué nodo tiene la propiedad en el volumen con copia de seguridad.**

```
metaset -s setname
```

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

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

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

```
lockfs -w mountpoint
```

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

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

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

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

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

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

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

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

---

**Nota** - Las lecturas se siguen efectuando desde otras subduplicaciones. Sin embargo, la subduplicación fuera de línea pierde la sincronización al realizarse la primera escritura en la duplicación. Esta incoherencia se corrige al poner la subduplicación en línea de nuevo. El comando `fsck` no debe ejecutarse.

---

6. **Desbloquee los sistemas y permita que se pueda seguir escribiendo mediante el comando `lockfs` con la opción `-u`.**

```
lockfs -u mountpoint
```

7. **Realice una comprobación del sistema de archivos.**

```
fsck /dev/md/diskset/rdisk/submirror
```

8. **Haga una copia de seguridad de la subduplicación fuera de línea en una cinta o en otro medio.**

---

**Nota** - Use el nombre del dispositivo básico (`/rdsk`) para la subduplicación, en lugar del nombre del dispositivo de bloqueo (`/dsk`).

---

9. **Use el comando `metattach` para volver a conectar el metadispositivo o volumen.**

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

Cuando se pone en línea un metadispositivo o un volumen, se vuelve a sincronizar automáticamente con la duplicación. Para obtener más información, consulte la página del comando `man metattach(1M)`.

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

```
metastat -s setname mirror
```

Consulte [Gestión de sistemas de archivos ZFS en Oracle Solaris 11.3](#) para obtener más información.

## ▼ 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 asuma un rol que le proporcione la autorización RBAC `solaris.cluster.read`.**

2. **Exporte la información de configuración del cluster a un archivo.**

```
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. Consulte la página del comando [man `clconfiguration\(5CL\)`](#) para obtener información sobre el archivo de configuración XML.

3. **Compruebe que la información de configuración del cluster se haya exportado correctamente al archivo XML.**

```
pfedit configfile
```

## Restauración de archivos de cluster

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

Puede restaurar un cluster o un nodo desde un archivo unificado, o puede restaurar archivos o sistemas de archivos específicos.

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 22** Mapa de tareas: restaurar los archivos del cluster

| Tarea                                                                     | Instrucciones                                                                                     |
|---------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------|
| Para Solaris Volume Manager, restaure el sistema de archivos raíz ZFS (/) | <a href="#">Cómo restaurar el sistema de archivos raíz ZFS (/) (Solaris Volume Manager) [305]</a> |

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

Siga este procedimiento para restaurar sistemas de archivos raíz ZFS (/) 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. Se admite UFS, excepto como un sistema de archivos raíz. UFS se puede usar en metadispositivos de metaconjuntos de Solaris Volume Manager de discos compartidos.

---

**Nota** - Como debe particionar el disco nuevo con el mismo formato que el defectuoso, identifique el esquema de partición antes de comenzar el proceso y vuelva a crear los sistemas de archivos de forma adecuada.

---

`phys-schost#` refleja un indicador de 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. **Asuma un rol que proporcione la autorización de RBAC `solaris.cluster.modify` en un nodo de cluster con acceso a los conjuntos de discos a los que también está conectado el nodo que se va a restaurar.**

Use otro nodo *que no sea* el que va a restaurar.

2. **Elimine de todos los metaconjuntos el nombre de host del nodo que se va a restaurar.**

Ejecute este comando desde un nodo del metaconjunto que no sea el que va a eliminar. Debido a que el nodo de recuperación está fuera de línea, el sistema mostrará un error de RPC: `Rpcbind failure - RPC: Timed out`. Ignore este error y continúe con el siguiente paso.

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

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

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

Para obtener más información, consulte [“Sustitución de discos en una agrupación raíz ZFS” de Gestión de sistemas de archivos ZFS en Oracle Solaris 11.3.](#)

Para recuperar las instantáneas de la agrupación raíz o la agrupación raíz ZFS, siga el procedimiento de [“Sustitución de discos en una agrupación raíz ZFS” de Gestión de sistemas de archivos ZFS en Oracle Solaris 11.3.](#)

---

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

---

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

**4. Rearranque el nodo en modo multiusuario.**

```
reboot
```

**5. Sustituya el ID del dispositivo.**

```
cldevice repair root-disk
```

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

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

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

`-f raw-disk-device` Dispositivo de disco básico en el que crear las réplicas.

`-a` Agrega réplicas.

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

**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 rearranca en modo de cluster. El cluster está preparado para utilizarse.

**ejemplo 89** Restauración del sistema de archivos raíz ZFS (/) (Solaris Volume Manager)

En el ejemplo siguiente, se muestra el sistema de archivos raíz (/) restaurado en el nodo `phys-schost-1`. El comando `metaset` se ejecuta desde otro nodo en el 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/rdisk/c0t0d0s0` y se vuelven a crear tres réplicas de la base de datos en `/dev/rdisk/c0t0d0s4`. Para obtener más información sobre la restauración de datos, consulte [“Resolución de problemas de datos en una agrupación de almacenamiento ZFS” de Gestión de sistemas de archivos ZFS en Oracle Solaris 11.3](#).

*Elimine el nodo del metaset*

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

*Reemplace el disco con fallos e inicie el nodo*

*Restaura el sistema de archivos root (/) y /usr siguiendo los procedimientos descritos en la documentación de Oracle Solaris*

*Reinicie el nodo*

```
reboot
```

*Reemplace el ID del disco*

```
cldevice repair /dev/dsk/c0t0d0
```

*Vuelva a crear las réplicas de la base de datos de estado*

```
metadb -c 3 -af /dev/rdisk/c0t0d0s4
```

*Agregue nuevamente el nodo al metaset*

```
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```



## Uso de la interfaz de explorador de Oracle Solaris Cluster Manager

---

En este capítulo, se proporciona una descripción de la interfaz de explorador de Oracle Solaris Cluster Manager, que puede utilizar para administrar varios aspectos de un cluster. Además, se incluyen procedimientos para acceder a Oracle Solaris Cluster Manager y utilizarlo.

---

**Nota** - Oracle Solaris Cluster Manager usa una versión privada del software de Oracle GlassFish Server, que se envía con el producto Oracle Solaris Cluster. No intente instalar ni actualizar a ninguno de los juegos de parches la versión pública del software de Oracle GlassFish Server. Hacerlo podría ocasionar problemas con el paquete cuando se actualiza el software Oracle Solaris Cluster o se instalan SRU de Oracle Solaris Cluster. Cualquier corrección de bug de la versión privada de Oracle GlassFish Server requerida por Oracle Solaris Cluster Manager se entrega en las SRU de Oracle Solaris Cluster.

---

Este capítulo incluye lo siguiente:

- [“Visión general de Oracle Solaris Cluster Manager” \[309\]](#)
- [“Acceso al software de Oracle Solaris Cluster Manager” \[310\]](#)
- [“Configuración de soporte de accesibilidad para Oracle Solaris Cluster Manager” \[313\]](#)
- [“Uso de la topología para supervisar el cluster” \[314\]](#)

### Visión general de Oracle Solaris Cluster Manager

La interfaz de explorador de Oracle Solaris Cluster Manager le permite visualizar gráficamente información del cluster, comprobar el estado de los componentes del cluster y supervisar los cambios de configuración. Oracle Solaris Cluster Manager también le permite realizar diversas tareas administrativas, incluida la administración de los siguientes componentes de Oracle Solaris Cluster:

- Servicios de datos
- Clusters de zona
- Nodos
- Conmutadores, cables y adaptadores privados

- Dispositivos de quórum
- Grupos de dispositivos
- Discos
- Dispositivos NAS
- Límites de carga de nodos
- Recursos y grupos de recursos
- Asociaciones de Geographic Edition

Oracle Solaris Cluster Manager actualmente no puede realizar todas las tareas administrativas de Oracle Solaris Cluster. En algunas operaciones debe utilizarse la interfaz de línea de comandos.

## Acceso al software de Oracle Solaris Cluster Manager

La interfaz de explorador de Oracle Solaris Cluster Manager ofrece una manera sencilla de administrar varias tareas en el software de Oracle Solaris Cluster. Consulte la ayuda en pantalla de Oracle Solaris Cluster Manager para obtener más información.

Common Agent Container se inicia automáticamente al iniciar el cluster. Si necesita verificar que Common Agent Container se esté ejecutando, consulte [“Resolución de problemas de Oracle Solaris Cluster Manager” \[311\]](#).

---

**Sugerencia** - No haga clic en Anterior en el explorador para salir de Oracle Solaris Cluster Manager.

---

### ▼ Cómo acceder a Oracle Solaris Cluster Manager

En este procedimiento, se muestra cómo acceder a Oracle Solaris Cluster Manager en el cluster.

1. **En un miembro del cluster, asuma el rol `root` .**
2. **Inicie un navegador desde la consola de administración u otro equipo fuera del cluster.**
3. **Compruebe que los tamaños del disco del navegador y de la antememoria estén configurados con un valor superior a 0.**
4. **Verifique que Java y Javascript estén activados en el explorador.**
5. **Conéctese al puerto de Oracle Solaris Cluster Manager en un nodo del cluster desde el explorador.**

El número de puerto predeterminado es 8998.

`https://node:8998/scm`

**6. Acepte los certificados que presente el navegador web.**

Aparece la página de inicio de sesión de Oracle Solaris Cluster Manager.

**7. Introduzca el nombre de un nodo del cluster que desea gestionar o acepte el valor predeterminado de localhost para gestionar el cluster actual.**

**8. Introduzca el nombre de usuario y la contraseña del nodo.**

**9. Haga clic en Iniciar sesión.**

Aparece la página de inicio de la aplicación Oracle Solaris Cluster Manager.

---

**Nota** - Si tiene más de un cluster configurado, puede seleccionar Otro en la lista desplegable e iniciar sesión en otro cluster para visualizar información de ese cluster. Si un cluster es miembro de una o varias asociaciones, una vez que accede a la carpeta Asociaciones, se agregan automáticamente los nombres de todos los socios a la lista desplegable. Una vez que realiza la autenticación, puede seleccionar Cambiar cluster.

---

Si no se puede conectar a Oracle Solaris Cluster Manager, consulte [“Resolución de problemas de Oracle Solaris Cluster Manager” \[311\]](#). Si elige un perfil de red restringido durante la instalación de Oracle Solaris, el acceso externo a Oracle Solaris Cluster Manager está restringido. Se requiere esta red para usar la interfaz de explorador de Oracle Solaris Cluster Manager.

## Resolución de problemas de Oracle Solaris Cluster Manager

- Verifique que se estén ejecutando los dos servicios del gestor.

```
svcs system/cluster/manager*
```

| STATE  | STIME  | FMRI                                           |
|--------|--------|------------------------------------------------|
| online | Oct_30 | svc:/system/cluster/manager-glassfish3:default |
| online | Oct_30 | svc:/system/cluster/manager:default            |

Ejecute el comando `svcadm` para desactivar o activar `system/cluster/manager-glassfish3`. Esta acción detiene y reinicia el servidor de aplicaciones respectivamente. Debe mantener `system/cluster/manager` en línea. No es necesario que lo desactive o lo active.

- Si no puede conectarse a Oracle Solaris Cluster Manager, introduzca `usr/sbin/cacaoadm status` para determinar si Common Agent Container se está ejecutando. Si Common Agent Container no se está ejecutando, aparecerá la página de inicio de sesión, pero no podrá autenticarse.

Puede introducir `usr/sbin/cacaoadm start` para iniciar manualmente Common Agent Container.

## ▼ Cómo configurar las claves de seguridad de Common Agent Container

Oracle Solaris Cluster Manager usa técnicas de cifrado eficaces para garantizar una comunicación segura entre el servidor web de Oracle Solaris Cluster Manager y todos los nodos del cluster.

Pueden surgir errores de conexión de cacao cuando se utilizan los asistentes de configuración del servicio de datos en Oracle Solaris Cluster Manager o se realizan otras tareas de Oracle Solaris Cluster Manager. Este procedimiento copia los archivos de seguridad para Common Agent Container en todos los nodos del cluster. Esto garantiza que los archivos de seguridad para Common Agent Container sean idénticos en todos los nodos del cluster y que los archivos copiados conserven los permisos de archivos correctos. Al realizar este procedimiento, se sincronizan las claves de seguridad.

### 1. En cada nodo, detenga el agente de archivo de seguridad.

```
phys-schost# /usr/sbin/cacaoadm stop
```

### 2. En un nodo, cambie al directorio `/etc/cacao/instances/default/`.

```
phys-schost-1# cd /etc/cacao/instances/default/
```

### 3. Cree un archivo tar del directorio `/etc/cacao/instances/default/`.

```
phys-schost-1# tar cf /tmp/SECURITY.tar security
```

### 4. Copie el archivo `/tmp/Security.tar` en todos los nodos del cluster.

### 5. Extraiga los archivos de seguridad en todos los nodos en los que copió el archivo `/tmp/SECURITY.tar`.

Se sobrescriben todos los archivos de seguridad que ya existen en el directorio `/etc/cacao/instances/default/`.

```
phys-schost-2# cd /etc/cacao/instances/default/
```

```
phys-schost-2# tar xf /tmp/SECURITY.tar
```

### 6. Suprima todas las copias del archivo tar para evitar riesgos de seguridad.

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
```

## 7. En cada nodo, inicie el agente de archivo de seguridad.

```
phys-schost# /usr/sbin/cacaoadm start
```

## ▼ Cómo comprobar la dirección de enlace de red

Si recibe un mensaje de error del sistema cuando intente ver información sobre un nodo distinto al nodo que está ejecutando Oracle Solaris Cluster Manager, compruebe si el parámetro "dirección de enlace a la red" del contenedor de agente común está definido en el valor correcto de 0.0.0.0.

Siga estos pasos en cada nodo del cluster.

### 1. Determine la dirección de enlace de red.

```
phys-schost# cacaoadm list-params | grep network
network-bind-address=0.0.0.0
```

Si la dirección de enlace de red está configurada con un valor distinto de 0.0.0.0, deberá cambiarlo a la dirección deseada.

### 2. Detenga e inicie cacao antes y después del cambio.

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

## Configuración de soporte de accesibilidad para Oracle Solaris Cluster Manager

Oracle Solaris Cluster Manager ofrece las siguientes opciones de configuración de accesibilidad:

- Lector de pantalla
- Contraste alto
- Fuentes grandes

Para configurar uno o varios de estos controles, haga clic en Accesibilidad en la barra de menú superior de Oracle Solaris Cluster Manager. Luego, haga clic en la casilla de control del control de accesibilidad que desee activar o desactivar.

Los cambios en la configuración del control de accesibilidad no se conservan después de la sesión. Los valores de configuración se conservan si los autentica en otro nodo de cluster durante la misma sesión.

## Uso de la topología para supervisar el cluster

La vista de topología le permite supervisar el cluster e identificar problemas. Puede ver rápidamente las relaciones entre los objetos y los grupos de recursos y recursos que pertenecen a cada nodo.

### ▼ Cómo usar la topología para supervisar y actualizar el cluster

Para acceder a la página Topología, inicie sesión en Oracle Solaris Cluster Manager, haga clic en Grupos de recursos y, luego, haga clic en el separador Topología. Las líneas representan relaciones de dependencia y colocación.

La ayuda en pantalla proporciona instrucciones detalladas sobre los elementos de la vista, además de cómo seleccionar un objeto para filtrar la vista y cómo hacer clic con el botón derecho para ver un menú contextual de acciones para ese objeto. Puede hacer clic en la flecha junto a la ayuda en pantalla para contraerla o restaurarla. También puede contraer o restaurar el filtro.

En la siguiente tabla, se proporciona una lista de los controles de la página de topología de recursos.

| Control                                                                             | Función        | Descripción                                                                                       |
|-------------------------------------------------------------------------------------|----------------|---------------------------------------------------------------------------------------------------|
|  | Zoom           | Amplíe o reduzca una parte de una página.                                                         |
|  | Visión general | Arrastre la ventanilla por el diagrama para recorrer la vista.                                    |
|  | Aislar         | Haga clic en un recurso o grupo de recursos para eliminar todos los demás objetos de la pantalla. |
|  | Detalles       | Haga clic en un grupo de recursos para obtener detalles sobre los recursos.                       |

| Control                                                                           | Función     | Descripción                                                                |
|-----------------------------------------------------------------------------------|-------------|----------------------------------------------------------------------------|
|  | Restablecer | Vuelva a la vista completa después del aislamiento o de obtener detalles.  |
|  | Filtrar     | Reduzca los resultados seleccionando objetos por tipo, instancia o estado. |

En el siguiente procedimiento, se muestra cómo supervisar los nodos del cluster para detectar errores críticos:

1. **En la ficha Topología, busque el área Posibles maestros.**
2. **Use la función de acercar para ver el estado de cada nodo del cluster.**
3. **Busque un nodo que tenga un ícono rojo de estado crítico, haga clic con el botón derecho en el nodo y seleccione Mostrar detalles.**
4. **En la página de estado del nodo, haga clic en *Archivo de registro del sistema* del sistema para ver y filtrar los mensajes de log.**



## Ejemplo de despliegue: configuración de replicación de datos basada en host entre clusters con el software Availability Suite

---

En este apéndice, se ofrece un método de uso de replicación de datos basada en host sin usar Oracle Solaris 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” \[102\]](#).

En el ejemplo de este apéndice, se muestra cómo configurar la replicación de datos basada en host entre clusters que usan el software Función Availability Suite de Oracle Solaris. 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 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 roles (RBAC) para acceder a los nodos del cluster, asegúrese de poder asumir un rol de RBAC que proporcione autorización para todos los comandos de Oracle Solaris Cluster. Esta serie de replicación de datos necesita las siguientes autorizaciones RBAC de Oracle Solaris Cluster:

- `solaris.cluster.modify`
- `solaris.cluster.admin`
- `solaris.cluster.read`

Consulte [Protección de los usuarios y los procesos en Oracle Solaris 11.3](#) para obtener más información sobre el uso de roles RBAC. Consulte las páginas del comando `man` de Oracle Solaris Cluster para saber la autorización de RBAC que requiere cada subcomando de Oracle Solaris Cluster.

## Comprensión del software Availability Suite en un cluster

En esta sección, se presenta la tolerancia ante desastres y se describen los métodos de replicación de datos que utiliza el software Availability Suite.

La tolerancia ante 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.

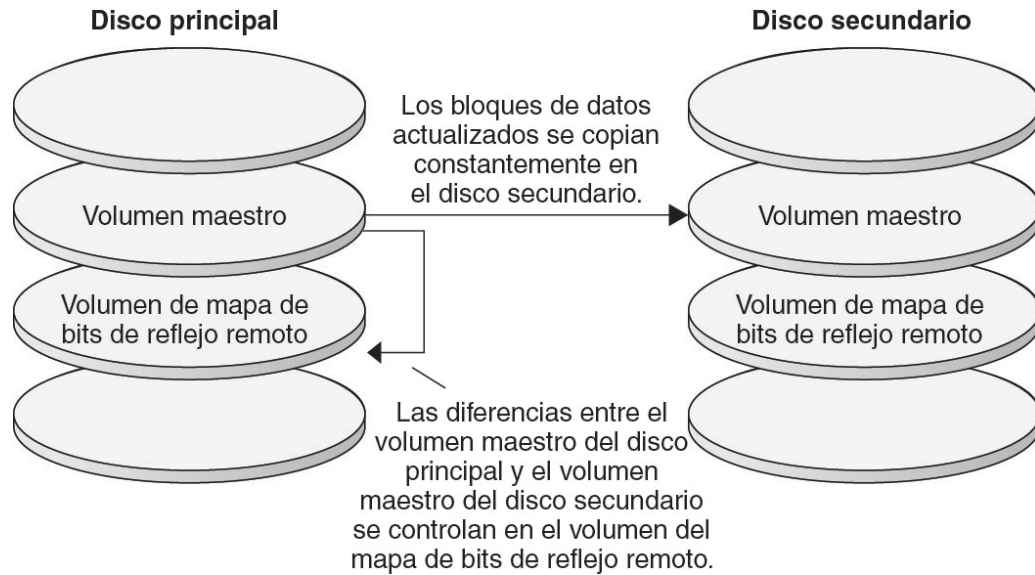
Si los datos se replican de manera síncrona entre el cluster principal y el secundario, no se pierde ningún dato comprometido cuando el sitio principal falla. Sin embargo, si los datos se replican de manera asíncrona, es posible que algunos datos no se repliquen en el cluster secundario antes del fallo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

## Métodos de replicación de datos utilizados por Availability Suite

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

## Replicación de reflejo remoto

La [Figura 2, “Replicación de reflejo remoto”](#) muestra la replicación de reflejo remoto. 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 2** Replicación de reflejo remoto

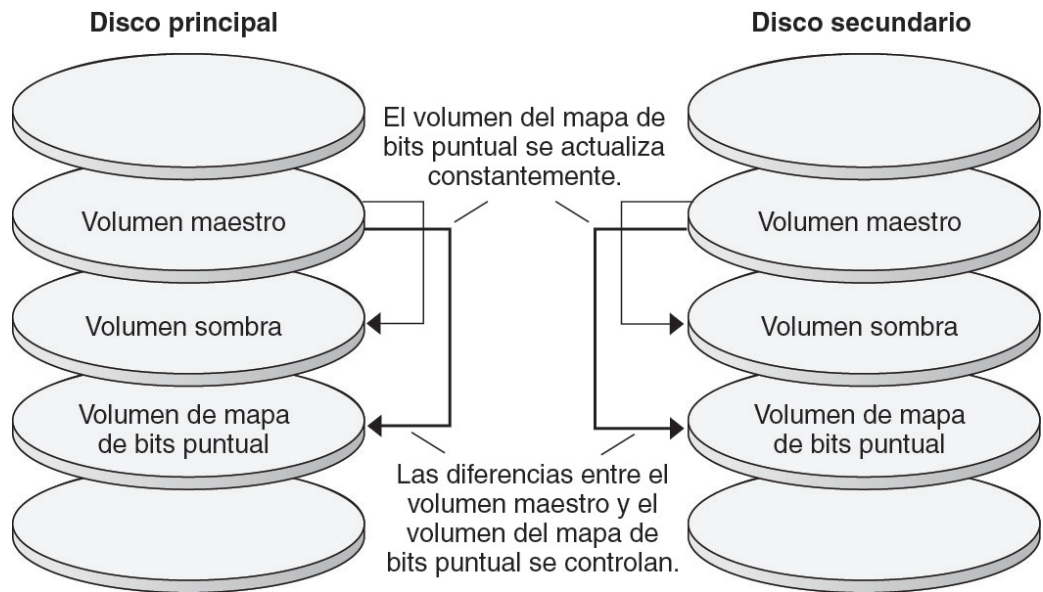
La replicación por duplicación remota se puede efectuar de manera sincrónica en tiempo real o de manera asincrónica. Cada volumen definido en cada cluster se puede configurar de manera individual, para replicación sincrónica o asincrónica.

- En una replicación de datos síncrona, no se confirma la finalización de una operación de escritura hasta que se haya actualizado el volumen remoto.
- En una replicación de datos asíncrona, se confirma la finalización de una operación de escritura antes de que se actualice el volumen remoto. La replicación asincrónica de datos proporciona una mayor flexibilidad en largas distancias y poco ancho de banda.

## Instantánea de un momento determinado

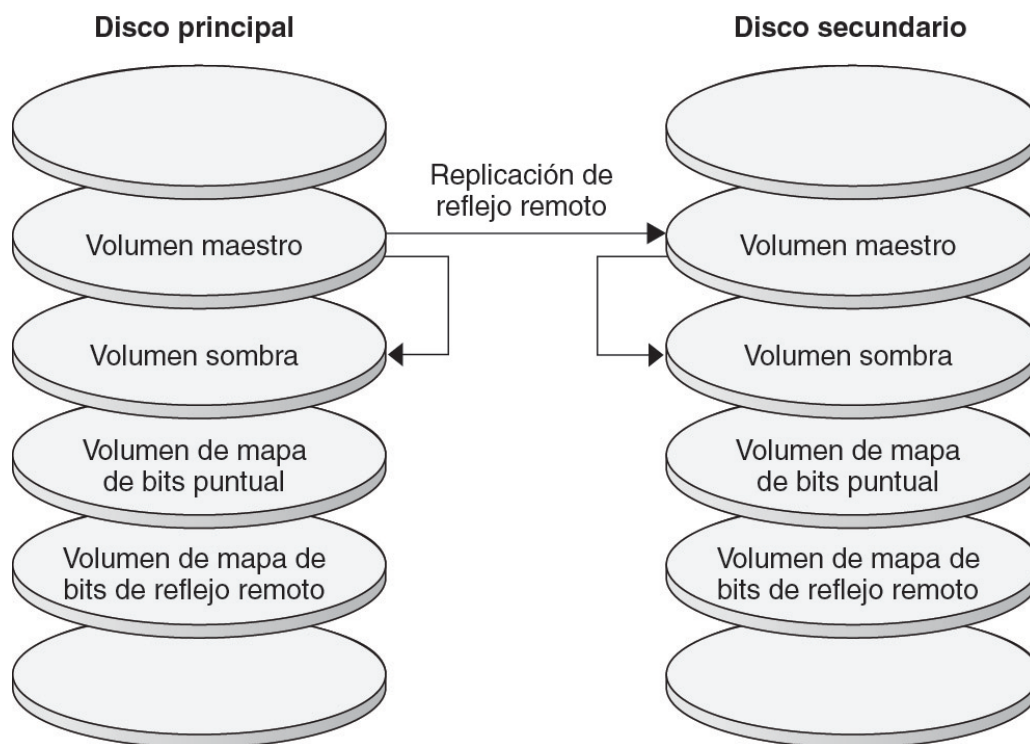
La [Figura 3, “Instantánea de un momento determinado”](#) muestra una instantánea de un momento determinado. Los datos del volumen primario de cada disco se copian en el volumen sombra del mismo disco. El mapa de bits instantáneo controla y detecta las diferencias entre el volumen primario y el volumen sombra. Cuando los datos se copian en el volumen sombra, el mapa de bits de un momento determinado se vuelve a configurar.

**FIGURA 3** Instantánea de un momento determinado



## Replicación en la configuración de ejemplo

La [Figura 4, “Replicación en la configuración de ejemplo”](#) ilustra la forma en que se usan la replicación de reflejo remoto y la instantánea de un momento determinado en este ejemplo de configuración.

**FIGURA 4** 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” \[322\]](#)
- [“Configuración de grupos de recursos de aplicaciones” \[323\]](#)
  - [“Configuración de grupos de recursos en una aplicación de migración tras error” \[323\]](#)

- [“Configuración de grupos de recursos en una aplicación escalable” \[326\]](#)
- [“Directrices para la gestión de una recuperación” \[327\]](#)

## 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 un recurso de nombre de host lógico. Debe existir un nombre de host lógico en cada extremo del flujo de replicación de datos y debe encontrarse en el mismo nodo del cluster que actúa como la ruta de E/S principal para el dispositivo. Un grupo de recursos de replicaciones debe tener las características siguientes:

- Ser un grupo de recursos de migración tras error.  
Un recurso de migración tras error sólo puede ejecutarse en un nodo a la vez. Cuando se produce la migración tras error, los recursos de recuperación participan en ella.
- Debe tener un recurso de nombre de host lógico.  
Un nombre de host lógico se aloja en un nodo de cada cluster (principal y secundario) y se utiliza para proporcionar direcciones de origen y destino para el flujo de replicación de datos del software Availability Suite.
- Tener un recurso 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 siguientes propiedades de extensión se deben definir en el recurso HASToragePlus:

- `GlobalDevicePaths`. La propiedad de esta extensión define el grupo de dispositivos al que pertenece un volumen.
- `AffinityOn property = True`. La propiedad de esta extensión provoca que el grupo de dispositivos se conmute o migre tras error si el grupo de recursos de replicaciones también lo hace. Esta función se denomina *conmutación de afinidad*.

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

- Recibir el nombre del grupo de dispositivos con el que se coloca, seguido de `-stor-rg`.  
Por ejemplo, `devgrp-stor-rg`.
- Estar en línea en el cluster primario y en el secundario

## Configuración de grupos de recursos de aplicaciones

Para que esté totalmente disponible, una aplicación se debe gestionar como un recurso en un grupo de recursos de aplicaciones. Un grupo de recursos de aplicaciones se puede configurar en una aplicación de migración tras error o en una aplicación escalable.

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

Los recursos de aplicaciones y los grupos de recursos de aplicaciones configurados en el 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” \[323\]](#)
- [“Configuración de grupos de recursos en una aplicación escalable” \[326\]](#)

## 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.

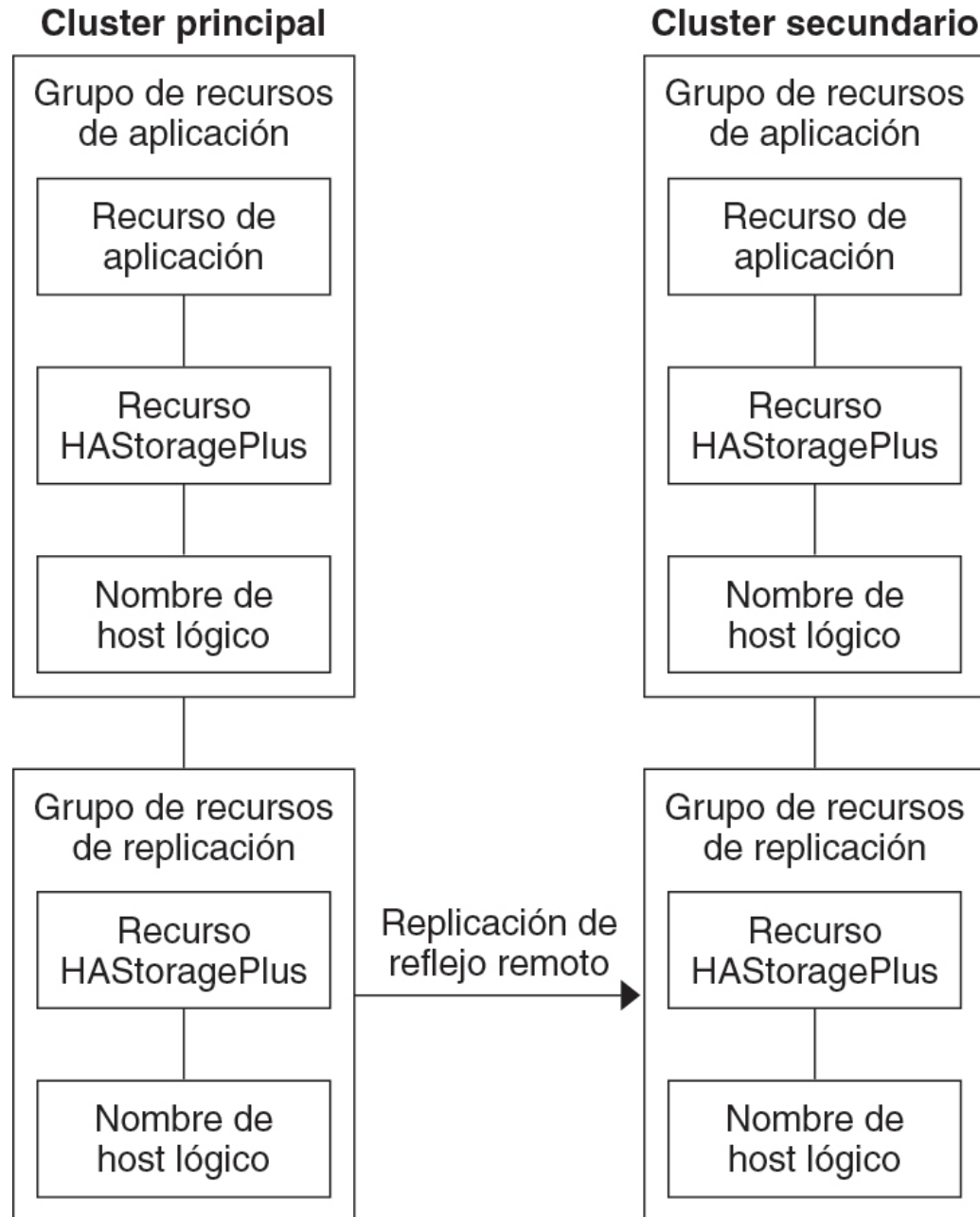
- Si los datos de la aplicación están montados de manera global, la presencia de un recurso de `HAStoragePlus` en el grupo de recursos de aplicaciones no es obligatoria, aunque sí aconsejable.
- Si los datos de la aplicación se montan de manera local, la presencia de un recurso de `HAStoragePlus` en 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 5, “Configuración de grupos de recursos en una aplicación de migración tras error”](#) ilustra la configuración de un grupo de recursos de aplicación y de un grupo de recursos de replicación en una aplicación de conmutación por error.

**FIGURA 5** Configuración de grupos de recursos en una aplicación de migración tras error



## Configuración de grupos de recursos en una aplicación escalable

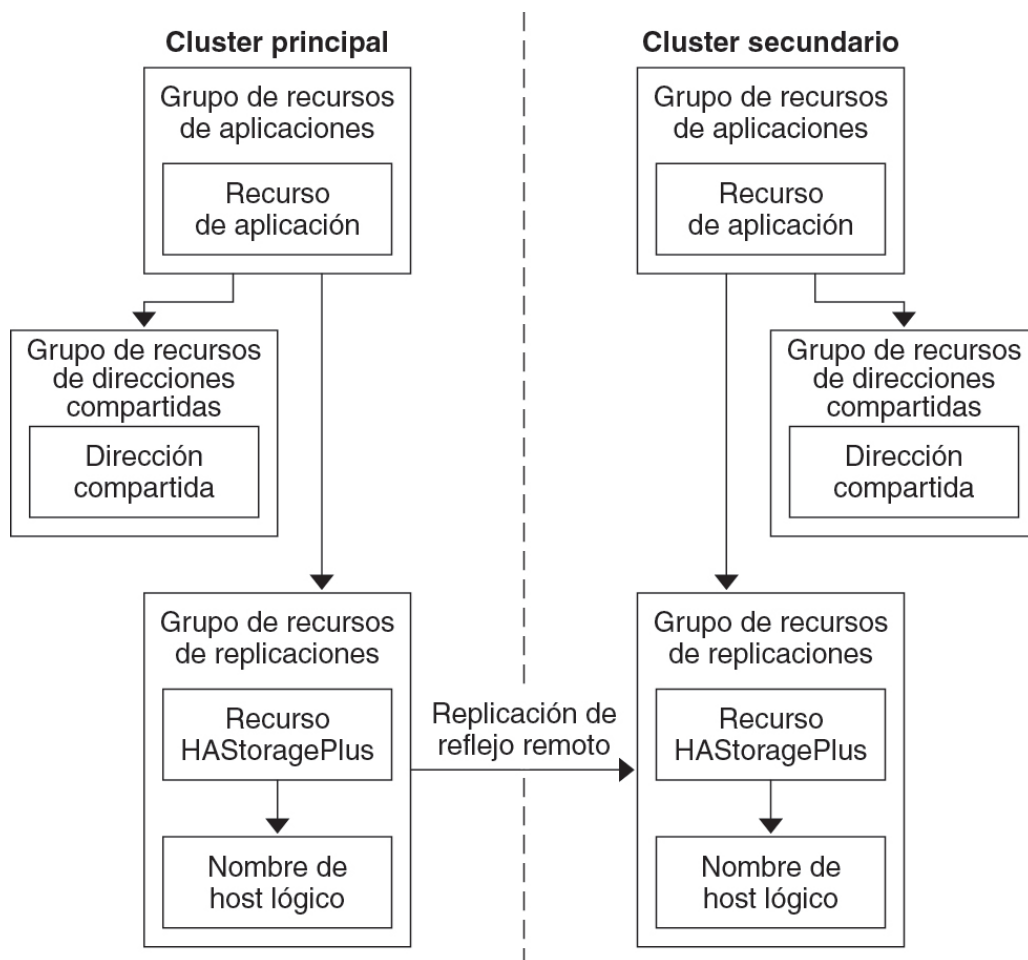
En una aplicación escalable, una aplicación se ejecuta en varios nodos con el fin de crear un único servicio lógico. Si se produce un error en un nodo que ejecuta una aplicación escalable, no habrá migración tras error. La aplicación continúa ejecutándose en los otros nodos.

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 6, “Configuración de grupos de recursos en una aplicación escalable”](#) ilustra la configuración de grupos de recursos en una aplicación escalable.

**FIGURA 6** Configuración de grupos de recursos en una aplicación escalable

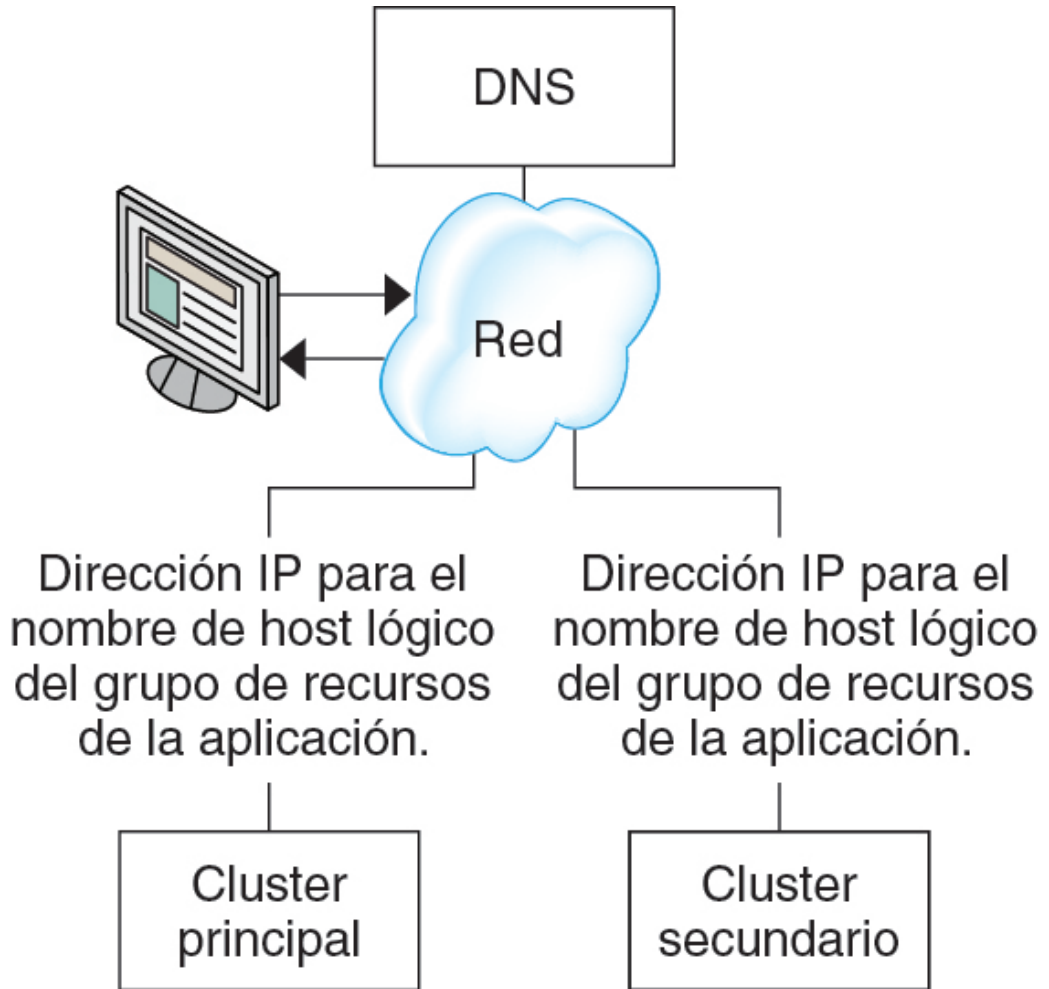
## Directrices para la gestión de una recuperación

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 usan el DNS para asignar el nombre de host lógico de una aplicación a una dirección IP. Después de una recuperación, donde la aplicación se mueve a un 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.

**FIGURA 7** Asignación del DNS de un cliente a un cluster



Si desea actualizar el DNS, utilice el comando `nsupdate`. Para obtener 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](#)” [356].

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

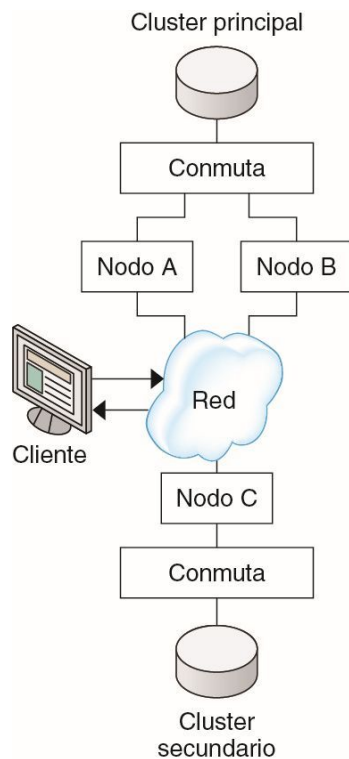
En la [Tabla 23](#), “Mapa de tareas: ejemplo de configuración de una replicación de datos”, se muestran las tareas de este ejemplo de configuración de replicación de datos para una aplicación NFS mediante el software de Availability Suite.

**TABLA 23** Mapa de tareas: ejemplo de configuración de una replicación de datos

| Tarea                                                                                                                                                        | Instrucciones                                                                                                                                            |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1. Conectar e instalar los clusters                                                                                                                          | <a href="#">“Conexión e instalación de clusters” [329]</a>                                                                                               |
| 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 | <a href="#">“Ejemplo de configuración de grupos de dispositivos y de recursos” [331]</a>                                                                 |
| 3. Activar la replicación de datos en el cluster primario y en el secundario                                                                                 | <a href="#">Activación de la replicación en el cluster primario [347]</a><br><a href="#">Activación de la replicación en el cluster secundario [349]</a> |
| 4. Efectuar la replicación de datos                                                                                                                          | <a href="#">Replicación por duplicación remota [350]</a><br><a href="#">Instantánea de un momento determinado [352]</a>                                  |
| 5. Comprobar la configuración de la replicación de datos                                                                                                     | <a href="#">Procedimiento para comprobar la configuración de la replicación [353]</a>                                                                    |

## Conexión e instalación de clusters

La [Figura 8](#), “Ejemplo de configuración del cluster” 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 8** Ejemplo de configuración del cluster

En la [Tabla 24, “Hardware y software necesarios”](#) se resume el hardware y el software que se requieren para la configuración de ejemplo. El sistema operativo Oracle Solaris, el software Oracle Solaris Cluster y el software Volume Manager deben instalarse en los nodos del cluster *antes* de instalar las actualizaciones del software y el software Availability Suite.

**TABLA 24** Hardware y software necesarios

| Hardware o software              | Requisito                                                                                                                                                                                                                                                                    |
|----------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Hardware de nodo                 | El software Availability Suite es compatible con todos los servidores que utilicen el sistema operativo Oracle Solaris.<br><br>Para obtener información sobre el hardware que se debe usar, consulte <a href="#">Oracle Solaris Cluster Hardware Administration Manual</a> . |
| Espacio en disco                 | Aproximadamente 15 MB.                                                                                                                                                                                                                                                       |
| Sistema operativo Oracle Solaris | Las versiones del sistema operativo Oracle Solaris compatibles con Oracle Solaris Cluster. Todos los nodos deben utilizar la misma versión del sistema operativo Oracle Solaris.                                                                                             |

| Hardware o software                               | Requisito                                                                                                                                                                                                                                                                                                                                                                          |
|---------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                                                   | Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software de Oracle Solaris Cluster 4.3</a>                                                                                                                                                                                                                                         |
| Software de Oracle Solaris Cluster                | Software de Oracle Solaris Cluster 4.1, como mínimo.<br><br>Para obtener más información sobre la instalación, consulte <a href="#">Guía de instalación del software de Oracle Solaris Cluster 4.3</a> .                                                                                                                                                                           |
| Software de Volume Manager                        | Software Solaris Volume Manager. Todos los nodos deben usar la misma versión del administrador de volúmenes.<br><br>Para obtener información sobre la instalación, consulte <a href="#">Capítulo 4, “Configuración del software de Solaris Volume Manager” de Guía de instalación del software de Oracle Solaris Cluster 4.3</a> .                                                 |
| Software de Availability Suite                    | Diferentes clusters pueden usar distintas versiones del sistema operativo Oracle Solaris y del software Oracle Solaris Cluster, pero usted debe usar la misma versión del software Availability Suite entre clusters.<br><br>Si desea obtener más información sobre cómo instalar el software, consulte los manuales de instalación de su versión del software Availability Suite: |
| Actualizaciones de software de Availability Suite | Para obtener información sobre las últimas actualizaciones de software, inicie sesión en <a href="#">My Oracle Support</a> .                                                                                                                                                                                                                                                       |

## Ejemplo de configuración de grupos de dispositivos y de recursos

Esta sección describe cómo se configuran los grupos de dispositivos y los de recursos en una aplicación NFS. Para obtener más información, consulte [“Configuración de grupos de recursos de replicaciones” \[322\]](#) y [“Configuración de grupos de recursos de aplicaciones” \[323\]](#).

Esta sección incluye los procedimientos siguientes:

- [Configuración de un grupo de dispositivos en el cluster primario \[332\]](#)
- [Configuración de un grupo de dispositivos en el cluster secundario \[334\]](#)
- [Configuración de sistemas de archivos en el cluster primario para la aplicación NFS \[335\]](#)
- [Configuración del sistema de archivos en el cluster secundario para la aplicación NFS \[336\]](#)
- [Creación de un grupo de recursos de replicaciones en el cluster primario \[337\]](#)
- [Creación de un grupo de recursos de replicaciones en el cluster secundario \[339\]](#)
- [Creación de un grupo de recursos de aplicaciones NFS en el cluster primario \[341\]](#)
- [Creación de un grupo de recursos de aplicaciones NFS en el cluster secundario \[344\]](#)
- [Procedimiento para comprobar la configuración de la replicación \[353\]](#)

La tabla siguiente enumera los nombres de los grupos y recursos creados para la configuración de ejemplo.

**TABLA 25** 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, 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, 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 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 *devicegroupname-stor-rg*.

Para obtener información sobre el software Solaris Volume Manager, consulte [Capítulo 4, “Configuración del software de Solaris Volume Manager” de Guía de instalación del software de Oracle Solaris Cluster 4.3](#).

## ▼ 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:
  - [“Comprensión del software Availability Suite en un cluster” \[318\]](#)
  - [“Directrices para la configuración de replicación de datos basada en host entre clusters” \[321\]](#)
- Configure los clusters primario y secundario como se describe en [“Conexión e instalación de clusters” \[329\]](#).

1. **Para acceder a nodeA, asuma el rol que proporciona la autorización RBAC `solaris.cluster.modify`.**

El nodo nodeA es el primero del cluster primario. Para recordar cuál es el nodo nodeA, consulte la [Figura 8, “Ejemplo de configuración del cluster”](#).

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

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

**3. Agregue discos al metaconjunto.**

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

**4. Agregue mediadores al metaconjunto.**

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

**5. Cree los volúmenes requeridos (o metadispositivos).**

Cree dos componentes de un reflejo:

```
nodeA# metainit -s nfsset d101 1 1 /dev/did/dsk/d6s2
nodeA# metainit -s nfsset d102 1 1 /dev/did/dsk/d7s2
```

Cree el reflejo con uno de los componentes:

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

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

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

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

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

```
nodeA# metainit -s nfsset d200 -p d100 50G
```

- *d201*: el volumen de copia puntual para los datos NFS.

```
nodeA# metainit -s nfsset d201 -p d100 50G
```

- *d202*: el volumen de mapa de bits puntual.

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

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

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

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

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

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

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d200
nodeA# yes | newfs /dev/md/nfsset/rdisk/d204
```

**Pasos siguientes** Vaya a [Configuración de un grupo de dispositivos en el cluster secundario \[334\]](#).

## ▼ 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 \[332\]](#).

1. **Para acceder a nodeC, asuma el rol que proporciona la autorización de RBAC `solaris.cluster.modify`.**

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

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

3. **Agregue discos al metaconjunto.**

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

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

---

**Nota** - Los mediadores no son requeridos en un único cluster de nodo.

---

4. **Cree los volúmenes requeridos (o metadispositivos).**

Cree dos componentes de un reflejo:

```
nodeC# metainit -s nfsset d101 1 1 /dev/did/dsk/d3s2
nodeC# metainit -s nfsset d102 1 1 /dev/did/dsk/d4s2
```

Cree el reflejo con uno de los componentes:

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

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

```
metattach -s nfsset d100 d102
```

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

- **d200:** el volumen maestro de los datos NFS.

```
nodeC# metainit -s nfsset d200 -p d100 50G
```

- **d201:** el volumen de copia puntual para los datos NFS.

```
nodeC# metainit -s nfsset d201 -p d100 50G
```

- **d202:** el volumen de mapa de bits puntual.

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

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

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

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

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

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

```
nodeC# yes | newfs /dev/md/nfsset/rdisk/d200
nodeC# yes | newfs /dev/md/nfsset/rdisk/d204
```

**Pasos siguientes** Vaya a [Configuración de sistemas de archivos en el cluster primario para la aplicación NFS \[335\]](#).

# ▼ 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 \[334\]](#).

1. **En nodeA y nodeB, asuma el rol que proporciona la autorización RBAC solaris.cluster.admin.**
2. **En el nodo nodeA y el nodo nodeB, cree un directorio de punto de montaje para el sistema de archivos NFS.**

Por ejemplo:

```
nodeA# mkdir /global/mountpoint
```

3. **En nodeA y nodeB, configure el volumen maestro para que *no* se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo nodeA y el nodo nodeB. El texto debe estar en una sola línea.

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \
/global/mountpoint ufs 3 no global,logging
```

4. **En nodeA y nodeB, cree un punto de montaje para el metadispositivo d204.**

El ejemplo siguiente crea el punto de montaje `/global/etc`.

```
nodeA# mkdir /global/etc
```

5. **En nodeA y nodeB, configure el metadispositivo d204 para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo nodeA y el nodo nodeB. El texto debe estar en una sola línea.

```
/dev/md/nfsset/dsk/d204 /dev/md/nfsset/rdisk/d204 \
/global/etc ufs 3 yes global,logging
```

**6. Monte el metadispositivo d204 en nodeA.**

```
nodeA# mount /global/etc
```

**7. Cree los archivos de configuración y la información para el servicio de datos NFS de Oracle Solaris Cluster HA.**

**a. Cree un directorio denominado /global/etc/SUNW.nfs en nodeA.**

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

**b. Cree el archivo /global/etc/SUNW.nfs/dfstab.nfs-rs en 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 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 \[336\]](#).

## ▼ 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 \[335\]](#).

**1. En nodeC, asuma el rol que proporciona la autorización RBAC solaris.cluster.admin.**

**2. En el nodo nodeC, cree un directorio de punto de montaje para el sistema de archivos de NFS.**

Por ejemplo:

```
nodeC# mkdir /global/mountpoint
```

**3. En el nodo nodeC, configure el volumen maestro para que se monte automáticamente en el punto de montaje.**

Agregue o sustituya el texto siguiente en el archivo `/etc/vfstab` del nodo `nodeC`. El texto debe estar en una sola línea.

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \
/global/mountpoint ufs 3 yes global,logging
```

#### 4. Monte el metadispositivo `d204` en `nodeA`.

```
nodeC# mount /global/etc
```

#### 5. Cree los archivos de configuración y la información para el servicio de datos NFS de Oracle Solaris Cluster HA.

##### a. Cree un directorio denominado `/global/etc/SUNW.nfs` en `nodeA`.

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

##### b. Cree el archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en `nodeA`.

```
nodeC# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

##### c. Agregue la línea siguiente al archivo `/global/etc/SUNW.nfs/dfstab.nfs-rs` en `nodeA`.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

**Pasos siguientes** Vaya a [Creación de un grupo de recursos de replications en el cluster primario \[337\]](#).

## ▼ Creación de un grupo de recursos de replications en el cluster primario

**Antes de empezar** ■ Complete el procedimiento [Configuración del sistema de archivos en el cluster secundario para la aplicación NFS \[336\]](#).

■ Asegúrese de que el archivo `/etc/netmasks` tenga las entradas de la máscara de red y la subred de la dirección IP para todos los nombres de host lógicos. Si es necesario, edite el archivo `/etc/netmasks` para agregar las entradas que faltan.

#### 1. Acceda a `nodeA` con el rol que proporciona las autorizaciones de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.

#### 2. Registre el tipo de recurso `SUNW.HAStoragePlus`.

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

#### 3. Cree un grupo de recursos de replications para el grupo de dispositivos.

```
nodeA# clresourcegroup create -n nodeA,nodeB devgrp-stor-rg
```

```
-n nodeA,nodeB
```

Especifica que los nodos del cluster nodeA y nodeB pueden controlar el grupo de recursos de replications.

```
devgrp-stor-rg
```

El nombre del grupo de recursos de replications. En este nombre, devgrp especifica el nombre del grupo de dispositivos.

#### 4. Agregue un recurso SUNW.HAStoragePlus al grupo de recursos de replications.

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HAStoragePlus \
-p GlobalDevicePaths=nfsset \
-p AffinityOn=True \
devgrp-stor
```

```
-g
```

Especifica el grupo de recursos en el que se agrega el recurso.

```
-p GlobalDevicePaths=
```

Especifica el grupo de dispositivos del que depende el software 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 de cluster definidos por -p GlobalDevicePaths=. Por ese motivo, si el grupo de recursos de replications migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HAStoragePlus(5)`.

#### 5. Agregue un recurso de nombre de host lógico al grupo de recursos de replications.

```
nodeA# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-prim
```

El nombre de host lógico para el grupo de recursos de replications en el cluster primario es lhost-reprg-prim.

#### 6. Active los recursos, administre el grupo de recursos y póngalo en línea.

```
nodeA# clresourcegroup online -emM -n nodeA devgrp-stor-rg
```

```
-e
```

Activa 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 replications esté en línea para el nodo nodeA.

**Pasos siguientes** Vaya a [Creación de un grupo de recursos de replications en el cluster secundario \[339\]](#).

## ▼ Creación de un grupo de recursos de replications en el cluster secundario

- Antes de empezar**
- Complete el procedimiento [Creación de un grupo de recursos de replications en el cluster primario \[337\]](#).
  - Asegúrese de que el archivo `/etc/netmasks` tenga las entradas de la máscara de red y la subred de la dirección IP para todos los nombres de host lógicos. Si es necesario, edite el archivo `/etc/netmasks` para agregar las entradas que faltan.

1. **Acceda a nodeC con el rol que proporciona las autorizaciones de RBAC `solaris.cluster.modify`, `solaris.cluster.admin` y `solaris.cluster.read`.**

2. **Registre `SUNW.HAStoragePlus` como un tipo de recurso.**

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

3. **Cree un grupo de recursos de replications 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 replications.

**4. Agregue un recurso `SUNW.HAStoragePlus` al grupo de recursos de aplicaciones.**

```
nodeC# clresource create \
-t SUNW.HAStoragePlus \
-p GlobalDevicePaths=nfsset \
-p AffinityOn=True \
devgrp-stor
```

`create`

Crea el recurso.

`-t`

Especifica el tipo de recurso.

`-p GlobalDevicePaths=`

Especifica el grupo de dispositivos en el que se basa el software 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 de cluster definidos por `-p GlobalDevicePaths=`. Por ese motivo, si el grupo de recursos de replications migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

`devgrp-stor`

El recurso de `HAStoragePlus` 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# clresourcegroup create -g devgrp-stor-rg lhost-reprg-sec
```

El nombre de host lógico para el grupo de recursos de replications en el cluster secundario es `lhost-reprg-sec`.

**6. Active los recursos, administre el grupo de recursos y póngalo en línea.**

```
nodeC# clresourcegroup online -eM -n nodeC devgrp-stor-rg
```

`online`

Lo establece en línea.

-e

Activa los recursos asociados.

-M

Administra el grupo de recursos.

-n

Especifica el nodo en el que poner el grupo de recursos en línea.

## 7. Compruebe que el grupo de recursos esté en línea.

```
nodeC# clresourcegroup status devgrp-stor-rg
```

Examine el campo de estado del grupo de recursos para confirmar que el grupo de recursos de 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 \[341\]](#).

## ▼ 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 \[339\]](#).

- Asegúrese de que el archivo `/etc/netmasks` tenga las entradas de la máscara de red y la subred de la dirección IP para todos los nombres de host lógicos. Si es necesario, edite el archivo `/etc/netmasks` para agregar las entradas que faltan.

### 1. Acceda a nodeA con el rol que proporciona 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 registró como un tipo de recurso, hágalo.

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

### 4. Cree un grupo de recursos de aplicaciones para el servicio NFS.

```
nodeA# clresourcegroup create \
-p Pathprefix=/global/etc \
-p Auto_start_on_new_cluster=False \
-p RG_affinities=+++devgrp-stor-rg \
nfs-rg
```

Pathprefix=/global/etc

Especifica el directorio en el que los recursos del grupo pueden guardar los archivos de administración.

Auto\_start\_on\_new\_cluster=False

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

RG\_affinities=+++devgrp-stor-rg

Especifica el grupo de recursos con el que el grupo de recursos de aplicaciones se debe colocar. En este ejemplo, el grupo de recursos de aplicaciones debe colocarse con el grupo de recursos de replications devgrp-stor-rg.

Si el grupo de recursos de replications se conmuta a un nodo principal nuevo, el grupo de recursos de aplicaciones se conmuta automáticamente. Sin embargo, los intentos de conmutar el grupo de recursos de aplicaciones a un nodo principal nuevo se bloquean porque la acción interrumpe el requisito de colocación.

nfs-rg

El nombre del grupo de recursos de aplicaciones.

## 5. Agregue un recurso 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/punto\_montaje

Especifica que el punto de montaje del sistema de archivos es global.

```
-p AffinityOn=True
```

Especifica que el recurso de aplicaciones debe efectuar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos del cluster definidos por `-p FileSystemMountPoints`. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

```
nfs-dg-rs
```

El nombre del recurso de `HAStoragePlus` para la aplicación NFS.

Para obtener más información sobre estas propiedades de extensión, consulte la página del comando `man SUNW.HAStoragePlus(5)`.

## 6. Agregue un recurso de nombre de host lógico al grupo de recursos de aplicaciones.

```
nodeA# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-prim
```

El nombre de host lógico del grupo de recursos de aplicaciones del cluster primario es `lhost-nfsrg-prim`.

## 7. Cree el archivo de configuración `dfstab.resource-name` y colóquelo en el subdirectorío `SUNW.nfs`, en el directorío `Pathprefix` del grupo de recursos que lo contiene.

### a. Cree un directorío denominado `SUNW.nfs` en `nodeA`.

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

### b. Cree un archivo `dfstab.resource-name` en `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 `nodeA`.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

## 8. Establezca en línea el grupo de recursos de aplicaciones.

```
nodeA# clresourcegroup online -M -n nodeA nfs-rg
```

```
online
```

Pone el grupo de recursos en línea.

-e

Activa 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.

## 9. 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 \[344\]](#).

## ▼ 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 \[341\]](#).
  - Asegúrese de que el archivo `/etc/netmasks` tenga las entradas de la máscara de red y la subred de la dirección IP para todos los nombres de host lógicos. Si es necesario, edite el archivo `/etc/netmasks` para agregar las entradas que faltan.

### 1. Acceda a nodeC con el rol que proporciona 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 registró como un 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_affinities=+++devgrp-stor-rg \
nfs-rg
```

create

Crea el grupo de recursos.

-p

Especifica una propiedad del grupo de recursos.

Pathprefix=/global/etc

Especifica un directorio en el que los recursos del grupo pueden guardar los archivos de administración.

Auto\_start\_on\_new\_cluster=False

Especifica que el grupo de recursos de aplicaciones no se inicie de forma automática.

RG\_affinities=+++devgrp-stor-rg

Especifica el grupo de recursos donde el grupo de recursos de aplicaciones debe colocarse. En este ejemplo, el grupo de recursos de aplicaciones debe colocarse con el grupo de recursos de replications devgrp-stor-rg.

Si el grupo de recursos de replications se conmuta a un nodo principal nuevo, el grupo de recursos de aplicaciones se conmuta automáticamente. Sin embargo, los intentos de conmutar el grupo de recursos de aplicaciones a un nodo principal nuevo se bloquean porque la acción interrumpe el requisito de colocación.

nfs-rg

El nombre del grupo de recursos de aplicaciones.

## 5. Agregue un recurso 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/mountpoint`

Especifica que el punto de montaje del sistema de archivos es `global`.

`AffinityOn=True`

Especifica que el recurso de aplicaciones debe efectuar una conmutación de afinidad para los dispositivos globales y los sistemas de archivos del cluster definidos por `-p FileSystemMountPoints=`. Por lo tanto, si el grupo de recursos de aplicaciones migra tras error o se conmuta, el grupo de dispositivos asociados también se conmuta.

`nfs-dg-rs`

El nombre del recurso de `HAStoragePlus` para la aplicación NFS.

**6. Agregue un recurso de nombre de host lógico al grupo de recursos de aplicaciones.**

```
nodeC# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-sec
```

El nombre de host lógico del grupo de recursos de aplicaciones del 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. 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 activar la replicación de datos” \[347\]](#).

## Ejemplo de cómo activar la replicación de datos

Esta sección describe cómo activar la replicación de datos en la configuración de ejemplo. Esta sección utiliza los comandos `sndradm` e `iiadm` del software de Availability Suite. Si desea más información sobre estos comandos, consulte la documentación de Availability Suite.

Esta sección incluye los procedimientos siguientes:

- [Activación de la replicación en el cluster primario \[347\]](#)
- [Activación de la replicación en el cluster secundario \[349\]](#)

### ▼ Activación de la replicación en el cluster primario

1. **Acceda a `nodeA` con el rol que proporciona la autorización de RBAC `solaris.cluster.modify`.**

2. **Vacíe 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. **Active la duplicación por duplicación remota del cluster primario al secundario.**

Este paso permite realizar la replicación del cluster principal al secundario. Este paso permite realizar la replicación del volumen maestro (`d200`) en el cluster principal al volumen maestro (`d200`) en el cluster secundario. Además, este paso permite la replicación en el mapa de bits de reflejo remoto de `d203`.

- Si los clusters primario y secundario no están sincronizados, ejecute este comando para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

- Si el cluster primario y el secundario están sincronizados, ejecute este comando para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -E lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

## 5. Active la sincronización automática.

Ejecute este comando para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -a on lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

Este paso activa 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 Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es logging y el estado activo de la sincronización automática es off. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

## 7. Active la instantánea de un momento determinado.

Utilice el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/iiadm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeA# /usr/sbin/iiadm -w \
/dev/md/nfsset/rdisk/d201
```

Este paso activa 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 d200, el volumen shadow es d201 y el volumen de mapa de bits de un momento determinado es d203.

## 8. Vincule la instantánea de un momento determinado con el grupo de duplicación remota.

Utilice el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

Este paso asocia la instantánea de un momento determinado con el conjunto duplicado remoto de volúmenes. El software Availability Suite comprueba que se tome una instantánea de un momento determinado antes de que pueda haber una replicación por duplicación remota.

**Pasos siguientes** Vaya a [Activación de la replicación en el cluster secundario \[349\]](#).

# ▼ Activación de la replicación en el cluster secundario

**Antes de empezar** Complete el procedimiento [Activación de la replicación en el cluster primario \[347\]](#).

1. **Acceda a nodeC como el rol root.**
2. **Vacíe todas las transacciones.**

```
nodeC# lockfs -a -f
```

3. **Active la duplicación por duplicación remota del cluster primario al secundario.**

Utilice el comando siguiente para el software Availability Suite:

```
nodeC# /usr/sbin/sndradm -n -e lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

El cluster primario detecta la presencia del cluster secundario y comienza la sincronización. Si desea información sobre el estado de los clusters, consulte el archivo de registro del sistema /var/adm de Availability Suite.

4. **Active la instantánea de un momento determinado independiente.**

Utilice el comando siguiente para el software Availability Suite:

```
nodeC# /usr/sbin/iiadm -e ind \
```

```
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeC# /usr/sbin/iiadm -w \
/dev/md/nfsset/rdisk/d201
```

**5. Vincule la instantánea de un momento determinado con el grupo de duplicación remota.**

Utilice el comando siguiente para el software Availability Suite:

```
nodeC# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

**Pasos siguientes** Vaya a [“Ejemplo de cómo efectuar una replicación de datos” \[350\]](#).

## Ejemplo de cómo efectuar una replicación de datos

Esta sección describe la ejecución de la replicación de datos en la configuración de ejemplo. Esta sección utiliza los comandos `sndradm` e `iiadm` del software de Availability Suite. Si desea más información sobre estos comandos, consulte la documentación de Availability Suite.

Esta sección incluye los procedimientos siguientes:

- [Replicación por duplicación remota \[350\]](#)
- [Instantánea de un momento determinado \[352\]](#)
- [Procedimiento para comprobar la configuración de la replicación \[353\]](#)

### ▼ Replicación por duplicación remota

En este procedimiento, el volumen maestro del disco primario se replica en el volumen maestro del disco secundario. El volumen maestro es `d200` y el volumen del mapa de bits de reflejo remoto es `d203`.

- 1. Acceda a nodeA como el rol root.**
- 2. Compruebe que el cluster esté en modo de registro.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

En modo de registro, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco.

### 3. Vacíe todas las transacciones.

```
nodeA# lockfs -a -f
```

### 4. Repita el [Paso 1](#) al [Paso 3](#) en nodeC.

### 5. Copie el volumen principal del nodo nodeA en el volumen principal del nodo nodeC.

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -m lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

### 6. Espere hasta que la termine replicación y los volúmenes se sincronicen.

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

### 7. Confirme que el cluster esté en modo de replicación.

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe en el volumen principal, el software Availability Suite, actualiza el volumen secundario.

**Pasos siguientes** Vaya a [Instantánea de un momento determinado \[352\]](#).

## ▼ 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 d200, el volumen de mapa de bits es d203 y el volumen shadow es d201.

**Antes de empezar** Realice el procedimiento [Replicación por duplicación remota \[350\]](#).

1. **Acceda a nodeA con el rol que proporciona las autorizaciones de RBAC solaris.cluster.modify y solaris.cluster.admin.**

2. **Desactive el recurso que se esté ejecutando en nodeA.**

```
nodeA# clresource disable nfs-rs
```

3. **Cambie el cluster primario a modo de registro.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/md/nfsset/rdsk/d200 \
/dev/md/nfsset/rdsk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdsk/d200 \
/dev/md/nfsset/rdsk/d203 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

4. **Sincronice el volumen sombra del cluster primario con el volumen maestro del cluster primario.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/iidm -u s /dev/md/nfsset/rdsk/d201
nodeA# /usr/sbin/iidm -w /dev/md/nfsset/rdsk/d201
```

5. **Sincronice el volumen sombra del cluster secundario con el volumen maestro del cluster secundario.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeC# /usr/sbin/iidm -u s /dev/md/nfsset/rdsk/d201
nodeC# /usr/sbin/iidm -w /dev/md/nfsset/rdsk/d201
```

6. **Reinicie la aplicación en el nodo nodeA.**

```
nodeA# clresource enable nfs-rs
```

## 7. Vuelva a sincronizar el volumen secundario con el volumen principal.

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

**Pasos siguientes** Vaya a [Procedimiento para comprobar la configuración de la replicación \[353\]](#).

# ▼ Procedimiento para comprobar la configuración de la replicación

**Antes de empezar** Complete el procedimiento [Instantánea de un momento determinado \[352\]](#).

1. **Acceda a nodeA y nodeC con el rol que proporciona 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 Availability Suite:

```
nodeA# /usr/sbin/sndradm -P
```

La salida debe ser similar a la siguiente:

```
/dev/md/nfsset/rdisk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdisk/d200
autosync: on, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: replicating
```

En modo de replicación, el estado es `logging` y el estado activo de la sincronización automática es `off`. Cuando se escribe en el volumen principal, el software 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 Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
```

```
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

**4. Cree un directorio en un equipo cliente.**

**a. Inicie sesión en una máquina cliente como rol root.**

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

**b. Cree un directorio en el equipo cliente.**

```
client-machine# mkdir /dir
```

**5. Monte el volumen principal en el directorio de aplicaciones y visualice el directorio montado.**

**a. Monte el volumen principal en el directorio de aplicaciones.**

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

**b. Visualice el directorio montado.**

```
client-machine# ls /dir
```

**6. Desmonte el volumen principal del directorio de aplicaciones.**

**a. Desmonte el volumen principal del directorio de aplicaciones.**

```
client-machine# umount /dir
```

**b. Ponga el grupo de recursos de aplicaciones fuera de línea en el cluster primario.**

```
nodeA# clresource disable -g nfs-rg +
nodeA# clresourcegroup offline nfs-rg
```

**c. Cambie el cluster primario a modo de registro.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -l lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

Cuando se escribe el volumen de datos en el disco, se actualiza el archivo de mapa de bits en el mismo disco. No se produce ninguna replicación.

- d. Asegúrese de que el directorio PathPrefix esté disponible.

```
nodeC# mount | grep /global/etc
```

- e. Confirme que el sistema de archivos se adecue para ser montado en el cluster secundario.

```
nodeC# fsck -y /dev/md/nfsset/rdisk/d200
```

- f. Establezca la aplicación en el estado gestionado y establézcala en línea en el cluster secundario.

```
nodeC# clresourcegroup online -eM nodeC nfs-rg
```

- g. Acceda a la máquina cliente como rol root.

Verá un indicador con un aspecto similar al siguiente:

```
client-machine#
```

- h. Monte el directorio de aplicaciones creado en el [Paso 4](#) para el directorio de aplicaciones en el volumen secundario.

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

- i. Visualice el directorio montado.

```
client-machine# ls /dir
```

7. Compruebe que el directorio mostrado en el [Paso 5](#) sea el mismo del [Paso 6](#).

8. Devuelva la aplicación del volumen principal al directorio de aplicaciones montado.

- a. Establezca sin conexión el grupo de recursos de aplicaciones en el volumen secundario.

```
nodeC# clresource disable -g nfs-rg +
nodeC# clresourcegroup offline nfs-rg
```

- b. Asegúrese de que el volumen global se desmonte del volumen secundario.

```
nodeC# umount /global/mountpoint
```

- c. Establezca el grupo de recursos de aplicaciones en estado gestionado y establézcalo en línea en el cluster principal.

```
nodeA# clresourcegroup online -eM nodeA nfs-rg
```

**d. Cambie el volumen principal al modo de replicación.**

Ejecute el comando siguiente para el software Availability Suite:

```
nodeA# /usr/sbin/sndradm -n -u lhost-reprg-prim \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 lhost-reprg-sec \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d203 ip sync
```

Cuando se escribe en el volumen principal, el software Availability Suite, actualiza el volumen secundario.

Véase también [“Ejemplo de cómo gestionar una recuperación” \[356\]](#)

## Ejemplo de cómo gestionar una recuperación

En esta sección, se describe cómo actualizar las entradas DNS. Para obtener más información, consulte [“Directrices para la gestión de una recuperación” \[327\]](#).

### ▼ Actualización de la entrada de DNS

Si desea ver una ilustración de cómo asigna el DNS un cliente a un cluster, consulte [Figura 7, “Asignación del DNS de un cliente a un cluster”](#).

**1. Inicie el comando nsupdate.**

Para obtener más 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

- accesibilidad
  - Oracle Solaris Cluster Manager, 313
- acceso
  - Oracle Solaris Cluster Manager, 310
- activación de cables de transporte, 207
- activación y desactivación de una MIB de eventos de SNMP, 259, 260
- actualización, 295
  - Ver también* aplicación de parches
  - a un paquete específico, 294
  - clusters de zona con marca solaris, 295
  - consejos, 293
  - restricción para Oracle GlassFish, 309
  - restricción para Oracle Solaris Cluster Manager, 309
  - servidor AI, 297
  - servidor de quórum, 297
  - visión general, 291, 291
- actualización del espacio de nombres global, 127
- actualizaciones de software
  - visión general, 291
- actualizar manualmente información de DID, 169
- adaptadores de transporte, agregación, 39, 43, 200, 205
- adaptadores, transporte, 205
- administración
  - clusters de zona, 24, 268
  - clusters globales, 24
  - configuración de cluster global, 237
  - dispositivos replicados basados en almacenamiento, 111
  - dispositivos replicados EMC SRDF, 112
  - interconexiones de cluster y redes públicas, 199
  - IPMP, 215
  - PNM, 199
  - quórum, 173
  - sistemas de archivos de cluster, 125
- AffinityOn property
  - propiedad de extensión para replicación de datos, 322
- agregación
  - cables, adaptadores y conmutadores de transporte, 39, 43, 200
  - dirección de red a un cluster de zona, 272
  - dispositivos de quórum, 177
    - NAS de Oracle ZFS Storage Appliance, 178
    - servidor de quórum, 179
  - dispositivos de quórum de disco compartido con conexión directa, 177
  - grupo de dispositivos, 133, 135
  - grupos de dispositivos de Solaris Volume Manager, 135
  - nodos, 219
    - a un cluster de zona, 220
    - a un cluster global, 220
  - roles (RBAC), 65
  - roles personalizados (RBAC), 66
  - sistemas de archivos de cluster, 159
- SNMP
  - hosts, 262
  - usuarios, 263
- ZFS
  - agrupaciones de almacenamiento sin HAStoragePlus, 137
  - grupos de dispositivos replicados, 136
- almacenamiento iSCSI
  - utilizado como dispositivo de quórum
    - con IPMI basadas en sondeos, 215
    - con IPMP basadas en enlaces, 215
- almacenamiento SATA, 177
  - admitido como dispositivo de quórum, 175
- anulación de registro

- grupos de dispositivos de Solaris Volume Manager, 139
- anulación de supervisión
  - rutas de discos, 167
- aplicación de parches, 295
  - Ver también* actualización
- clusters de zona con marca solaris10
  - con clzonecluster, 296
  - con patchadd, 297
  - entorno de inicio alternativo, 297
- aplicaciones de failover para replicación de datos
  - AffinityOn property, 322
  - directrices
    - gestión de recuperación, 327
    - grupos de recursos, 323
  - gestión, 356
  - GlobalDevicePaths, 322
  - ZPoolsSearchDir, 322
- aplicaciones escalables para replicación de datos, 326
- archivo /etc/vfstab, 57
  - agregación de puntos de montaje, 160
  - verificación de la configuración, 161
- archivo /var/adm/messages, 98
- archivo lofi
  - desinstalación, 257
- archivo md.tab, 30
- archivo ntp.conf.sc, 247
- archivo unificado
  - configuración de un cluster de zona, 269
  - instalación de un cluster de zona, 270
  - restauración de un nodo del cluster, 222
- archivo vfstab
  - agregación de puntos de montaje, 160
  - verificación de la configuración, 161
- archivos
  - /etc/vfstab, 57
  - md.conf, 133
  - md.tab, 30
  - ntp.conf.sc, 247
- atributos *Ver* propiedades
- Automated Installer
  - manifiesto, 224
- Availability Suite
  - uso para replicación de datos, 317

## B

- BUI *Ver* Oracle Solaris Cluster Manager
- búsqueda
  - ID de nodo para un cluster de zona, 240
  - ID de nodo para un cluster global, 240

## C

- cables de transporte
  - activación, 207
  - agregación, 39, 43, 200, 205
  - desactivación, 208
- cables, transporte, 205
- cambio
  - nodo principal para un grupo de dispositivos, 150
  - nodos principales, 150
  - nombre de cluster, 238
  - nombres de host privados, 245
  - propiedad numsecondaries, 146
  - propiedades, 145
  - protocolo de MIB de eventos de SNMP, 260
- cambio de nodo principal para un grupo de dispositivos, 150
- cambio de nombre de nodos
  - en un cluster de zona, 248
  - en un cluster global, 248
- cambio de versión
  - zonas de conmutación por error del tipo de marca solaris, 291
  - zonas de conmutación por error del tipo de marca solaris10, 291
- cierre
  - cluster de zona, 69
  - cluster global, 69
  - nodos, 84
  - nodos de cluster de zona, 84
  - nodos del cluster global, 84
- claves de seguridad
  - configuración, 312
- clsetup, 27, 27, 33
  - administración de conmutadores de transporte, 200
  - administración de dispositivos de quórum, 173
  - administración de grupos de dispositivos, 125
  - administración de un cluster de zona, 268

- agregación de una dirección de red a un cluster de zona, 272
- creación de un cluster de zona, 25
- cluster
  - ámbito, 34
  - autenticación de nodos, 240
  - cambio de nombre, 238
  - configuración de hora del día, 242
  - copia de seguridad, 30, 301
  - restauración de archivos, 304
- cluster check
  - comando, 51
- cluster de zona
  - administración, 237
  - cierre, 69
  - clonación, 268
  - configuración desde el archivo unificado, 269
  - creación, 24
  - eliminación de un sistema de archivos, 268
  - estado de componente, 38
  - inicio, 69
  - instalación desde el archivo unificado, 270
  - montajes directos admitidos, 274
  - movimiento de una ruta de zona, 268
  - preparación para aplicaciones, 268
  - validación de configuración, 51
  - visualización de información de configuración, 50
- cluster de zona con marca labeled, 23
- cluster de zona con marca solaris, 23
- cluster de zona con marca solaris10, 23
- cluster de zone
  - agregación de direcciones de red, 272
  - reinicio, 77
- cluster global
  - administración, 237
  - cierre, 69
  - eliminación de nodos, 229
  - estado de componente, 38
  - inicio, 69
  - reinicio, 77
  - validación de configuración, 51
  - visualización de información de configuración, 43
- clusters de campus
  - recuperación con replicación de datos basada en almacenamiento, 107
  - replicación de datos basada en almacenamiento, 104
- clusters de zona
  - con marca solaris
    - actualización, 295
  - con marca solaris10
    - aplicación de parches con clzonecluster, 296
    - aplicación de parches con patchadd, 297
    - aplicación de parches en entorno de inicio alternativo, 297
  - definición, 24
  - marcas, 23
- clusters de zona con marca solaris
  - actualización, 295
- clusters de zona con marca solaris10
  - aplicación de parches con clzonecluster, 296
  - aplicación de parches con patchadd, 297
  - aplicación de parches en un entorno de inicio alternativo, 297
- clusters globales
  - definición, 24
- comando boot, 73
- comando cconsole Ver utilidad pconsole
- comando claccess, 27
- comando cldevice, 27
- comando cldevicegroup, 27
- comando clinterconnect, 27
- comando clnasdevice, 27
- comando clnode, 265, 266
- comando clquorum, 27
- comando clreslogicalhostname, 27
- comando clresource, 27
  - supresión de grupos de recursos, 273
  - supresión de recursos y grupos de recursos, 273
- comando clresourcegroup, 27, 266
- comando clresourcetype, 27
- comando clressharedaddress, 27
- comando clsnmp host, 27
- comando clsnmpmib, 27
- comando clsnmpuser, 27
- comando cltelemattribute, 27
- comando cluster check, 27
  - verificación de archivo vfstab, 161
- comando cluster shutdown, 69
- comando clzonecluster, 27

- aplicación de parches en clusters de zona con marca solaris10, 296
  - descripción, 33
  - detención, 69
  - inicio, 73
- comando metaset, 109
- comando scinstall
  - restauración de un nodo del cluster, 222
- comando showrev -p, 34
- comandos, 27
  - boot, 73
  - cconsole, 32
  - claccess, 27
  - cldevice, 27
  - cldevicegroup, 27
  - clinterconnect, 27
  - clnasdevice, 27
  - clquorum, 27
  - clreslogicalhostname, 27
  - clresource, 27
  - clresourcegroup, 27
  - clresourcetype, 27
  - clressharedaddress, 27
  - clsetup, 27
  - clsnmpghost, 27
  - clsnmpmib, 27
  - clsnmpuser, 27
  - cltelemetryattribute, 27
  - cluster check, 27, 30, 51, 57
  - cluster shutdown, 69
  - clzonecluster, 27, 69
  - clzonecluster boot, 73
  - clzonecluster verify, 51
  - metaset, 109
- comprobación
  - puntos de montaje globales, 57, 164
- conexiones seguras a consolas del cluster, 32
- configuración
  - agrupaciones de almacenamiento ZFS locales sin HAStoragePlus, 137
  - claves de seguridad, 312
  - dispositivos replicados basados en almacenamiento, 111
  - grupo de dispositivos ZFS replicado, 136
  - límites de carga en nodos, 266
  - replicación de datos, 317
  - roles (RBAC), 63
- configuración de hora del cluster, 242
- configuraciones de ejemplo (clusters de campus)
  - dos salas, replicación basada en almacenamiento, 104
- conmutadores de transporte, agregación, 39, 43, 200, 205
- conmutadores, transporte, 205
- consola de administración, 29
- consolas
  - conexión a, 32
- contenedor de agente común
  - configuración de claves de seguridad, 312
- control de acceso basado en roles *Ver* RBAC
- convención de denominación
  - grupos de recursos de replications, 322
- convenciones de denominación
  - dispositivos de disco raw, 160
- copia de seguridad
  - cluster, 30, 301
  - reflejos en línea, 301
- CPU
  - configuración, 288
- creación de reflejo local *Ver* replicación basada en almacenamiento
- creación de reflejo remoto *Ver* replicación basada en almacenamiento
- D**
- desactivación de cables de transporte, 208
- desasociación de una interfaz de red de un grupo IPMP, 216
- desinstalación
  - archivo de dispositivo Iofi, 257
  - paquetes, 298
  - software de Oracle Solaris Cluster, 254
- detención
  - cluster de zona, 77
  - cluster global, 77
  - nodos de cluster de zona, 84
  - nodos del cluster global, 84
- dirección de enlace a la red
  - comprobación, 313

dirección de red  
   agregación a un cluster de zona, 272  
 disco SCSI compartido  
   admitido como dispositivo de quórum, 175  
 dispositivos  
   global, 109  
 dispositivos de disco raw  
   convenciones de denominación, 160  
 dispositivos de quórum  
   agregación, 177  
     dispositivos de quórum de disco compartido con  
     conexión directa, 177  
     NAS de Oracle ZFS Storage Appliance, 178  
     servidor de quórum, 179  
   cambio del timeout predeterminado, 192  
   eliminación, 175, 182  
     último dispositivo de quórum, 183  
   enumeración de configuración, 190  
   estado de mantenimiento, colocación de un  
   dispositivo en, 187  
   estado de mantenimiento, sacar un dispositivo, 189  
   modificación de listas de nodos, 186  
   reconfiguración dinámica de dispositivos, 175  
   reparación, 192  
   sustitución, 185  
   tipos admitidos, 175  
   y replicación basada en almacenamiento, 107  
 dispositivos de quórum de disco compartido con  
 conexión directa  
   agregación, 177  
 dispositivos de quórum de servidor de quórum  
   agregación, 179  
   requisitos de instalación, 179  
   resolución de problemas de eliminación de  
   dispositivos, 183  
 dispositivos replicados basados en almacenamiento  
   administración, 111  
   configuración, 111  
 DLPI, 204  
 dominio invitado, 89

## E

ejemplos  
   ejecución de una comprobación de validación  
   funcional, 55

  enumeración de comprobaciones de validación  
   interactivas, 55  
 ejemplos de configuración (clusters de campus)  
   dos salas, replicación de datos basada en  
   almacenamiento, 104  
 eliminación  
   cables, adaptadores y conmutadores de transporte,  
   205  
   de un cluster de zona, 228  
   dispositivos de quórum, 175, 182  
   grupos de dispositivos de Solaris Volume Manager,  
   139  
   matrices de almacenamiento, 232  
   nodos, 227, 229  
   nodos de todos los grupos de dispositivos, 139  
   recursos y grupos de recursos de un cluster de zona,  
   273  
   sistemas de archivos de cluster, 162  
   SNMP  
     hosts, 263  
     usuarios, 265  
   último dispositivo de quórum, 183  
 EMC SRDF  
   administración, 112  
   configuración de dispositivos DID, 115  
   configuración de grupo de replicación, 112  
   ejemplo de configuración, 118  
   mejores prácticas, 108  
   modo Domino, 105  
   recuperación después de un fallo completo de la sala  
   primaria de un cluster de campus, 122  
   requisitos, 106  
   restricciones, 106  
   verificación de la configuración, 117  
 entorno de inicio alternativo  
   aplicación de parches en clusters de zona con marca  
   solaris10, 297  
 enumeración  
   configuración de quórum, 190  
 espacio de nombres  
   global, 109  
   migración, 129  
 espacio de nombres de dispositivos globales  
   migración, 129  
 estado  
   componente de cluster de zona, 38

- componente de cluster global, 38
- estado de mantenimiento
  - colocación de un dispositivo de quórum en, 187
  - nodos, 250
  - sacar un dispositivo de quórum, 189

## F

- fence\_level *Ver* durante replicación
- fuera de servicio
  - dispositivo de quórum, 187

## G

- gestión de energía, 237
- global
  - dispositivos, 109
    - reconfiguración dinámica, 110
  - espacio de nombres, 109, 127
  - puntos de montaje, comprobación, 57, 164
- GlobalDevicePaths
  - propiedad de extensión para replicación de datos, 322
- globales
  - dispositivos
    - configuración de permisos, 110
- grupo de dispositivos de discos raw
  - agregación, 135
- grupos de dispositivos
  - agregación, 135
  - cambio de propiedades, 145
  - configuración para replicación de datos, 332
  - disco raw
    - agregación, 135
  - eliminación
    - y anulación de registro, 139
  - estado de mantenimiento, 151
  - propiedad principal, 145
  - Solaris Volume Manager
    - agregación, 133
  - visión general de la administración, 125
  - visualización de configuración, 149
- grupos de dispositivos de aplicación
  - configuración para la replicación de datos, 341
- grupos de recursos
  - replicación de datos

- configuración, 322
- directrices para configurar, 321
- rol en failover, 322
- grupos de recursos de aplicaciones
  - directrices, 323
- grupos de recursos de dirección compartida para replicación de datos, 326
- GUI *Ver* Oracle Solaris Cluster Manager

## H

- herramienta de administración de la línea de comandos, 26
- herramientas de administración
  - Oracle Solaris Cluster Manager, 309
- hosts
  - agregación y eliminación de SNMP, 262, 263

## I

- impresión
  - rutas de disco erróneas, 168
- información de DID
  - actualizar manualmente, 169
- información de versión, 34
- inicio
  - cluster de zona, 69, 73
  - cluster global, 69, 73
  - modo sin cluster, 95
  - nodos de cluster de zona, 84, 84
  - nodos del cluster global, 84, 84
  - utilidad pconsole, 225
- inicio de sesión
  - remoto, 32
- inicio de sesión remoto, 32
- inicio en modo sin cluster, 95
- instantánea
  - momento determinado, 319
- instantánea de un momento determinado
  - definición, 319
  - realización, 352
- instantáneas *Ver* replicación basada en almacenamiento
- interconexiones
  - activación, 213
  - resolución de problemas, 213
- interconexiones de cluster

- administración, 199
- interconexiones de clusters
  - reconfiguración dinámica, 201
- IPMP
  - administración, 215
  - comprobación de estado, 42

## K

- /kernel/drv/
  - archivo md.conf, 133

## L

- límites de carga
  - configuración en nodos, 265, 266
  - propiedad `concentrate_load`, 265
  - propiedad `preemption_mode`, 265

## M

- manifiesto
  - Automated Installer, 224
- mantenimiento
  - dispositivo de quórum, 187
- mapa de bits
  - replicación de reflejo remoto, 318
  - una instantánea de un momento determinado, 319
- matrices de almacenamiento
  - eliminación, 232
- mejores prácticas, 108
  - EMC SRDF, 108
  - replicación de datos basada en almacenamiento, 108
- mensajes de error
  - archivo `/var/adm/messages`, 98
  - eliminación de nodos, 234
- MIB
  - activación y desactivación de eventos de SNMP, 259, 260
  - cambio del protocolo de eventos de SNMP, 260
- MIB de eventos
  - activación y desactivación de SNMP, 259, 260
  - cambio de `log_number`, 260
  - cambio de `min_severity`, 260
  - cambio de protocolo SNMP, 260

- migración
  - espacio de nombres de dispositivos globales, 129
- modificación
  - listas de nodos de dispositivo de quórum, 186
  - usuarios (RBAC), 66
- montaje de bucle de retorno
  - exportación de un sistema de archivos a un cluster de zona, 274
- montaje directo
  - exportación de un sistema de archivos a un cluster de zona, 274

## N

- nodos
  - agregación, 219
  - autenticación, 240
  - búsqueda de ID, 240
  - cambio de nombre en un cluster de zona, 248
  - cambio de nombre en un cluster global, 248
  - cierre, 84
  - colocación en estado de mantenimiento, 250
  - conexión a, 32
  - configuración de límites de carga, 266
  - eliminación
    - de grupos de dispositivos, 139
    - de un cluster de zona, 228
    - de un cluster global, 229
    - mensajes de error, 234
  - inicio, 84
  - principal, 110
  - principales, 145
  - secundarios, 145
- nodos de cluster de zona
  - cierre, 84
  - inicio, 84
- nodos del cluster de zona
  - especificación de dirección IP y NIC, 219
  - reinicio, 92
- nodos del cluster global
  - cierre, 84
  - inicio, 84
  - recursos compartidos de CPU, 288
  - reinicio, 92
- nodos del cluster globalglobal
  - verificación

- estado, 226
- nombres de host privados
  - cambio, 245

## O

- opciones de montaje para sistemas de archivos de cluster

- requisitos, 161

- OpenBoot PROM (OBP), 244

- Oracle GlassFish

- restricción de actualización, 309

- Oracle Solaris Cluster Manager, 26, 309

- acceso, 310

- cómo iniciar sesión, 309

- configuración de controles de accesibilidad, 313

- resolución de problemas, 309

- restricción de actualización, 309

- tareas que puede realizar

- acceso a utilidades de configuración de cluster, 33

- activación de la supervisión de una ruta de disco, 166

- activación de un cable, 208

- activación de un dispositivo de quórum, 189

- activación de una interconexión de cluster, 214

- agregación de adaptadores de transporte, 203

- agregación de adaptadores privados, 203

- agregación de almacenamiento compartido a un cluster de zona, 220

- agregación de cables, 203

- agregación de un almacenamiento compartido, 238

- agregación de un dispositivo de almacenamiento a un cluster de zona, 268

- agregación de un dispositivo NAS de Oracle ZFS Storage Appliance, 178

- agregación de un sistema de archivos, 160

- agregación de un sistema de archivos a un cluster de zona, 268

- agregación de una dirección de red a un cluster de zona, 272

- apagado de un cluster de zona, 25

- cierre de un nodo de cluster de zona, 86

- cierre de un nodo de cluster global, 86

- cierre de un nodo del cluster de zona, 228

- colocación de un grupo de dispositivos en línea, 125, 150

- colocación de un grupo de dispositivos fuera de línea, 125, 152

- colocación de un nodo en estado de

- mantenimiento, 251

- comprobación de estado de componente de cluster, 39

- comprobación del estado de la interconexión de cluster, 202

- creación de límites de carga en un nodo del cluster de zona, 266

- creación de límites de carga en un nodo del cluster global, 266

- creación de un cluster de zona, 268

- creación de un dispositivo de quórum, 177

- creación de un servidor de quórum, 180

- desactivación de dispositivo de quórum, 188

- desactivación de la supervisión de una ruta de disco, 168

- desactivación de un cable, 209

- desinstalación de software de un nodo del cluster de zona, 25

- edición de la propiedad de un nodo, 172

- edición de la propiedad Resource Security del cluster de zona, 268

- eliminación de adaptadores de transporte, 205

- eliminación de adaptadores privados, 205

- eliminación de cables, 205

- eliminación de un dispositivo de almacenamiento de un cluster de zona, 277

- eliminación de un dispositivo de quórum, 183

- eliminación de un sistema de archivos, 162

- eliminación de un sistema de archivos de un cluster de zona, 275

- evacuación de un nodo, 70, 250

- inicio de un nodo del cluster de zona, 89

- reinicio de un nodo del cluster de zona, 92

- restablecimiento de dispositivos de quórum, 188

- supresión de un cluster de zona, 273

- visualización de configuración de cluster, 42

- visualización de configuración de zona, 25

- visualización de información de quórum, 191

- visualización de mensajes del sistema del nodo, 32, 33

- visualización de recursos y grupos de recursos, 36
- visualización del estado de un nodo, 41
- tareas que se pueden realizar
  - creación de un cluster de zona, 25
- topología, 309
- uso, 309
- vista de topología, 314
- Oracle ZFS Storage Appliance
  - admitido como dispositivo de quórum, 175
  - agregación como dispositivo de quórum, 178

## P

- paquetes
  - actualización a un paquete específico, 294
  - desinstalación, 298
- patchadd
  - aplicación de parches en clusters de zona con marca `solaris10`, 297
- perfiles
  - derechos de RBAC, 63
- perfiles de derechos
  - RBAC, 63
- permisos, dispositivo global, 110
- planificador por reparto equitativo
  - configuración de recursos compartidos de CPU, 288
- propiedad `failback`, 145
- propiedad `numsecondaries`, 146
- propiedad principal de grupos de dispositivos, 145
- propiedades
  - `failback`, 145
  - `numsecondaries`, 146
  - `preferenced`, 145
- propiedades de extensión para replicación de datos
  - recurso de aplicación, 340, 342, 345
- puntos de montaje
  - global, 57
  - modificación del archivo `/etc/vfstab`, 160

## Q

- quórum
  - administración, 173

- visión general, 173

## R

- RBAC, 63
  - perfiles de derechos (descripción), 63
  - tareas
    - agregación de roles, 65
    - agregación de roles personalizados, 66
    - configuración, 63
    - modificación de usuarios, 66
    - uso, 63
- reconfiguración dinámica, 110, 110
  - dispositivos de quórum, 175
  - interconexiones de clusters, 201
  - interfaces de red pública, 217
- recuperación
  - clusters con replicación de datos basada en almacenamiento, 107
- recurso de nombre de host lógico
  - rol en recuperación de replicación de datos, 322
- recursos
  - supresión, 273
  - visualización de información de configuración, 37
- recursos compartidos de CPU
  - configuración, 287
  - control, 287
  - nodos del cluster global, 288
- red pública
  - administración, 199, 215
  - reconfiguración dinámica, 217
- reflejos, copia de seguridad en línea, 301
- reinicio
  - cluster de zona, 77
  - cluster global, 77
  - nodos del cluster de zona, 92, 92
  - nodos del cluster global, 92, 92
- reparación
  - dispositivo de quórum, 192
- reparación del archivo `/var/adm/messages` completo, 98
- replicación *Ver* replicación de datos
- replicación asíncrona de datos, 105
- replicación de datos, 101
  - activación, 347
  - actualización de una entrada de DNS, 356

- asíncrona, 319
- basada en almacenamiento, 102, 104
- basada en host, 102, 317
- configuración
  - grupos de dispositivos, 332
  - grupos de recursos de aplicación NFS, 341
  - sistemas de archivos para una aplicación NFS, 335
  - switchover de afinidad, 322, 338
- configuración de ejemplo, 329
- definición, 102
- directrices
  - configurar grupos de recursos, 321
  - gestión de recuperación, 327
  - gestión de switchover, 327
- ejemplo, 350
- gestión de recuperación, 356
- grupos de recursos
  - aplicación, 323
  - aplicaciones de failover, 323
  - aplicaciones escalables, 326
  - configuración, 322
  - convención de denominación, 322
  - creación, 337
  - dirección compartida, 326
- hardware y software requeridos, 330
- instantánea de un momento determinado, 319, 352
- introducción, 318
- reflejo remoto, 318, 350
- síncrona, 319
- verificación de configuración, 353
- replicación de datos asíncrona, 319
- replicación de datos basada en almacenamiento, 104
  - definición, 102
  - mejores prácticas, 108
  - recuperación, 107
  - requisitos, 106
  - restricciones, 106
  - y dispositivos de quórum, 107
- replicación de datos basada en host
  - definición, 102
  - ejemplo, 317
- replicación de datos síncrona, 319
- replicación de reflejo remoto
  - definición, 318
  - realización, 350

- replicación remota *Ver* replicación basada en almacenamiento
- replicación síncrona de datos, 105
- replicación, basada en almacenamiento, 104
- resolución de problemas
  - cables, adaptadores y conmutadores de transporte, 213
  - dirección de enlace a la red, 313
  - Oracle Solaris Cluster Manager, 311
- restauración
  - archivos de cluster, 304
  - nodo del cluster
    - con `scinstall`, 222
    - desde un archivo unificado, 222
  - sistema de archivos raíz, 305
- rol
  - agregación de roles, 65
  - agregación de roles personalizados, 66
  - configuración, 63
- ruta de disco
  - anulación de supervisión, 167
  - resolver error de estado, 169
  - supervisión, 109, 165, 166
    - imprimir rutas de disco erróneas, 168
- ruta de zona
  - movimiento, 268
- rutas de disco compartido
  - activación de reinicio automático, 171
  - desactivar reinicio automático, 172
  - supervisión, 165

## S

- secundarias
  - número predeterminado, 145
- secundarios
  - configuración del número deseado, 146
- servidor de quórum de Oracle Solaris Cluster
  - admitido como dispositivo de quórum, 175
- servidores de quórum *Ver* dispositivos de quórum
- sistema de archivos
  - aplicación NFS
    - configuración para replicación de datos, 335
  - eliminación en un cluster de zona, 268
  - restauración de sistema de archivos raíz
    - descripción, 305

- sistema de archivos de red (NFS)
  - configuración de sistemas de archivos de aplicaciones para replicación de datos, 335
  - recurso en el cluster principal, 343
- sistema de nombres de dominio (DNS)
  - actualización en replicación de datos, 356
  - directrices para actualizar, 327
- sistema operativo Oracle Solaris
  - comando `svcadm`, 245
  - control de CPU, 287
  - definición de cluster de zona, 23
  - definición de cluster global, 23
  - instrucciones especiales
    - inicio de nodos, 89
    - reinicio de un nodo, 92
  - replicación basada en almacenamiento, 103
  - replicación basada en host, 103
  - tareas administrativas para un cluster global, 24
- sistema operativo Solaris *Ver* sistema operativo Oracle Solaris
- sistemas de archivos de cluster, 109
  - administración, 125
  - agregación, 159
  - eliminación, 162
  - opciones de montaje, 161
  - verificación de la configuración, 161
- SNMP
  - activación de hosts, 262
  - activación y desactivación de una MIB de eventos, 259, 260
  - agregación de usuarios, 263
  - cambio de protocolo, 260
  - desactivación de hosts, 263
  - eliminación de usuarios, 265
- Solaris Volume Manager
  - nombres de dispositivo de disco raw, 160
- SRDF *Ver* EMC SRDF
- ssh, 32
- supervisión
  - rutas de disco compartido, 171
  - rutas de discos, 166
- supervisión del cluster
  - con la topología de Oracle Solaris Cluster Manager, 314
- sustitución de dispositivos de quórum, 185
- switchback

- directrices para aplicar en replicación de datos, 328
- switchover de afinidad
  - configuración para replicación de datos, 338
- switchover para replicación de datos
  - realización, 356
  - switchover de afinidad, 322

## T

- timeout
  - cambio del valor predeterminado de un dispositivo de quórum, 192
- tipos de dispositivos de quórum admitidos, 175
- tolerancia ante desastres
  - definición, 318
- topología
  - uso en Oracle Solaris Cluster Manager, 314

## U

- último dispositivo de quórum
  - eliminación, 183
- uso
  - roles (RBAC), 63
- usuarios
  - agregación de SNMP, 263
  - eliminación de SNMP, 265
  - modificación de propiedades, 66
- utilidad `clsetup`
  - creación de un cluster de zona, 23
- utilidad `pconsole`, 32
  - conexiones seguras, 32
  - uso, 225

## V

- validación
  - configuración de cluster de zona, 51
  - configuración de cluster global, 51
- verificación
  - configuración de replicación de datos, 353
  - configuración de `vfstab`, 161
  - estado del nodo del cluster, 226
- visión general

- quórum, 173
- visualización
  - configuración de grupo de dispositivos, 149
- volumen *Ver* replicación basada en almacenamiento

## **Z**

### **ZFS**

- agregación
  - agrupaciones de almacenamiento sin HAStoragePlus, 137
  - grupos de dispositivos replicados, 136
- configuración
  - agrupaciones de almacenamiento locales sin HAStoragePlus, 137
  - grupos de dispositivos replicados, 136
- eliminación de un sistema de archivos, 274
- replicación, 136
- restricciones para los sistemas de archivo raíz, 125

ZFS Storage Appliance *Ver* dispositivos de quórum

- zpool
  - configuración local sin HAStoragePlus, 136, 137

ZPoolsSearchDir

- propiedad de extensión para replicación de datos, 322