

Oracle® Solaris Cluster 4.3 시스템 관리 설명 서



부품 번호: E62255
2016년 7월

부품 번호: E62255

Copyright © 2000, 2016, Oracle and/or its affiliates. All rights reserved.

본 소프트웨어와 관련 문서는 사용 제한 및 기밀 유지 규정을 포함하는 라이선스 합의서에 의거해 제공되며, 지적 재산법에 의해 보호됩니다. 라이선스 합의서 상에 명시적으로 허용되어 있는 경우나 법규에 의해 허용된 경우를 제외하고, 어떠한 부분도 복사, 재생, 번역, 방송, 수정, 라이선스, 전송, 배포, 진열, 실행, 발행, 또는 전시될 수 없습니다. 본 소프트웨어를 리버스 엔지니어링, 디스어셈블리 또는 디컴파일하는 것은 상호 운용에 대한 법규에 의해 명시된 경우를 제외하고는 금지되어 있습니다.

이 안의 내용은 사전 공지 없이 변경될 수 있으며 오류가 존재하지 않음을 보증하지 않습니다. 만일 오류를 발견하면 서면으로 통지해 주시기 바랍니다.

만일 본 소프트웨어나 관련 문서를 미국 정부나 또는 미국 정부를 대신하여 라이선스한 개인이나 법인에게 배송하는 경우, 다음 공지사항이 적용됩니다.

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.

본 소프트웨어 혹은 하드웨어는 다양한 정보 관리 애플리케이션의 일반적인 사용을 목적으로 개발되었습니다. 본 소프트웨어 혹은 하드웨어는 개인적인 상해를 초래할 수 있는 애플리케이션을 포함한 본질적으로 위험한 애플리케이션에서 사용할 목적으로 개발되거나 그 용도로 사용될 수 없습니다. 만일 본 소프트웨어 혹은 하드웨어를 위험한 애플리케이션에서 사용할 경우, 라이선스 사용자는 해당 애플리케이션의 안전한 사용을 위해 모든 적절한 비상-안전, 백업, 대비 및 기타 조치를 반드시 취해야 합니다. Oracle Corporation과 그 자회사는 본 소프트웨어 혹은 하드웨어를 위험한 애플리케이션에서의 사용으로 인해 발생하는 어떠한 손해에 대해서도 책임지지 않습니다.

Oracle과 Java는 Oracle Corporation 및/또는 그 자회사의 등록 상표입니다. 기타의 명칭들은 각 해당 명칭을 소유한 회사의 상표일 수 있습니다.

Intel 및 Intel Xeon은 Intel Corporation의 상표 내지는 등록 상표입니다. SPARC 상표 일체는 라이선스에 의거하여 사용되며 SPARC International, Inc.의 상표 내지는 등록 상표입니다. AMD, Opteron, AMD 로고, 및 AMD Opteron 로고는 Advanced Micro Devices의 상표 내지는 등록 상표입니다. UNIX는 The Open Group의 등록상표입니다.

본 소프트웨어 혹은 하드웨어와 관련문서(설명서)는 제3자로부터 제공되는 콘텐츠, 제품 및 서비스에 접속할 수 있거나 정보를 제공합니다. 사용자와 오라클 간의 합의서에 별도로 규정되어 있지 않는 한 Oracle Corporation과 그 자회사는 제3자의 콘텐츠, 제품 및 서비스와 관련하여 어떠한 책임도 지지 않으며 명시적으로 모든 보증에 대해서도 책임을 지지 않습니다. Oracle Corporation과 그 자회사는 제3자의 콘텐츠, 제품 및 서비스에 접속하거나 사용으로 인해 초래되는 어떠한 손실, 비용 또는 손해에 대해 어떠한 책임도 지지 않습니다. 단, 사용자와 오라클 간의 합의서에 규정되어 있는 경우는 예외입니다.

설명서 접근성

오라클의 접근성 개선 노력에 대한 자세한 내용은 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=docacc>에서 Oracle Accessibility Program 웹 사이트를 방문하십시오.

오라클 고객지원센터 액세스

지원 서비스를 구매한 오라클 고객은 My Oracle Support를 통해 온라인 지원에 액세스할 수 있습니다. 자세한 내용은 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=info>를 참조하거나, 청각 장애가 있는 경우 <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=trs>를 방문하십시오.

목차

이 설명서 사용	21
1 Oracle Solaris Cluster 관리 방법 소개	23
Oracle Solaris Cluster 관리 개요	23
영역 클러스터 작업	24
Oracle Solaris OS 기능 제한 사항	25
관리 도구	26
Oracle Solaris Cluster Manager 브라우저 인터페이스	26
명령줄 인터페이스	26
클러스터 관리 준비	28
Oracle Solaris Cluster 하드웨어 구성 문서화	28
관리 콘솔 사용	29
클러스터 백업	29
클러스터 관리	30
클러스터에 원격 로그인	31
클러스터 콘솔에 보안 연결을 설정하는 방법	31
▼ 클러스터 구성 유틸리티에 액세스하는 방법	32
▼ Oracle Solaris Cluster 릴리스 및 버전 정보를 표시하는 방법	33
▼ 구성된 리소스 유형, 리소스 그룹 및 리소스를 표시하는 방법	35
▼ 클러스터 구성요소의 상태를 확인하는 방법	37
▼ 공용 네트워크의 상태를 확인하는 방법	40
▼ 클러스터 구성을 보는 방법	40
▼ 기본 클러스터 구성을 검증하는 방법	50
▼ 전역 마운트 지점을 확인하는 방법	55
▼ Oracle Solaris Cluster 명령 로그의 내용을 보는 방법	57
2 Oracle Solaris Cluster 및 RBAC	59
Oracle Solaris Cluster와 함께 RBAC 설정 및 사용	59
Oracle Solaris Cluster RBAC 권한 프로파일	59
Oracle Solaris Cluster Management 권한 프로파일을 사용하여 RBAC 역할 만 들기 및 지정	61

▼ 명령줄을 사용하여 역할을 만드는 방법	61
사용자의 RBAC 등록 정보 수정	62
▼ 사용자 계정 도구를 사용하여 사용자의 RBAC 등록 정보를 수정하는 방법	62
▼ 명령줄에서 사용자의 RBAC 등록 정보를 수정하는 방법	63
3 클러스터 종료 및 부트	65
클러스터 종료 및 부트 개요	65
▼ 클러스터를 종료하는 방법	66
▼ 클러스터를 부트하는 방법	69
▼ 클러스터를 재부트하는 방법	73
클러스터의 단일 노드 종료 및 부트	79
▼ 노드 종료 방법	80
▼ 노드 부트 방법	83
▼ 노드 재부트 방법	87
▼ 비클러스터 모드로 노드를 부트하는 방법	90
꼭 찬 /var 파일 시스템 복구	92
▼ 꼭 찬 /var 파일 시스템을 복구하는 방법	92
4 데이터 복제 접근 방식	93
데이터 복제 이해	93
지원되는 데이터 복제 방법	94
캠퍼스 클러스터 내에서 저장소 기반 데이터 복제 사용	95
캠퍼스 클러스터 내에서 저장소 기반 데이터 복제를 사용할 경우 요구사항 및 제한 사항	97
캠퍼스 클러스터 내에서 저장소 기반 데이터 복제를 사용할 경우 수동 복구 문제	98
저장소 기반 데이터 복제를 사용할 경우 최고 사례	98
5 전역 장치, 디스크 경로 모니터링 및 클러스터 파일 시스템 관리	99
전역 장치 및 전역 이름 공간 관리 개요	99
Solaris Volume Manager에 대한 전역 장치 사용 권한	100
전역 장치 동적 재구성	100
저장소 기반의 복제된 장치 구성 및 관리	101
EMC Symmetrix Remote Data Facility 복제된 장치 관리	101
클러스터 파일 시스템 관리 개요	113
클러스터 파일 시스템 제한	113
장치 그룹 관리	113
▼ 전역 장치 이름 공간 업데이트 방법	115

▼ 전역 장치 이름 공간에 사용되는 lofi 장치의 크기 변경 방법	116
전역 장치 이름 공간 마이그레이션	117
▼ 전용 파티션에서 lofi 장치로 전역 장치 이름 공간 마이그레이션 방 법	117
▼ lofi 장치에서 전용 파티션으로 전역 장치 이름 공간 마이그레이션 방 법	118
장치 그룹 추가 및 등록	120
▼ 장치 그룹 추가 및 등록 방법(Solaris Volume Manager)	120
▼ 장치 그룹 추가 및 등록 방법(원시 디스크)	122
▼ 복제된 장치 그룹 추가 및 등록 방법(ZFS)	123
▼ HAStoragePlus 없이 로컬 ZFS 저장소 풀을 구성하는 방법	124
장치 그룹 유지 보수	125
장치 그룹 제거 및 등록 해제 방법(Solaris Volume Manager)	126
▼ 모든 장치 그룹에서 노드를 제거하는 방법	126
▼ 장치 그룹에서 노드를 제거하는 방법(Solaris Volume Manager)	127
▼ 원시 디스크 장치 그룹에서 노드를 제거하는 방법	129
▼ 장치 그룹 등록 정보를 변경하는 방법	130
▼ 장치 그룹에 원하는 보조 수를 설정하는 방법	132
▼ 장치 그룹 구성을 나열하는 방법	134
▼ 장치 그룹에 대한 기본 노드를 전환하는 방법	135
▼ 장치 그룹을 유지 보수 상태로 전환하는 방법	136
저장 장치에 대한 SCSI 프로토콜 설정 관리	138
▼ 모든 저장 장치에 대한 기본 전역 SCSI 프로토콜 설정을 표시하는 방 법	138
▼ 단일 저장 장치의 SCSI 프로토콜을 표시하는 방법	139
▼ 모든 저장 장치에 대한 기본 전역 보호(fencing) 프로토콜 설정을 변경하 는 방법	140
▼ 단일 저장 장치에 대한 보호(fencing) 프로토콜을 변경하는 방법	141
클러스터 파일 시스템 관리	143
▼ 클러스터 파일 시스템을 추가하는 방법	143
▼ 클러스터 파일 시스템을 제거하는 방법	146
▼ 클러스터의 전역 마운트를 확인하는 방법	148
디스크 경로 모니터링 관리	148
▼ 디스크 경로를 모니터링하는 방법	149
▼ 디스크 경로의 모니터를 해제하는 방법	150
▼ 오류 디스크 경로를 인쇄하는 방법	151
▼ 디스크 경로 상태 오류를 해결하는 방법	152
▼ 파일의 디스크 경로를 모니터링하는 방법	152
▼ 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 활성화하는 방법	154

▼ 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 비활성화하는 방법	155
6 퀴럼 관리	157
퀴럼 장치 관리	157
퀴럼 장치 동적 재구성	158
퀴럼 장치 추가	159
퀴럼 장치 제거 또는 교체	165
퀴럼 장치 유지 보수	168
퀴럼 기본 시간 초과 변경	174
Oracle Solaris Cluster 퀴럼 서버 관리	175
퀴럼 서버 소프트웨어 시작 및 중지	175
▼ 퀴럼 서버를 시작하는 방법	175
▼ 퀴럼 서버를 중지하는 방법	176
퀴럼 서버에 대한 정보 표시	177
오래된 퀴럼 서버 클러스터 정보 정리	178
7 클러스터 상호 연결 및 공용 네트워크 관리	181
클러스터 상호 연결 관리	181
클러스터 상호 연결 동적 재구성	182
▼ 클러스터 상호 연결 상태를 확인하는 방법	183
▼ 클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치를 추가하는 방법	184
▼ 클러스터 전송 케이블, 전송 어댑터, 전송 스위치를 제거하는 방법	186
▼ 클러스터 전송 케이블을 활성화하는 방법	188
▼ 클러스터 전송 케이블을 비활성화하는 방법	189
▼ 전송 어댑터의 인스턴스 번호를 결정하는 방법	191
▼ 기존 클러스터의 개인 네트워크 주소 또는 주소 범위를 변경하는 방법	192
클러스터 상호 연결 문제 해결	194
공용 네트워크 관리	195
클러스터에서 IP Network Multipathing 그룹을 관리하는 방법	195
공용 네트워크 인터페이스 동적 재구성	197
8 클러스터 노드 관리	199
클러스터 또는 영역 클러스터에 노드 추가	199
▼ 기존 클러스터 또는 영역 클러스터에 노드를 추가하는 방법	200
클러스터 노드 복원	202
▼ 통합 아카이브에서 노드를 복원하는 방법	202

클러스터에서 노드 제거	206
▼ 영역 클러스터에서 노드를 제거하는 방법	207
▼ 클러스터 소프트웨어 구성에서 노드를 제거하는 방법	208
▼ 2 노드보다 많은 연결이 있는 클러스터에서 어레이와 단일 노드 사이의 연결을 제거하는 방법	211
▼ 오류 메시지를 수정하는 방법	213
9 클러스터 관리	215
클러스터 관리 개요	215
▼ 클러스터 이름을 변경하는 방법	216
▼ 노드 ID를 노드 이름에 매핑하는 방법	217
▼ 새 클러스터 노드에 대한 인증 작업 방법	218
▼ 클러스터에서 시간을 설정하는 방법	220
▼ SPARC: 노드에서 OBP(OpenBoot PROM)를 표시하는 방법	221
▼ 노드 개인 호스트 이름을 변경하는 방법	222
▼ 노드의 이름을 바꾸는 방법	225
▼ 기존 Oracle Solaris Cluster 논리 호스트 이름 리소스에서 사용하는 논 리 호스트 이름을 변경하는 방법	226
▼ 노드를 유지 보수 상태로 전환하는 방법	227
▼ 노드를 유지 보수 상태에서 해제하는 방법	229
▼ 클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방 법	231
노드 제거 문제 해결	233
Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관리	234
로드 한계 구성	240
영역 클러스터 관리 작업 수행	243
▼ 통합 아카이브에서 영역 클러스터를 구성하는 방법	244
▼ 통합 아카이브에서 영역 클러스터를 설치하는 방법	245
▼ 영역 클러스터에 네트워크 주소를 추가하는 방법	246
▼ 영역 클러스터를 제거하는 방법	247
▼ 영역 클러스터에서 파일 시스템을 제거하는 방법	248
▼ 영역 클러스터에서 저장 장치를 제거하는 방법	251
테스트 목적으로 사용할 수 있는 문제 해결 절차	253
전역 클러스터 외부에서 응용 프로그램 실행	253
손상된 디스크 세트 복원	255
10 CPU 사용 제어 구성	259
CPU 제어 소개	259
시나리오 선택	259

페어 웨어 스케줄러	260
CPU 제어 구성	260
▼ 전역 클러스터 노드에서 CPU 사용을 제어하는 방법	260
11 소프트웨어 업데이트	263
Oracle Solaris Cluster 소프트웨어 업데이트 개요	263
Oracle Solaris Cluster 소프트웨어 업데이트	264
업데이트 팁	265
특정 패키지 업데이트	265
영역 클러스터 업데이트 또는 패치 적용	267
쿼럼 서버 또는 AI 설치 서버 업데이트	269
패키지 제거	269
▼ 패키지를 제거하는 방법	269
▼ 쿼럼 서버 또는 AI 설치 서버 패키지를 제거하는 방법	270
12 클러스터 백업 및 복원	271
클러스터 백업	271
▼ 미러를 온라인으로 백업하는 방법(Solaris Volume Manager)	271
▼ 클러스터 구성을 백업하는 방법	273
클러스터 파일 복원	273
▼ ZFS 루트(/) 파일 시스템을 복원하는 방법(Solaris Volume Manager)	274
13 Oracle Solaris Cluster Manager 브라우저 인터페이스 사용	277
Oracle Solaris Cluster Manager 개요	277
Oracle Solaris Cluster Manager 소프트웨어 액세스	278
▼ Oracle Solaris Cluster Manager에 액세스하는 방법	278
Oracle Solaris Cluster Manager 문제 해결	279
Oracle Solaris Cluster Manager에 대한 접근성 지원 구성	281
토폴로지를 사용하여 클러스터 모니터링	281
▼ 토폴로지를 사용하여 클러스터를 모니터링하고 업데이트하는 방법	282
A 배치 예제: Availability Suite 소프트웨어를 사용하여 클러스터 간 호스트 기반 데이터 복제 구성	285
클러스터의 Availability Suite 소프트웨어 이해	285
Availability Suite 소프트웨어에서 사용하는 데이터 복제 방법	286
원격 미러 복제	286
포인트 인 타임 스냅샷	287
구성 예에서의 복제	287

클러스터 간 호스트 기반 데이터 복제 구성 지침	288
복제 리소스 그룹 구성	289
응용 프로그램 리소스 그룹 구성	289
테이크오버 관리 지침	292
작업 맵: 데이터 복제 구성의 예	293
클러스터 연결 및 설치	294
장치 그룹 및 리소스 그룹을 구성하는 방법의 예	296
▼ 기본 클러스터에 장치 그룹을 구성하는 방법	297
▼ 보조 클러스터에 장치 그룹을 구성하는 방법	298
▼ NFS 응용 프로그램에서 기본 클러스터 파일 시스템을 구성하는 방법	300
▼ NFS 응용 프로그램에서 보조 클러스터 파일 시스템을 구성하는 방법	301
▼ 기본 클러스터에서 복제 리소스 그룹을 만드는 방법	302
▼ 보조 클러스터에서 복제 리소스 그룹을 만드는 방법	303
▼ 기본 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법	305
▼ 보조 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법	308
데이터 복제 활성화 방법의 예	310
▼ 기본 클러스터에서 복제를 활성화하는 방법	311
▼ 보조 클러스터에서 복제를 활성화하는 방법	312
데이터 복제 수행 방법의 예	313
▼ 원격 미러 복제를 수행하는 방법	314
▼ 포인트 인 타임 스냅샷을 수행하는 방법	315
▼ 복제가 올바르게 구성되었는지 확인하는 방법	316
테이크오버를 관리하는 방법의 예	319
▼ DNS 항목 업데이트 방법	319
 색인	 321

그림

그림 1	저장소 기반 데이터 복제를 사용한 2룸 구성	96
그림 2	원격 미러 복제	286
그림 3	포인트 인 타임 스냅샷	287
그림 4	구성 예에서의 복제	288
그림 5	페일오버 응용 프로그램에서 리소스 그룹 구성	291
그림 6	확장 가능 응용 프로그램에서 리소스 그룹 구성	292
그림 7	클라이언트를 클러스터에 DNS 매핑	293
그림 8	클러스터 구성 예	295

표

표 1	Oracle Solaris Cluster 서비스	25
표 2	Oracle Solaris Cluster 관리 도구	30
표 3	작업 목록: 클러스터 종료 및 부트	66
표 4	작업 맵: 노드 종료 및 부트	80
표 5	작업 맵: 디스크 및 테이프 장치 동적 재구성	101
표 6	작업 맵: EMC SRDF 저장소 기반의 복제된 장치 관리	101
표 7	작업 맵: 장치 그룹 관리	114
표 8	작업 맵: 클러스터 파일 시스템 관리	143
표 9	작업 맵: 디스크 경로 모니터링 관리	148
표 10	작업 목록: 쿼럼 관리	158
표 11	작업 맵: 쿼럼 장치 동적 재구성	159
표 12	작업 목록: 클러스터 상호 연결 관리	182
표 13	작업 맵: 공용 네트워크 인터페이스 동적 재구성	183
표 14	작업 맵: 공용 네트워크 인터페이스 동적 재구성	198
표 15	작업 맵: 기존 전역 또는 영역 클러스터에 노드 추가	200
표 16	작업 맵: 노드 제거	206
표 17	작업 목록: 클러스터 관리	215
표 18	작업 맵: Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관 리	235
표 19	기타 영역 클러스터 작업	243
표 20	Oracle Solaris Cluster 소프트웨어 업데이트	264
표 21	작업 맵: 클러스터 파일 백업	271
표 22	작업 맵: 클러스터 파일 복원	274
표 23	작업 맵: 데이터 복제 구성의 예	294
표 24	필수 하드웨어 및 소프트웨어	295
표 25	구성 예의 그룹 및 리소스 요약	296

코드 예

예 1	Oracle Solaris Cluster 릴리스 및 버전 정보 표시	33
예 2	구성된 리소스 유형, 리소스 그룹 및 리소스 표시	35
예 3	클러스터 구성요소의 상태 확인	37
예 4	공용 네트워크 상태 확인	40
예 5	전역 클러스터 구성 보기	41
예 6	영역 클러스터 구성 보기	48
예 7	모든 기본 검사를 통과한 상태에서 전역 클러스터 구성 검사	52
예 8	대화식 유효성 검사 나열	53
예 9	기능 유효성 검사 실행	53
예 10	실패한 검사가 있는 상태에서 전역 클러스터 구성 검사	54
예 11	전역 마운트 지점 검사	55
예 12	Oracle Solaris Cluster 명령 로그의 내용 보기	58
예 13	smrole 명령을 사용하여 사용자정의 운영자 역할 만들기	61
예 14	영역 클러스터 종료	67
예 15	SPARC: 전역 클러스터 종료	68
예 16	x86: 전역 클러스터 종료	68
예 17	SPARC: 전역 클러스터 부트	70
예 18	x86: 클러스터 부트	70
예 19	영역 클러스터 재부트	74
예 20	SPARC: 전역 클러스터 재부트	75
예 21	x86: 클러스터 재부트	76
예 22	SPARC: 전역 클러스터 노드 종료	81
예 23	x86: 전역 클러스터 노드 종료	82
예 24	영역 클러스터 노드 종료	82
예 25	SPARC: 전역 클러스터 노드 부트	85
예 26	x86: 클러스터 노드 부트	85
예 27	SPARC: 전역 클러스터 노드 재부트	88
예 28	영역 클러스터 노드 재부트	89
예 29	SPARC: 비클러스터 모드로 전역 클러스터 노드 부트	91
예 30	복제본 쌍 만들기	106

예 31	데이터 복제 설정 확인	107
예 32	사용된 디스크에 해당하는 DID 표시	108
예 33	DID 인스턴스 조합	110
예 34	조합된 DID 표시	110
예 35	기본 사이트 페일오버 후 수동으로 EMC SRDF 데이터 복구	112
예 36	전역 장치 이름 공간 업데이트	115
예 37	Solaris Volume Manager 장치 그룹 추가	122
예 38	장치 그룹에서 노드 제거(Solaris Volume Manager)	128
예 39	원시 장치 그룹에서 노드 제거	129
예 40	장치 그룹의 등록 정보 변경	131
예 41	필요한 보조 노드 수 변경(Solaris Volume Manager)	133
예 42	원하는 보조 노드의 수를 기본값으로 설정	134
예 43	모든 장치 그룹의 상태 표시	135
예 44	특정 장치 그룹의 구성 표시	135
예 45	장치 그룹에 대한 기본 노드 전환	136
예 46	장치 그룹을 유지 보수 상태로 만들기	138
예 47	모든 저장 장치에 대한 기본 전역 SCSI 프로토콜 설정 표시	139
예 48	단일 장치의 SCSI 프로토콜 표시	140
예 49	모든 저장 장치에 대한 기본 전역 보호(fencing) 프로토콜 설정 지정	141
예 50	단일 장치의 보호(fencing) 프로토콜 설정	142
예 51	클러스터 파일 시스템 제거	147
예 52	단일 노드의 디스크 경로 모니터링	150
예 53	모든 노드의 디스크 경로 모니터링	150
예 54	CCR의 디스크 구성 다시 읽기	150
예 55	디스크 경로 모니터링 취소	151
예 56	오류 디스크 경로 인쇄	151
예 57	파일을 사용한 디스크 경로 모니터	153
예 58	마지막 퀴럼 장치 제거	167
예 59	퀴럼 장치를 유지 보수 상태로 만들기	171
예 60	퀴럼 투표 수 재설정(퀴럼 장치)	172
예 61	퀴럼 구성 표시	173
예 62	구성된 모든 퀴럼 서버 시작	176
예 63	특정 퀴럼 서버 시작	176
예 64	구성된 모든 퀴럼 서버 중지	176
예 65	특정 퀴럼 서버 중지	177
예 66	한 개의 퀴럼 서버 구성 표시	177
예 67	여러 개의 퀴럼 서버 구성 표시	178
예 68	실행 중인 모든 퀴럼 서버 구성 표시	178
예 69	퀴럼 서버 구성에서 오래된 클러스터 정보 정리	179

예 70	클러스터 상호 연결 상태 확인	184
예 71	클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치 추가 확인	185
예 72	전송 케이블, 전송 어댑터 또는 전송 스위치 제거 확인	188
예 73	클러스터 소프트웨어 구성에서 노드 제거	210
예 74	클러스터 이름 변경	217
예 75	노드 ID를 노드 이름에 매핑	218
예 76	클러스터에 새 시스템 추가 방지	219
예 77	전역 클러스터에 모든 새 시스템 추가 허용	219
예 78	전역 클러스터에 추가할 새 시스템 지정	219
예 79	인증을 표준 UNIX로 설정	219
예 80	인증을 DES로 설정	219
예 81	개인 호스트 이름 변경	224
예 82	전역 클러스터 노드를 유지 보수 상태로 전환	228
예 83	클러스터 노드의 유지 보수 상태 해제 및 쿼럼 투표 수 재설정	230
예 84	전역 클러스터에서 영역 클러스터 제거	248
예 85	영역 클러스터에서 고가용성 로컬 파일 시스템 제거	250
예 86	영역 클러스터에서 고가용성 ZFS 파일 시스템 제거	251
예 87	영역 클러스터에서 Solaris Volume Manager 디스크 세트 제거	252
예 88	영역 클러스터에서 DID 장치 제거	252
예 89	ZFS 루트(/) 파일 시스템 복원(Solaris Volume Manager)	276

이 설명서 사용

Oracle Solaris Cluster 시스템 관리 설명서에서는 SPARC 및 x86 기반 시스템에서 Oracle Solaris Cluster 구성을 관리하는 절차에 대해 설명합니다.

- 개요 - Oracle Solaris Cluster 구성 방법 설명
- 대상 - 기술자, 시스템 관리자 및 공인 서비스 공급자
- 필요한 지식 - 전문적인 하드웨어 문제 해결 및 교체 경력

제품 설명서 라이브러리

이 제품과 관련 제품들에 대한 설명서 및 리소스는 <http://www.oracle.com/pls/topic/lookup?ctx=E62282>에서 사용할 수 있습니다.

피드백

<http://www.oracle.com/goto/docfeedback>에서 이 설명서에 대한 피드백을 보낼 수 있습니다.

Oracle Solaris Cluster 관리 방법 소개

이 장에서는 전역 클러스터 및 영역 클러스터 관리에 대해 다음 정보를 제공하며 Oracle Solaris Cluster 관리 도구 사용 절차를 설명합니다.

- “Oracle Solaris Cluster 관리 개요” [23]
- “영역 클러스터 작업” [24]
- “Oracle Solaris OS 기능 제한 사항” [25]
- “관리 도구” [26]
- “클러스터 관리 준비” [28]
- “클러스터 관리” [30]

이 설명서의 모든 절차는 Oracle Solaris 11 운영체제에서 사용할 수 있습니다.

전역 클러스터는 전역 클러스터 노드 하나 이상으로 구성됩니다. 또한 노드가 아니라 HA for Zones 데이터 서비스로 구성된 solaris 또는 solaris10 브랜드 비전역 영역이 전역 클러스터에 포함될 수도 있습니다.

영역 클러스터는 cluster 속성으로 설정된 solaris, solaris10 또는 labeled 브랜드의 비전역 영역 하나 이상으로 구성됩니다. 다른 브랜드 유형은 영역 클러스터에서 허용되지 않습니다. labeled 브랜드 영역 클러스터는 Oracle Solaris 소프트웨어의 Trusted Extensions 기능 전용입니다. clzonecluster 명령, clsetup 유틸리티 또는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 영역 클러스터를 만듭니다.

Oracle Solaris 영역에서 제공하는 격리를 사용하여 전역 클러스터와 유사한 영역 클러스터에서 지원되는 서비스를 실행할 수 있습니다. 영역 클러스터는 전역 클러스터를 사용하므로 전역 클러스터가 필요합니다. 전역 클러스터에는 영역 클러스터가 포함되지 않습니다. 영역 클러스터에는 시스템에 적어도 하나 이상의 영역 클러스터 노드가 있습니다. 영역 클러스터 노드는 동일한 시스템의 노드가 작동하는 동안에만 작동합니다. 한 시스템의 전역 클러스터 노드가 실패하면 해당 시스템의 모든 영역 클러스터 노드도 실패합니다. 영역 클러스터에 대한 일반적인 정보는 [Oracle Solaris Cluster 4.3 Concepts Guide](#)를 참조하십시오.

Oracle Solaris Cluster 관리 개요

Oracle Solaris Cluster 고가용성 환경에서는 최종 사용자가 중요한 응용 프로그램을 사용할 수 있습니다. 시스템 관리자의 임무는 Oracle Solaris Cluster 구성이 안정적으로 작동하는지 확인하는 것입니다.

관리 작업을 시작하기 전에 다음 설명서의 계획 정보에 익숙해져야 합니다.

- [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서 의 1 장, “Oracle Solaris Cluster 구성 계획”](#)
- [Oracle Solaris Cluster 4.3 Concepts Guide](#)

Oracle Solaris Cluster 관리는 다음과 같은 설명서 간 작업들로 구성됩니다.

주 - 이러한 작업 중 일부는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 수행할 수 있습니다. 이에 대해서는 개별 작업 절차에 설명되어 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

- 정기적으로 또는 매일 전역 클러스터 또는 영역 클러스터를 관리하고 유지 보수하는 데 사용되는 표준 작업. 이 작업은 이 안내서에서 설명합니다.
- 설치, 구성 및 등록 정보 변경과 같은 데이터 서비스 작업. 이러한 작업에 대해서는 [Oracle Solaris Cluster 4.3 데이터 서비스 계획 및 관리 설명서](#)에서 설명합니다.
- 저장소나 네트워크 하드웨어를 추가하거나 복구하는 것 같은 서비스 작업. 이러한 작업에 대해서는 [Oracle Solaris Cluster Hardware Administration Manual](#)에서 설명합니다.

일반적으로 클러스터가 작동하는 동안 Oracle Solaris Cluster 관리 작업을 수행할 수 있습니다. 클러스터 외부 노드가 필요하거나 노드를 종료한 경우에도 나머지 노드가 클러스터 작업을 계속하는 동안에는 관리 작업을 수행할 수 있습니다. 달리 명시되지 않은 경우 전역 클러스터 노드에서 Oracle Solaris Cluster 관리 작업을 수행해야 합니다. 전체 클러스터를 종료해야 하는 절차의 경우에는 정상적인 근무 시간 외로 작동 중지 시간을 예약하여 시스템에 대한 영향을 최소화합니다. 클러스터나 클러스터 노드를 종료할 경우에는 사용자에게 미리 알리십시오.

영역 클러스터 작업

두 개의 Oracle Solaris Cluster 관리 명령(cluster 및 cnode)을 하나의 영역 클러스터에서 실행할 수도 있습니다. 그러나 이러한 명령의 범위는 해당 명령이 실행된 영역 클러스터로 제한됩니다. 예를 들어, 전역 클러스터 노드에서 cluster 명령을 사용하면 전역 클러스터 및 모든 영역 클러스터에 대한 정보가 모두 검색됩니다. 영역 클러스터에서 cluster 명령을 사용하면 특정 영역 클러스터에 대한 정보가 검색됩니다.

전역 클러스터 노드에서 clzonecluster 명령을 사용하는 경우 명령이 전역 클러스터의 모든 영역 클러스터에 영향을 줍니다. 또한 영역 클러스터 명령은 명령 실행 시 해당 영역 클러스터 노드의 작동이 중지된 경우에도 영역 클러스터의 모든 노드에 영향을 줍니다.

영역 클러스터는 RGM(Resource Group Manager)에 의해 제어되는 리소스의 위임 관리를 지원합니다. 따라서 영역 클러스터 관리자는 영역 클러스터 경계를 벗어난 영역 클러스터 종속성을 볼 수 있지만 변경할 수는 없습니다. 전역 클러스터 노드의 관리자만 영역 클러스터 경계를 벗어난 종속성을 만들거나 수정 또는 삭제할 수 있습니다.

다음 목록에는 영역 클러스터에서 수행되는 중요한 관리 작업이 포함되어 있습니다.

- **영역 클러스터 시작 및 재부트** - 3장, [클러스터 종료 및 부트](#)를 참조하십시오. 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터를 부트하고 재부트할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.
- **영역 클러스터에 노드 추가** - 8장, [클러스터 노드 관리](#)를 참조하십시오.
- **영역 클러스터에서 노드 제거** - [영역 클러스터에서 노드를 제거하는 방법 \[207\]](#)을 참조하십시오. 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터 노드에서 소프트웨어를 제거할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.
- **영역 클러스터의 구성 보기** - [클러스터 구성을 보는 방법 \[40\]](#)을 참조하십시오. 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터의 구성을 볼 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.
- **영역 클러스터 구성 검증** - [기본 클러스터 구성을 검증하는 방법 \[50\]](#)을 참조하십시오.
- **영역 클러스터 중지** - 3장, [클러스터 종료 및 부트](#)를 참조하십시오. 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터를 종료할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

Oracle Solaris OS 기능 제한 사항

SMF(서비스 관리 기능) 관리 인터페이스를 사용하여 다음 Oracle Solaris Cluster 서비스를 사용 또는 사용 안함으로 설정하지 마십시오.

표 1 Oracle Solaris Cluster 서비스

Oracle Solaris Cluster 서비스	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

Oracle Solaris Cluster 서비스	FMRI
cl_svc_enable	svc:/system/cluster/cl_svc_enable:default
network-multipathing	svc:/system/cluster/network-multipathing

관리 도구

명령줄 또는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 Oracle Solaris Cluster 구성에 대한 관리 작업을 수행할 수 있습니다. 다음 절에서는 Oracle Solaris Cluster Manager 및 명령줄 도구에 대해 간략하게 설명합니다.

Oracle Solaris Cluster Manager 브라우저 인터페이스

Oracle Solaris Cluster 소프트웨어는 클러스터에서 다양한 관리 작업을 수행하는 데 사용할 수 있는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 지원합니다. 자세한 내용은 [13장. Oracle Solaris Cluster Manager 브라우저 인터페이스 사용](#)을 참조하십시오. [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)에서 Oracle Solaris Cluster Manager 로그인 지침을 확인할 수도 있습니다.

다음은 Oracle Solaris Cluster Manager 에서 수행할 수 있는 몇 가지 작업입니다.

- 영역 클러스터 만들기 및 업데이트
- 리소스 및 리소스 그룹 만들기
- 영역 클러스터에 파일 시스템, 논리 호스트 또는 공유 저장소 추가
- Oracle 데이터베이스 데이터 서비스 만들기
- 전역 클러스터 또는 영역 클러스터에서 노드 관리
- 쿼럼 장치와 서버 추가 및 관리
- NAS 저장 장치 추가 및 관리와 디스크 및 장치 그룹 관리
- Geographic Edition 파트너십 관리

명령줄 인터페이스

clsetup 유틸리티를 통해 대부분의 Oracle Solaris Cluster 관리 작업을 대화식으로 수행할 수 있습니다. 가능한 경우, 이 설명서의 관리 절차에서는 clsetup 유틸리티를 사용합니다.

clsetup 유틸리티를 사용하여 다음과 같은 기본 메뉴 항목을 관리할 수 있습니다.

- 쿼럼

- 리소스 그룹
- 데이터 서비스
- 클러스터 상호 연결
- 장치 그룹 및 볼륨
- 개인 호스트 이름
- 새 노드
- 영역 클러스터
- 기타 클러스터 작업

Oracle Solaris Cluster 구성 관리에 사용되는 기타 명령은 다음 목록에서 제공합니다. 자세한 내용은 매뉴얼 페이지를 참조하십시오.

`if_mpadm(1M)`

IP Network Multipathing 그룹에서 한 어댑터에서 다른 어댑터로 IP 주소를 전환합니다.

`claccess(1CL)`

노드를 추가하기 위한 Oracle Solaris Cluster 액세스 정책을 관리합니다.

`cldevice(1CL)`

Oracle Solaris Cluster 장치를 관리합니다.

`cldevicegroup(1CL)`

Oracle Solaris Cluster 장치 그룹을 관리합니다.

`clinterconnect(1CL)`

Oracle Solaris Cluster 상호 연결을 관리합니다.

`clnasdevice(1CL)`

Oracle Solaris Cluster 구성을 위한 NAS 장치 액세스를 관리합니다.

`clnode(1CL)`

Oracle Solaris Cluster 노드를 관리합니다.

`clquorum(1CL)`

Oracle Solaris Cluster 쿼럼을 관리합니다.

`clreslogicalhostname(1CL)`

논리 호스트 이름에 대한 Oracle Solaris Cluster 리소스를 관리합니다.

`clresource(1CL)`

Oracle Solaris Cluster 데이터 서비스에 대한 리소스를 관리합니다.

`clresourcegroup(1CL)`

Oracle Solaris Cluster 데이터 서비스에 대한 리소스를 관리합니다.

`clresourcetype(1CL)`

Oracle Solaris Cluster 데이터 서비스에 대한 리소스를 관리합니다.

`clressharedaddress(1CL)`

공유 주소에 대한 Oracle Solaris Cluster 리소스를 관리합니다.

`clsetup(1CL)`

영역 클러스터를 만들고 Oracle Solaris Cluster 구성을 대화식으로 구성합니다.

`clsnmphost(1CL)`

Oracle Solaris Cluster SNMP 호스트를 관리합니다.

`clsnmpmib(1CL)`

Oracle Solaris Cluster SNMP MIB를 관리합니다.

`clsnmpuser(1CL)`

Oracle Solaris Cluster SNMP 사용자를 관리합니다.

`cltelemetryattribute(1CL)`

시스템 리소스 모니터링을 구성합니다.

`cluster(1CL)`

Oracle Solaris Cluster 구성의 전역 상태 및 전역 구성을 관리합니다.

`clzonecluster(1CL)`

영역 클러스터를 만들고 수정합니다.

또한 Oracle Solaris Cluster 구성의 볼륨 관리자 부분을 관리하는 명령을 사용할 수 있습니다. 이러한 명령은 클러스터에서 사용하는 특정 볼륨 관리자에 따라 다릅니다.

클러스터 관리 준비

이 절에서는 클러스터 관리를 준비하는 방법에 대해 설명합니다.

Oracle Solaris Cluster 하드웨어 구성 문서화

Oracle Solaris Cluster를 구성하는 규모에 따라 사이트에 맞는 하드웨어 구성을 문서화하십시오. 클러스터를 변경하거나 업그레이드할 때 관리 작업을 줄이려면 하드웨어 문서를 참조하십시오.

하십시오. 여러 클러스터 구성요소 사이의 케이블 및 연결에 레이블을 지정하면 더 쉽게 관리할 수 있습니다.

서비스 담당자가 클러스터에 대한 서비스를 제공할 때 소요 시간을 줄일 수 있도록 클러스터의 초기 구성과 이후의 변경사항에 대한 기록을 유지하십시오.

관리 콘솔 사용

전용 워크스테이션 또는 관리 네트워크를 통해 연결된 워크스테이션을 관리 콘솔로 사용하여 활성 클러스터를 관리할 수 있습니다.

pconsole 유틸리티를 사용해서 각 클러스터 노드에 대해 터미널 창을 실행하고, 동시에 모든 노드에 대해 사용자가 입력하는 명령을 실행하는 마스터 창을 실행할 수 있습니다. 관리 콘솔에서 pconsole 소프트웨어 설치에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “관리 콘솔에 pconsole 소프트웨어를 설치하는 방법”을 참조하십시오.

Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터 및 클러스터 구성요소를 구성, 모니터 및 관리할 수 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

관리 콘솔은 클러스터 노드가 아닙니다. 관리 콘솔은 공용 네트워크 또는 네트워크 기반 단말기 집중 장치를 통해 클러스터 노드에 원격 액세스하는 데 사용됩니다.

Oracle Solaris Cluster에는 전용 관리 콘솔이 필요 없지만 콘솔을 사용하면 다음과 같은 이점이 있습니다.

- 동일한 시스템에서 콘솔과 관리 도구를 그룹화하여 중앙에서 클러스터를 관리할 수 있습니다.
- 엔터프라이즈 서비스 또는 서비스 제공업체에서 더욱 신속하게 문제를 해결할 수 있습니다.

클러스터 백업

정기적으로 클러스터를 백업합니다. Oracle Solaris Cluster 소프트웨어가 저장 장치에 데이터 사본을 미리하여 능률적인 환경을 제공하지만 Oracle Solaris Cluster 소프트웨어가 이것으로 정기적인 백업을 대신할 수는 없습니다. Oracle Solaris Cluster 구성은 장애가 여러 번 발생한 후에도 작동할 수 있지만 사용자나 프로그램 오류 또는 심각한 장애로부터 보호할 수는 없습니다. 따라서 치명적인 데이터 손실로부터 보호할 백업 절차가 있어야 합니다.

백업할 때 다음 정보를 포함해야 합니다.

- 모든 파일 시스템 분할 영역
- DBMS 데이터 서비스를 실행하고 있는 경우에는 모든 데이터베이스 데이터

- 모든 클러스터 디스크에 대한 디스크 분할 영역 정보

클러스터 관리

표 2. “Oracle Solaris Cluster 관리 도구”에서는 클러스터 관리의 시작 위치를 제공합니다.

표 2 Oracle Solaris Cluster 관리 도구

작업	도구	지침
클러스터에 원격 로그인	클러스터에 원격으로 로그인하려면 명령줄에서 Oracle Solaris pconsole 유틸리티를 사용합니다.	“클러스터에 원격 로그인” [31] “클러스터 콘솔에 보안 연결을 설정하는 방법” [31]
대화식으로 클러스터 구성	clzonecluster 명령 또는 clsetup 유틸리티를 사용합니다.	클러스터 구성 유틸리티에 액세스하는 방법 [32]
Oracle Solaris Cluster 릴리스 번호 및 버전 정보 표시	clnode 명령을 show-rev -v -node 하위 명령 및 옵션과 함께 사용합니다.	Oracle Solaris Cluster 릴리스 및 버전 정보를 표시하는 방법 [33]
설치된 리소스, 리소스 그룹 및 리소스 유형 표시	다음 명령을 사용하여 리소스 정보를 표시합니다. ■ clresource ■ clresourcegroup ■ clresourcetype	구성된 리소스 유형, 리소스 그룹 및 리소스를 표시하는 방법 [35]
그래픽으로 클러스터 구성요소 모니터링	Oracle Solaris Cluster Manager를 사용합니다.	온라인 도움말을 참조하십시오.
그래픽으로 일부 클러스터 구성요소 관리	Oracle Solaris Cluster Manager를 사용합니다.	온라인 도움말을 참조하십시오.
클러스터 구성요소의 상태 확인	cluster 명령을 status 하위 명령과 함께 사용합니다.	클러스터 구성요소의 상태를 확인하는 방법 [37]
공용 네트워크에서 IPMP 그룹의 상태 확인	전역 클러스터의 경우 clnode status 명령을 -m 옵션과 함께 사용합니다. 영역 클러스터의 경우 clzonecluster 명령을 show 하위 명령과 함께 사용합니다.	공용 네트워크의 상태를 확인하는 방법 [40]
클러스터 구성을 보십시오.	전역 클러스터의 경우 cluster 명령을 show 하위 명령과 함께 사용합니다. 영역 클러스터의 경우 clzonecluster 명령을 show 하위 명령과 함께 사용합니다.	클러스터 구성을 보는 방법 [40]
구성된 NAS 장치 보기 및 표시	전역 클러스터 또는 영역 클러스터의 경우 clzonecluster 명령을 show 하위 명령과 함께 사용합니다.	clnasdevice(1CL)

작업	도구	지침
전역 마운트 지점 또는 클러스터 구성 확인	전역 클러스터의 경우 <code>cluster</code> 명령을 <code>check</code> 하위 명령과 함께 사용합니다. 영역 클러스터의 경우 <code>clzonecluster verify</code> 명령을 사용합니다.	기본 클러스터 구성을 검증하는 방법 [50]
Oracle Solaris Cluster 명령 로그의 내용 보기	<code>/var/cluster/logs/commandlog</code> 파일을 검사합니다.	Oracle Solaris Cluster 명령 로그의 내용을 보는 방법 [57]
Oracle Solaris Cluster 시스템 메시지 보기	<code>/var/adm/messages</code> 파일을 검사합니다.	Oracle Solaris 11.3의 시스템 관리 문제 해결의 "시스템 메시지 형식" Oracle Solaris Cluster Manager 브라우저 인터페이스에서 노드의 시스템 메시지를 볼 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 Oracle Solaris Cluster Manager에 액세스하는 방법 [278] 을 참조하십시오.
Solaris Volume Manager 상태 모니터링	<code>metastat</code> 명령을 사용합니다.	Solaris Volume Manager Administration Guide

클러스터에 원격 로그인

명령줄에서 병렬 콘솔 액세스(pconsole) 유틸리티를 사용해서 클러스터에 원격으로 로그인할 수 있습니다. pconsole 유틸리티는 Oracle Solaris terminal/pconsole 패키지의 일부입니다. `pkg install terminal/pconsole`을 실행하여 패키지를 설치합니다. pconsole 유틸리티는 사용자가 명령줄에 지정하는 각 원격 호스트에 대해 호스트 단말기 창을 만듭니다. 또한 중앙(마스터) 콘솔 창을 열어서 거기에 입력한 내용을 각 연결로 전파합니다.

pconsole 유틸리티는 X Windows 내에서 또는 콘솔 모드로 실행할 수 있습니다. 클러스터의 관리 콘솔로 사용할 시스템에 pconsole을 설치합니다. 터미널 서버에서 IP 주소에 특정 포트 번호를 연결할 수 있는 경우 호스트 이름이나 IP 주소 외에도 `terminal-server:portnumber`로 포트 번호를 지정할 수 있습니다.

자세한 내용은 `pconsole(1)` 매뉴얼 페이지를 참조하십시오.

클러스터 콘솔에 보안 연결을 설정하는 방법

단말기 집중 장치 또는 시스템 컨트롤러가 ssh를 지원하는 경우 pconsole 유틸리티를 사용하여 해당 시스템의 콘솔에 연결할 수 있습니다. pconsole 유틸리티는 Oracle Solaris terminal/pconsole 패키지의 일부로, 패키지를 설치할 때 함께 설치됩니다. pconsole 유틸리티는 사용자가 명령줄에 지정하는 각 원격 호스트에 대해 호스트 단말기 창을 만듭니다. 또한 중앙(마스터) 콘솔 창을 열어서 거기에 입력한 내용을 각 연결로 전파합니다. 자세한 내용은 `pconsole(1)` 매뉴얼 페이지를 참조하십시오.

▼ 클러스터 구성 유틸리티에 액세스하는 방법

clsetup 유틸리티를 사용하면 대화식으로 영역 클러스터를 만들고 전역 클러스터에 대한 쿼럼, 리소스 그룹, 클러스터 전송, 개인 호스트 이름, 장치 그룹 및 새 노드 옵션을 구성할 수 있습니다. clzonecluster 유틸리티는 영역 클러스터에 대해 이와 유사한 구성 작업을 수행합니다. 자세한 내용은 [clsetup\(1CL\)](#) 및 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 이 절차를 수행할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 전역 클러스터의 활성 구성원 노드에서 root 역할을 수행합니다.
전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.
2. 구성 유틸리티를 시작합니다.

```
phys-schost# clsetup
```

- 전역 클러스터의 경우 clsetup 명령을 사용하여 유틸리티를 시작합니다.

```
phys-schost# clsetup
```

q가 표시됩니다.

- 영역 클러스터의 경우 clzonecluster 명령을 사용하여 유틸리티를 시작합니다. 이 예에서 영역 클러스터는 sczone입니다.

```
phys-schost# clzonecluster configure sczone
```

다음 옵션을 사용하여 유틸리티에서 사용 가능한 작업을 볼 수 있습니다.

```
clzc:sczone> ?
```

또한 대화식 clsetup 유틸리티 또는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 클러스터 범위에서 영역 클러스터를 만들거나 파일 시스템 또는 저장 장치를 추가할 수 있습니다. 다른 모든 영역 클러스터 구성 작업은 `clzonecluster configure` 명령으로 수행합니다. clsetup 유틸리티 또는 Oracle Solaris Cluster Manager를 사용해서 영역 클러스터를 구성하는 방법에 대한 지침은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)를 참조하십시오.

3. 메뉴에서 구성을 선택합니다.

화면의 지시에 따라 작업을 완료하십시오. 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “영역 클러스터 만들기 및 구성”에 설명된 지침을 참조하십시오.

참조 자세한 내용은 `clsetup` 또는 `clzonecluster` 온라인 도움말 매뉴얼 페이지를 참조하십시오.

▼ Oracle Solaris Cluster 릴리스 및 버전 정보를 표시하는 방법

이 절차를 수행하기 위해 `root` 역할로 로그인할 필요는 없습니다. 전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

● Oracle Solaris Cluster 릴리스 및 버전 정보 표시

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

이 명령은 모든 Oracle Solaris Cluster 패키지의 Oracle Solaris Cluster 릴리스 번호와 버전 문자열을 표시합니다.

예 1 Oracle Solaris Cluster 릴리스 및 버전 정보 표시

다음 예에서는 클러스터의 릴리스 정보 및 Oracle Solaris Cluster 4.2와 함께 제공된 패키지의 버전 정보를 표시합니다.

```
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

```

▼ 구성된 리소스 유형, 리소스 그룹 및 리소스를 표시하는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 통해 리소스 및 리소스 그룹을 볼 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

시작하기 전에 root 역할이 아닌 사용자로 이 하위 명령을 사용하려면 solaris.cluster.read RBAC 권한 부여가 필요합니다.

● 클러스터에 구성된 리소스 유형, 리소스 그룹 및 리소스를 표시하십시오.

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

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다. 개인 리소스, 리소스 그룹 및 리소스 유형에 대한 내용을 보려면 다음 명령 중 하나와 함께 show 하위 명령을 사용합니다.

- resource
- resource group
- resourcetype

예 2 구성된 리소스 유형, 리소스 그룹 및 리소스 표시

다음 예에서는 클러스터 schost에 대해 구성된 리소스 유형(RT Name), 리소스 그룹(RG Name) 및 리소스(RS Name)을 보여 줍니다.

```
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

```

▼ 클러스터 구성요소의 상태를 확인하는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터 구성요소의 상태를 확인할 수도 있습니다. 자세한 내용은 Oracle Solaris Cluster Manager 온라인 도움말을 참조하십시오.

Oracle Solaris Cluster Manager 및 `cluster status` 명령도 영역 클러스터의 상태를 표시합니다.

시작하기 전에 root 역할이 아닌 사용자로 `status` 하위 명령을 사용하려면 `solaris.cluster.read` RBAC 권한 부여가 필요합니다.

● 클러스터 구성요소의 상태를 확인하십시오.

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

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

예 3 클러스터 구성요소의 상태 확인

다음 예에서는 `cluster status` 명령으로 반환된 클러스터 구성요소에 대한 상태 정보 샘플을 제공합니다.

```

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
```

▼ 공용 네트워크의 상태를 확인하는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

IP Network Multipathing 그룹의 상태를 확인하려면 `clnode status` 명령과 함께 사용합니다.

시작하기 전에 root 역할이 아닌 사용자로 이 하위 명령을 사용하려면 `solaris.cluster.read` RBAC 권한 부여가 필요합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 노드의 상태를 확인할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

● 클러스터 구성요소의 상태를 확인하십시오.

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

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

예 4 공용 네트워크 상태 확인

다음 예에서는 `clnode status` 명령으로 반환된 클러스터 구성요소에 대한 상태 정보 샘플을 제공합니다.

```
% 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
```

▼ 클러스터 구성을 보는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 통해 클러스터의 구성을 볼 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

시작하기 전에 root 역할이 아닌 사용자로 status 하위 명령을 사용하려면 solaris.cluster.read RBAC 권한 부여가 필요합니다.

● 전역 클러스터 또는 영역 클러스터의 구성을 봅니다.

```
% cluster show
```

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

전역 클러스터 노드에서 cluster show 명령을 실행하면 클러스터에 대한 세부 구성 정보 및 영역 클러스터를 구성한 경우 해당 클러스터에 대한 정보가 표시됩니다.

clzonecluster show 명령을 사용하여 영역 클러스터에 대한 구성 정보만 볼 수도 있습니다. 영역 클러스터의 등록 정보에는 영역 클러스터 이름, IP 유형, 자동 부트 및 영역 경로가 포함됩니다. show 하위 명령은 영역 클러스터 내에서 실행되며 특정 영역 클러스터에만 적용됩니다. 영역 클러스터 노드에서 clzonecluster show 명령을 실행하면 특정 영역 클러스터에 표시되는 객체에 대한 상태만 검색됩니다.

cluster 명령에 대한 추가 정보를 표시하려면 세부 정보 표시 옵션을 사용합니다. 자세한 내용은 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오. clzonecluster에 대한 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

예 5 전역 클러스터 구성 보기

다음 예에서는 전역 클러스터에 대한 구성 정보를 나열합니다. 영역 클러스터가 구성되어 있으면 해당 정보도 나열됩니다.

```
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

```

예 6 영역 클러스터 구성 보기

다음 예에서는 RAC가 포함된 영역 클러스터 구성의 등록 정보를 나열합니다.

```

% 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
```

`clnasdevice show` 하위 명령 또는 Oracle Solaris Cluster Manager를 사용하여 전역 또는 영역 클러스터에 대해 구성된 NAS 장치를 볼 수도 있습니다. 자세한 내용은 [clnasdevice\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 기본 클러스터 구성을 검증하는 방법

`cluster` 명령은 `check` 하위 명령을 사용하여 전역 클러스터가 제대로 작동하는 데 필요한 기본 구성을 검증합니다. 모든 검사가 성공하면 `cluster check`가 셸 프롬프트로 돌아갑니다. 검사가 실패하면 `cluster check`에서 지정된 출력 디렉토리 또는 기본 출력 디렉토리에 보고서를 생성합니다. 둘 이상의 노드에 대해 `cluster check`를 실행하면 `cluster check`에서 각 노드에 대한 보고서와 다중 노드 검사에 대한 보고서를 생성합니다. `cluster list-checks` 명령을 사용하여 사용 가능한 모든 클러스터 검사 목록을 표시할 수도 있습니다.

사용자 상호 작용 없이 실행되는 기본 검사뿐만 아니라 대화식 검사 및 기능 검사도 실행할 수 있습니다. 기본 검사는 `-k keyword` 옵션을 지정하지 않으면 실행됩니다.

- 대화식 검사에는 검사에서 확인할 수 없는 사용자 정보가 필요합니다. 검사에서 펌웨어 버전 번호와 같이 필요한 정보를 사용자에게 요구하는 메시지를 표시합니다. `-k interactive` 키워드를 사용하여 하나 이상의 대화식 검사를 지정합니다.
- 기능 검사에서는 클러스터의 특정 기능이나 동작을 검사합니다. 검사에서는 검사를 시작하거나 계속할지에 대한 확인뿐만 아니라 페일오버할 노드 등에 대한 사용자 입력을 요청합니다. `-k functional check-id` 키워드를 사용하여 기능 검사를 지정합니다. 한 번에 하나의 기능 검사만 수행합니다.

주 - 일부 기능 검사에서는 클러스터 서비스를 중단시키므로 검사에 대한 자세한 설명을 읽고 생산에서 클러스터를 먼저 가져와야 하는지 여부를 결정한 다음에 기능 검사를 시작해야 합니다. 이 정보를 표시하려면 다음 명령을 사용하십시오.

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

`-v` 플래그를 사용하여 자세한 표시 모드로 `cluster check` 명령을 실행하면 진행률 정보를 표시할 수 있습니다.

주 - 장치, 볼륨 관리 구성요소 또는 Oracle Solaris Cluster 구성을 변경할 수 있는 관리 절차를 수행한 후 `cluster check` 명령을 실행합니다.

전역 클러스터 노드에서 `clzonecluster(1CL)` 명령을 실행하면 검사 세트가 실행되어 영역 클러스터가 제대로 작동하는 데 필요한 구성이 검증됩니다. 모든 검사를 통과하면 `clzonecluster verify`가 셸 프롬프트로 돌아가며 영역 클러스터를 안전하게 설치할 수 있습니다.

니다. 검사가 실패하면 `clzonecluster verify`에서 확인이 실패한 전역 클러스터 노드에 대해 보고합니다. 둘 이상의 노드에 대해 `clzonecluster verify` 명령을 실행하면 각 노드에 대한 보고서와 다중 노드 검사에 대한 보고서가 생성됩니다. `verify` 하위 명령은 영역 클러스터 내에서 사용할 수 없습니다.

1. 전역 클러스터의 활성 구성원 노드에서 **root** 역할을 수행합니다.

```
phys-schost# su
```

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

2. 최신 검사가 있는지 확인합니다.

a. [My Oracle Support](#)의 Patches & Updates(패치 & 업데이트) 탭으로 이동합니다.

b. Advanced Search(고급 검색)에서 Product(제품)로 **Solaris Cluster**를 선택하고 Description(설명) 필드에 **check**를 입력합니다.

검색에서 **check**가 포함된 Oracle Solaris Cluster 소프트웨어 업데이트를 찾습니다.

c. 아직 클러스터에 설치되지 않은 소프트웨어 업데이트를 적용합니다.

3. 기본 유효성 검사를 실행합니다.

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

`-v` 상세 정보 모드입니다.

`-o outputdir` `outputdir` 하위 디렉토리로 출력을 리디렉션합니다.

이 명령은 사용 가능한 모든 기본 검사를 실행합니다. 클러스터 기능은 영향을 받지 않습니다.

4. 대화식 유효성 검사를 실행합니다.

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

`-k interactive` 실행 중인 대화식 유효성 검사를 지정합니다.

이 명령은 사용 가능한 대화식 검사를 모두 실행하고 클러스터에 대해 필요한 정보를 요구합니다. 클러스터 기능은 영향을 받지 않습니다.

5. 기능 유효성 검사를 실행합니다.

a. 간단한 표시 모드로 사용 가능한 기능 검사를 모두 나열합니다.

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

- b. 기능 검사에서 생산 환경의 클러스터 가용성 또는 서비스를 방해하는 작업을 수행하는 지 확인합니다.

예를 들어 기능 검사를 수행하면 노드 패닉이 발생하거나 다른 노드로 페일오버될 있습니다.

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

-C check-ID 특정 검사를 지정합니다.

- c. 수행하려는 기능 검사에서 클러스터 작동을 방해할 수 있는 경우 클러스터가 생산 환경에 없는지 확인합니다.

- d. 기능 검사를 시작합니다.

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

-k functional 실행 중인 기능 유효성 검사를 지정합니다.

검사에서 표시되는 메시지에 응답하여 검사가 실행되도록 하고 모든 정보 또는 수행해야 하는 작업을 확인합니다.

- e. 나머지 각 기능 검사를 실행하려면 Step c 와 Step d 를 반복합니다.

주 - 레코드 유지를 위해 실행한 각 검사에 대해 고유한 *outputdir* 하위 디렉토리 이름을 지정합니다. *outputdir* 이름을 다시 사용하면 새 검사의 출력이 다시 사용된 *outputdir* 하위 디렉토리의 기존 내용을 덮어씁니다.

6. 영역 클러스터가 구성된 경우 영역 클러스터의 구성을 검사하여 영역 클러스터를 설치할 수 있는지 여부를 확인합니다.

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

7. 향후 진단을 위해 클러스터 구성을 기록해 둡니다.

[Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “클러스터 구성의 진단 데이터를 기록하는 방법”을 참조하십시오.

예 7 모든 기본 검사를 통과한 상태에서 전역 클러스터 구성 검사

다음 예에서는 모든 검사를 통과한 phys-schost-1 및 phys-schost-2 노드에 대해 상세 정보 모드로 실행되는 `cluster check`를 보여 줍니다.

```
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
```

예 8 대화식 유효성 검사 나열

다음 예에는 클러스터에서 실행하는 데 사용할 수 있는 모든 대화식 검사가 나열되어 있습니다. 예제 출력에서는 가능한 검사의 표본 추출을 보여 주며 실제로 사용 가능한 검사는 각 구성에 따라 다릅니다.

```
# 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?
```

예 9 기능 유효성 검사 실행

다음 예에서는 먼저 자세한 기능 검사 목록을 보여 줍니다. 그런 다음 F6968101 검사에 대한 자세한 설명을 나열하며, 검사에 의해 클러스터 서비스가 중단됨을 나타냅니다. 클러스터는 생산 환경에서 가져옵니다. 그러면 기능 검사가 실행되고 자세한 출력이 funct.test.F6968101.12Jan2011 하위 디렉토리에 기록됩니다. 예제 출력에서는 가능한 검사의 표본 추출을 보여 주며 실제로 사용 가능한 검사는 각 구성에 따라 다릅니다.

```
# 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
```

클러스터를 생산 환경에서 가져옵니다.

```
# 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 <<<

화면의 지시를 따릅니다.
```

예 10 실패한 검사가 있는 상태에서 전역 클러스터 구성 검사

다음 예에서는 /global/phys-schost-1 마운트 지점이 없는 suncluster 클러스터의 phys-schost-2 노드를 보여 줍니다. 보고서는 /var/cluster/logs/cluster_check/<timestamp> 출력 디렉토리에 만들어집니다.

```
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.
...
#
```

▼ 전역 마운트 지점을 확인하는 방법

cluster 명령에는 /etc/vfstab 파일에서 클러스터 파일 시스템과 전역 마운트 지점의 구성 오류를 검사하는 검사가 포함됩니다. 자세한 내용은 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - 장치 또는 볼륨 관리 구성요소에 영향을 주는 클러스터 구성을 변경한 후 cluster check를 실행합니다.

1. 전역 클러스터의 활성 구성원 노드에서 root 역할을 수행합니다.
전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

```
% su
```

2. 전역 클러스터 구성을 확인합니다.

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

예 11 전역 마운트 지점 검사

다음 예에서는 /global/schost-1 마운트 지점이 없는 suncluster 클러스터의 phys-schost-2 노드를 보여 줍니다. 보고서는 /var/cluster/logs/cluster_check/<timestamp>/출력 디렉토리로 전송됩니다.

```
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.
...
#

```

▼ Oracle Solaris Cluster 명령 로그의 내용을 보는 방법

/var/cluster/logs/commandlog ASCII 텍스트 파일에는 클러스터에서 실행되는 선택된 Oracle Solaris Cluster 명령이 기록되어 있습니다. 명령의 기록은 클러스터를 설정할 때 자동으로 시작되어 클러스터를 종료할 때 끝납니다. 명령은 클러스터 모드로 실행되고 부트된 모든 노드에 기록됩니다.

이 파일에 기록되지 않은 명령에는 클러스터의 구성 및 현재 상태를 표시하는 명령이 포함되어 있습니다.

이 파일에 기록된 명령에는 클러스터의 현재 상태를 구성하고 변경하는 명령이 포함되어 있습니다.

- claccess
- cldevice
- cldevicegroup
- clinterconnect
- clnasdevice
- clnode
- clquorum
- clreslogicalhostname
- clresource
- clresourcegroup
- clresourcetype
- clressharedaddress
- clsetup
- clsnmp host
- clsnmpmib
- clsnmpuser
- cltelemetryattribute
- cluster
- clzonecluster

commandlog 파일의 기록에는 다음 요소가 포함될 수 있습니다.

- 날짜 및 시간 기록
- 명령이 실행된 호스트 이름
- 명령의 프로세스 ID
- 명령을 실행한 사용자의 로그인 이름
- 모든 옵션 및 피연산자를 포함하여 사용자가 실행한 명령

주 - 명령 옵션은 셸에서 쉽게 식별하고 복사, 붙여넣기 및 실행할 수 있도록 commandlog 파일에서 따옴표로 묶여 있습니다.

■ 실행된 명령의 종료 상태

주 - 명령이 알 수 없는 결과로 비정상적으로 중단된 경우 Oracle Solaris Cluster 소프트웨어는 commandlog 파일에서 종료 상태로 나타나지 않습니다.

기본적으로 commandlog 파일은 일주일에 한 번씩 저장됩니다. commandlog 파일에 대한 아카이빙 정책을 변경하려면, 클러스터의 각 노드에서 crontab 명령을 사용합니다. 자세한 내용은 [crontab\(1\)](#) 매뉴얼 페이지를 참조하십시오.

Oracle Solaris Cluster 소프트웨어는 모든 지정된 시간에 최대 8개의 사전 저장된 commandlog 파일을 각 클러스터 노드에 유지합니다. 현재 주간의 commandlog 파일 이름은 commandlog입니다. 가장 최근 주간의 파일 이름은 commandlog.0입니다. 가장 오래된 주간의 파일 이름은 commandlog.7입니다.

● 현재 주간의 commandlog 파일 항목을 한 번에 한 화면씩 봅니다.

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

예 12 Oracle Solaris Cluster 명령 로그의 내용 보기

다음 예에서는 more 명령을 실행하여 표시된 commandlog 파일의 내용을 보여 줍니다.

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

Oracle Solaris Cluster 및 RBAC

이 장에서는 Oracle Solaris Cluster와 관련된 역할 기반 액세스 제어(Role-Based Access Control, RBAC)에 대해 설명합니다. 주요 내용은 다음과 같습니다.

- “Oracle Solaris Cluster와 함께 RBAC 설정 및 사용” [59]
- “Oracle Solaris Cluster RBAC 권한 프로파일” [59]
- “Oracle Solaris Cluster Management 권한 프로파일을 사용하여 RBAC 역할 만들기 및 지정” [61]
- “사용자의 RBAC 등록 정보 수정” [62]

Oracle Solaris Cluster와 함께 RBAC 설정 및 사용

다음 표에는 RBAC 설정 및 사용 방법에 관한 참조 문서가 나열되어 있습니다. Oracle Solaris Cluster 소프트웨어와 함께 RBAC를 설정 및 사용하기 위해 수행할 특정 단계는 이 장 후반부에서 설명합니다.

작업	지침
RBAC에 대한 자세한 내용	Oracle Solaris 11.3의 사용자 및 프로세스 보안 의 1 장, “권한을 사용하여 사용자 및 프로세스 제어 정보”
RBAC의 설정, 요소 관리 및 사용	Oracle Solaris 11.3의 사용자 및 프로세스 보안 의 3 장, “Oracle Solaris에서 권한 지정”
RBAC 요소과 도구에 대한 자세한 내용	Oracle Solaris 11.3의 사용자 및 프로세스 보안 의 8 장, “Oracle Solaris 권한에 대한 참조”

Oracle Solaris Cluster RBAC 권한 프로파일

명령줄에서 실행하는 선택한 Oracle Solaris Cluster 명령과 옵션에서 권한 부여에 RBAC를 사용합니다. RBAC 권한 부여가 필요한 Oracle Solaris Cluster 명령 및 옵션에는 다음과 같은 권한 부여 레벨이 하나 이상 필요합니다. Oracle Solaris Cluster RBAC 권한 프로파일은 전역 클러스터의 노드에 적용됩니다.

`solaris.cluster.read`

list, show 및 다른 읽기 작업에 필요한 인증

`solaris.cluster.admin`

클러스터 객체의 상태를 변경하기 위한 인증

`solaris.cluster.modify`

클러스터 객체의 등록 정보를 만들고, 삭제 및 변경하기 위한 권한 부여

Oracle Solaris Cluster 명령에 필요한 RBAC 권한 부여에 대한 자세한 내용은 명령 매뉴얼 페이지를 참조하십시오.

RBAC 권한 프로파일에는 하나 이상의 RBAC 인증이 포함됩니다. 이 권한 프로파일을 사용자나 역할에 지정하여 Oracle Solaris Cluster에 대한 서로 다른 수준의 액세스를 제공할 수 있습니다. Oracle에서는 Oracle Solaris Cluster 소프트웨어에 다음과 같은 권한 프로파일을 제공합니다.

권한 프로파일	권한 부여 및 보안 속성	부여된 권한
Oracle Solaris Cluster 명령	euid=0 보안 속성으로 실행되는 Oracle Solaris Cluster 명령 목록입니다.	모든 Oracle Solaris Cluster 명령에 대해 다음 하위 명령을 포함하여 클러스터를 구성 및 관리하는 데 사용하는 선택한 Oracle Solaris Cluster 명령을 실행합니다. ■ list ■ show ■ status scha_control scha_resource_get scha_resource_setstatus scha_resourcegroup_get scha_resourcetype_get
기본 Oracle Solaris 사용자	이 기존 Oracle Solaris 권한 프로파일에는 다음 권한 부여를 비롯하여 Oracle Solaris 권한 부여가 포함되어 있습니다.	solaris.cluster.read - Oracle Solaris Cluster Manager 브라우저 인터페이스 액세스뿐만 아니라 Oracle Solaris Cluster 명령에 대해 나열, 표시 및 기타 읽기 작업을 수행합니다.
클러스터 작업	이 권한 프로파일은 Oracle Solaris Cluster 소프트웨어에만 해당되며 다음 권한 부여를 포함합니다.	solaris.cluster.read - Oracle Solaris Cluster Manager 브라우저 인터페이스 액세스뿐만 아니라 나열, 표시, 내보내기, 상태 및 기타 읽기 작업을 수행합니다. solaris.cluster.admin - 클러스터 객체의 상태를 변경합니다.
시스템 관리자	이 기존 Oracle Solaris 권한 프로파일에는 클러스터 관리 프로파일에 포함되어 있는 것과 동일한 인증이 포함되어 있습니다.	다른 시스템 관리 작업 외에 클러스터 관리 역할 ID가 수행할 수 있는 동일한 작업을 수행합니다.
클러스터 관리	이 권한 프로파일에는 solaris.cluster.modify 권한 부여뿐 아니라 클러스터 작업 프로파일에 포함되	클러스터 객체의 등록 정보를 변경할뿐만 아니라 클러스터 작업 역할 ID가 실행할 수 있는 작업과 동일한 작업을 수행합니다.

권한 프로파일	권한 부여 및 보안 속성	부여된 권한
	어 있는 것과 동일한 권한 부여가 포함되어 있습니다.	

Oracle Solaris Cluster Management 권한 프로파일을 사용하여 RBAC 역할 만들기 및 지정

이 작업을 사용하면 Oracle Solaris Cluster Management 권한 프로파일을 사용하여 새 RBAC 역할을 만들고 이 새 역할에 사용자를 지정할 수 있습니다.

▼ 명령줄을 사용하여 역할을 만드는 방법

1. 역할을 만들 방법을 선택합니다.

- 로컬 범위의 역할인 경우 `roleadd` 명령을 사용하여 새 로컬 역할과 해당 속성을 지정합니다. 자세한 내용은 [roleadd\(1M\)](#) 매뉴얼 페이지를 참조하십시오.
- 또는 로컬 범위의 역할인 경우 `user_attr` 파일을 편집하여 `type=role`인 사용자를 추가합니다. 자세한 내용은 [user_attr\(4\)](#) 매뉴얼 페이지를 참조하십시오.

이 방법은 긴급 상황인 경우에만 사용합니다.

- 이름 서비스의 역할인 경우 `roleadd` 및 `rolemod` 명령을 사용하여 새 역할과 해당 속성을 지정합니다. 자세한 내용은 [roleadd\(1M\)](#) 및 [rolemod\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

이 명령을 사용하려면 다른 역할을 만들 수 있는 `root` 역할의 인증이 필요합니다. `roleadd` 명령을 모든 이름 서비스에 적용할 수 있습니다. 이 명령은 Solaris Management Console 서버의 클라이언트로 실행됩니다.

2. 이름 서비스 캐시 데몬을 시작 및 중지합니다.

이름 서비스 캐시 데몬을 다시 시작할 때까지 새 역할이 적용되지 않습니다. `root`로 다음 텍스트를 입력합니다.

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

예 13 smrole 명령을 사용하여 사용자정의 운영자 역할 만들기

다음 순서에서는 `smrole` 명령을 사용하여 역할을 만드는 방법을 보여 줍니다. 이 예에서는 표준 운영자 권한 프로파일과 매체 복원 권한 프로파일이 할당된 새로운 버전의 운영자 역할을 만듭니다.

```
% 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 ::      <primaryadmin 암호 입력>

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 ::      <oper2 암호 입력>

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

새로 만든 역할 및 다른 역할을 보려면 다음과 같이 smrole을 list 옵션과 함께 사용합니다.

# /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 ::      <primaryadmin 암호 입력>

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

사용자의 RBAC 등록 정보 수정

사용자 계정 도구나 명령줄을 사용하여 사용자의 RBAC 등록 정보를 수정할 수 있습니다. 사용자의 RBAC 등록 정보를 수정하려면 [명령줄에서 사용자의 RBAC 등록 정보를 수정하는 방법 \[63\]](#)을 참조하십시오. 다음 절차 중 하나를 선택합니다.

- [사용자 계정 도구를 사용하여 사용자의 RBAC 등록 정보를 수정하는 방법 \[62\]](#)
- [명령줄에서 사용자의 RBAC 등록 정보를 수정하는 방법 \[63\]](#)

▼ 사용자 계정 도구를 사용하여 사용자의 RBAC 등록 정보를 수정하는 방법

시작하기 전에 사용자의 등록 정보를 수정하려면 User Tool Collection(사용자 도구 모음)을 root 사용자로 실행하거나 System Administrator 권한 프로파일이 지정된 역할을 보유하고 있어야 합니다.

1. User Accounts(사용자 계정) 도구를 시작합니다.

User Accounts(사용자 계정) 도구를 실행하려면 [Oracle Solaris 11.3의 사용자 및 프로세스 보안의 “지정된 관리 권한 사용”](#)에 설명된 대로 Solaris Management Console을 시작합니다. User Tool Collection(사용자 도구 모음)을 열고 User Accounts(사용자 계정) 아이콘을 누릅니다.

User Accounts(사용자 계정) 도구가 시작되면 기존 사용자 계정에 대한 아이콘이 보기 창에 표시됩니다.

2. 변경할 User Account(사용자 계정) 아이콘을 누르고 Action(작업) 메뉴에서 Properties(등록 정보)를 선택합니다(또는 사용자 계정 아이콘을 두 번 누릅니다).

3. 다음과 같이 변경할 등록 정보에 대해 대화 상자에서 적당한 탭을 누릅니다.

- 사용자에게 지정된 역할을 변경하려면 Roles(롤) 탭을 누르고 변경할 역할 지정을 Available Roles(이용 가능 롤) 또는 Assigned Roles(할당된 롤) 열로 이동합니다.
- 사용자에게 지정된 권한 프로파일을 변경하려면 Rights(권한) 탭을 누르고 Available Rights(사용 가능한 권한) 또는 Granted Rights(부여된 권한) 열로 이동합니다.

주 - 가급적이면 사용자에게 직접 권한 프로파일을 지정하지 마십시오. 권한이 부여된 응용 프로그램을 수행하기 위해 사용자가 역할을 가지도록 하는 것이 더 좋은 방법입니다. 이 전력은 사용자가 권한을 남용하지 않도록 합니다.

▼ 명령줄에서 사용자의 RBAC 등록 정보를 수정하는 방법

1. `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 해당하는 명령을 선택합니다.

- 로컬 범위에 또는 LDAP 저장소에 정의된 사용자에게 지정된 사용자 등록 정보를 변경하려면 `usermod` 명령을 사용합니다. 자세한 내용은 [usermod\(1M\)](#) 매뉴얼 페이지를 참조하십시오.
- 또는 로컬 범위에 정의된 사용자에게 지정된 인증, 역할 또는 권한 프로파일을 변경하려면 `user_attr` 파일을 편집합니다.
이 방법은 긴급 상황인 경우에만 사용합니다.
- 로컬로 또는 LDAP 저장소와 같은 이름 서비스에서 역할을 관리하려면 `roleadd` 또는 `rolemod` 명령을 사용합니다. 자세한 내용은 [roleadd\(1M\)](#) 또는 [rolemod\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

이러한 명령을 사용하려면 사용자 파일을 변경할 수 있는 `root` 역할과 같은 인증이 필요합니다. 이러한 명령을 모든 이름 서비스에 적용할 수 있습니다. [Oracle Solaris 11.3의 사용자 계정 및 사용자 환경 관리의 “사용자, 역할 및 그룹을 관리하는 데 사용되는 명령”](#)을 참조하십시오.

Oracle Solaris 11과 함께 제공된 Forced Privilege 및 Stop 권한 프로파일은 수정할 수 없습니다.

클러스터 종료 및 부트

이 장에서는 클러스터와 전역 클러스터, 영역 클러스터 및 노드를 종료하고 부트하는 절차 및 정보를 제공합니다.

- “클러스터 종료 및 부트 개요” [65]
- “클러스터의 단일 노드 종료 및 부트” [79]
- “꼭 찬 /var 파일 시스템 복구” [92]

이 장에 있는 관련 절차에 대한 자세한 내용은 [비클러스터 모드로 노드를 부트하는 방법](#) [90] 및 [표 4. “작업 맵: 노드 종료 및 부트”](#)를 참조하십시오.

클러스터 종료 및 부트 개요

Oracle Solaris Cluster `cluster shutdown` 명령은 전역 클러스터 서비스를 순차적으로 중지하고 전체 전역 클러스터를 정상적으로 종료합니다. 전역 클러스터의 위치를 이동할 때 또는 응용 프로그램 오류로 인해 데이터 손상이 발생하는 경우 전역 클러스터를 종료하기 위해 `cluster shutdown` 명령을 사용할 수 있습니다. `clzonecluster halt` 명령은 특정 노드에서 실행되는 영역 클러스터를 중지하거나 구성된 모든 노드에서 전체 영역 클러스터를 중지합니다. 영역 클러스터 내에서 `cluster shutdown` 명령을 사용할 수도 있습니다. 자세한 내용은 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

이 장의 절차에서 `phys-schost#`는 전역 클러스터 프롬프트를 반영합니다. `clzonecluster` 대화식 셸 프롬프트는 `clzc:schost>`입니다.

주 - 전체 전역 클러스터를 올바르게 종료하려면 `cluster shutdown` 명령을 사용합니다. Oracle Solaris shutdown 명령은 `clnode evacuate` 명령과 함께 사용되어 개별 노드를 종료합니다. 자세한 내용은 [클러스터를 종료하는 방법](#) [66], “클러스터의 단일 노드 종료 및 부트” [79] 또는 [clnode\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 노드를 비울 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법](#) [278]을 참조하십시오.

`cluster shutdown` 및 `clzonecluster halt` 명령은 다음 작업을 수행하여 각각 전역 클러스터 또는 영역 클러스터의 모든 노드를 중지합니다.

1. 실행 중인 모든 리소스 그룹을 오프라인으로 전환합니다.
2. 전역 클러스터 또는 영역 클러스터의 모든 클러스터 파일 시스템을 마운트 해제합니다.
3. `cluster shutdown` 명령은 전역 클러스터 또는 영역 클러스터의 활성 장치 서비스를 종료합니다.
4. `cluster shutdown` 명령은 `init 0`을 실행하고 클러스터의 모든 노드를 SPARC 기반 시스템의 OpenBoot PROM ok 프롬프트에 표시하거나 Press any key to continue 메시지를 x86 기반 시스템의 GRUB 메뉴에 표시합니다. GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”](#)를 참조하십시오. `clzonecluster halt` 명령은 `zoneadm -z zone-cluster-name halt` 명령을 수행하여 영역 클러스터의 영역을 중지합니다(종료하지는 않음).

주 - 필요하다면 노드가 클러스터 구성원에 포함되지 않도록 비클러스터 모드에서 노드를 부트할 수 있습니다. 클러스터 소프트웨어를 설치하거나 특정 관리 절차를 수행할 경우에는 비클러스터 모드가 유용합니다. 자세한 내용은 [비클러스터 모드로 노드를 부트하는 방법 \[90\]](#)을 참조하십시오.

표 3 작업 목록: 클러스터 종료 및 부트

작업	지침
클러스터를 중지합니다.	클러스터를 종료하는 방법 [66]
모든 노드를 부트하여 클러스터를 시작합니다. 클러스터 멤버십을 얻으려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.	클러스터를 부트하는 방법 [69]
클러스터를 재부트합니다.	클러스터를 재부트하는 방법 [73]

▼ 클러스터를 종료하는 방법

한 개의 전역 클러스터 또는 영역 클러스터나 모든 영역 클러스터를 종료할 수 있습니다.



주의 - 클러스터 콘솔에서 `send brk` 명령을 사용하여 전역 클러스터 노드 또는 영역 클러스터 노드를 종료하지 마십시오. 클러스터에서는 이 명령을 사용할 수 없습니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. (x86만 해당) 전역 클러스터 또는 영역 클러스터에서 Oracle RAC(Real Application Clusters)를 실행하는 경우 종료 중인 클러스터에서 모든 데이터베이스 인스턴스를 종료합니다.

종료 절차에 대한 내용은 Oracle RAC 제품 설명서를 참조하십시오.

2. 클러스터의 임의 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

3. 해당 전역 클러스터 또는 영역 클러스터나 모든 영역 클러스터를 종료합니다.

- 전역 클러스터를 종료합니다. 이 작업은 영역 클러스터도 모두 종료합니다.

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

- 특정 영역 클러스터를 종료합니다.

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

- 모든 영역 클러스터를 종료합니다.

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

영역 클러스터 내에서 `cluster shutdown` 명령을 사용하여 특정 영역 클러스터를 종료할 수도 있습니다.

4. 전역 클러스터 또는 영역 클러스터의 모든 노드에 `ok` 프롬프트(SPARC 기반 시스템) 또는 GRUB 메뉴(x86 기반 시스템)가 표시되는지 확인합니다.

모든 노드가 `ok` 프롬프트(SPARC 기반 시스템) 또는 부트 서브시스템(x86 기반 시스템)에 있을 때까지 노드의 전원을 끄지 마십시오.

- 클러스터에서 아직 가동하여 실행되고 있는 다른 전역 클러스터 노드에서 하나 이상의 전역 클러스터 노드 상태를 확인합니다.

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

- `status` 하위 명령을 사용하여 영역 클러스터가 종료되었는지 확인합니다.

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

5. 필요한 경우 전역 클러스터의 노드 전원을 끕니다.

예 14 영역 클러스터 종료

다음 예에서는 `sczone`이라는 영역 클러스터를 종료합니다.

```
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#
```

예 15 SPARC: 전역 클러스터 종료

다음은 정상적인 전역 클러스터 작업이 중지되고 모든 노드가 종료되어 ok 프롬프트가 표시되는 콘솔 출력의 예입니다. -g 0 옵션은 종료 유예 시간을 0으로 설정하고, -y 옵션은 확인 질문에 자동으로 yes 응답을 제공합니다. 전역 클러스터에 있는 다른 노드의 콘솔에도 종료 메시지가 나타납니다.

```
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
```

예 16 x86: 전역 클러스터 종료

다음은 정상적인 전역 클러스터 작업이 중지되고 모든 노드가 종료될 때 표시되는 콘솔 출력의 예입니다. 이 예에서는 모든 노드에 ok 프롬프트가 표시되지는 않습니다. -g 0 옵션은 종료 유예 시간을 0으로 설정하고, -y 옵션은 확인 질문에 자동으로 yes 응답을 제공합니다. 전역 클러스터에 있는 다른 노드의 콘솔에도 종료 메시지가 나타납니다.

```
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
```

참조 종료된 전역 클러스터 또는 영역 클러스터를 재시작하려면 [클러스터를 부트하는 방법 \[69\]](#)을 참조하십시오.

▼ 클러스터를 부트하는 방법

이 절차에서는 노드가 종료된 전역 클러스터 또는 영역 클러스터를 시작하는 방법에 대해 설명합니다. 전역 클러스터 노드의 경우 ok 프롬프트(SPARC 시스템) 또는 Press any key to continue 메시지(GRUB 기반 x86 시스템)가 표시됩니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - 영역 클러스터를 만들려면 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “[영역 클러스터 만들기 및 구성](#)”의 지침을 따르거나 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용합니다.

1. 각 노드를 클러스터 모드로 부트하십시오.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

■ SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot
```

■ x86 기반 시스템에서는 다음 명령을 실행합니다.

GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다. GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료](#)의 “[시스템 부트](#)”를 참조하십시오.

주 - 클러스터 구성원이 되려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.

■ 한 개의 영역 클러스터가 있는 경우 전체 영역 클러스터를 부트할 수 있습니다.

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

■ 영역 클러스터가 두 개 이상 있는 경우 모든 영역 클러스터를 부트할 수 있습니다.

zone-cluster-name 대신 더하기 기호(+)를 사용합니다.

2. 노드가 오류 없이 부트되고 온라인 상태인지 확인합니다.

cluster status 명령은 전역 클러스터 노드의 상태를 보고합니다.

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

전역 클러스터 노드에서 clzonecluster status 상태 명령을 실행하면 영역 클러스터 노드의 상태가 보고됩니다.

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

주 - 노드의 /var 파일 시스템이 꽉 차면 해당 노드에서 Oracle Solaris Cluster를 다시 시작하지 못할 수도 있습니다. 이런 문제가 발생하면 [꼭 찬 /var 파일 시스템을 복구하는 방법 \[92\]](#)을 참조하십시오. 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

예 17 SPARC: 전역 클러스터 부트

다음 예에서는 phys-schost-1 노드를 전역 클러스터로 부트할 때 표시되는 콘솔 출력을 보여 줍니다. 전역 클러스터에 있는 다른 노드의 콘솔에 유사한 메시지가 나타납니다. 영역 클러스터의 autoboot 등록 정보가 true로 설정된 경우 영역 클러스터 노드는 해당 시스템에서 전역 클러스터 노드가 부트된 후 자동으로 부트됩니다.

전역 클러스터 노드가 재부트되면 해당 시스템의 모든 영역 클러스터 노드가 중지됩니다. autoboot 등록 정보가 true로 설정된 동일한 시스템의 모든 영역 클러스터 노드는 전역 클러스터 노드가 다시 시작된 후 부트됩니다.

```
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.
...
```

예 18 x86: 클러스터 부트

다음 예에서는 phys-schost-1 노드를 클러스터로 부트할 때 표시되는 콘솔 출력을 보여 줍니다. 클러스터에 있는 다른 노드의 콘솔에 유사한 메시지가 나타납니다.

```
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.

```

▼ 클러스터를 재부트하는 방법

전역 클러스터를 종료하려면 `cluster shutdown` 명령을 실행한 다음 각 노드에서 `boot` 명령을 사용하여 전역 클러스터를 부트합니다. 영역 클러스터를 종료하려면 `clzonecluster halt` 명령을 사용한 다음 `clzonecluster boot` 명령을 사용하여 영역 클러스터를 부트합니다. `clzonecluster reboot` 명령을 사용할 수도 있습니다. 자세한 내용은 [cluster\(1CL\)](#), [boot\(1M\)](#), [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터에서 Oracle RAC를 실행하는 경우 종료 중인 클러스터에서 모든 데이터베이스 인스턴스를 종료합니다.

종료 절차에 대한 내용은 Oracle RAC 제품 설명서를 참조하십시오.

2. 클러스터의 임의 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

3. 클러스터를 종료합니다.

■ 전역 클러스터를 종료합니다.

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

■ 영역 클러스터가 있는 경우 전역 클러스터 노드에서 영역 클러스터를 종료합니다.

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

각 노드가 종료됩니다. 영역 클러스터 내에서 `cluster shutdown` 명령을 사용하여 영역 클러스터를 종료할 수도 있습니다.

주 - 클러스터 구성원이 되려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.

4. 각 노드를 부트합니다.

시스템 종료 사이에 구성을 변경하지 않으면 노드의 부트 순서는 중요하지 않습니다. 종료 중간에 구성을 변경하려면 먼저 최근 구성을 사용하여 노드를 시작하십시오.

■ SPARC 기반 시스템의 전역 클러스터 노드인 경우 다음 명령을 실행합니다.

```
ok boot
```

■ x86 기반 시스템의 전역 클러스터 노드인 경우 다음 명령을 실행합니다.

GRUB 메뉴가 표시되면 적절한 Oracle Solaris OS 항목을 선택하고 Enter 키를 누릅니다.

GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”](#)를 참조하십시오.

- 영역 클러스터의 경우 전역 클러스터의 단일 노드에서 다음 명령을 입력하여 영역 클러스터를 부트합니다.

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

주 - 클러스터 구성원이 되려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.

클러스터 구성요소가 활성화되면 부트된 노드의 콘솔에 메시지가 나타납니다.

5. 노드가 오류 없이 부트되고 온라인 상태인지 확인합니다.

- **clnode status** 명령은 전역 클러스터의 노드 상태를 보고합니다.

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

- 전역 클러스터 노드에서 **clzonecluster status** 명령을 실행하면 영역 클러스터 노드의 상태가 보고됩니다.

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

영역 클러스터 내에서 **cluster status** 명령을 실행하여 노드의 상태를 표시할 수도 있습니다.

주 - 노드의 /var 파일 시스템이 꽉 차면 해당 노드에서 Oracle Solaris Cluster를 다시 시작하지 못할 수도 있습니다. 이런 문제가 발생하면 [꼭 찬 /var 파일 시스템을 복구하는 방법 \[92\]](#)을 참조하십시오.

예 19 영역 클러스터 재부트

다음 예에서는 *sparse-sczzone*이라는 영역 클러스터를 정지하고 부트하는 방법을 보여 줍니다. **clzonecluster reboot** 명령을 사용할 수도 있습니다.

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

```

phys-schost#
phys-schost# clzonecluster boot sparse-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#

```

예 20 SPARC: 전역 클러스터 재부트

다음 예에서는 정상적인 전역 클러스터 작업이 중지되고 모든 노드가 종료되어 ok 프롬프트가 표시된 다음 전역 클러스터가 다시 시작될 때 표시되는 콘솔 출력을 보여 줍니다. -g 0 옵션은 유예 기간을 0으로 설정하고, -y 옵션은 확인 질문에 자동으로 yes 응답을 제공합니다. 전역 클러스터에 있는 다른 노드의 콘솔에도 종료 메시지가 나타납니다.

```

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:

```

예 21 x86: 클러스터 재부트

다음 예에서는 정상적인 클러스터 작업이 중지되고 모든 노드가 종료되어 클러스터가 다시 시작될 때 표시되는 콘솔 출력을 보여 줍니다. `-g 0` 옵션은 유예 기간을 0으로 설정하고, `-y`는 확인 질문에 자동으로 yes 응답을 제공합니다. 클러스터에 있는 다른 노드의 콘솔에도 종료 메시지가 나타납니다.

```

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

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

```

```

AMIBIOS (C)1985-2002 American Megatrends Inc.,
Copyright 1996-2002 Intel Corporation
SCB20.86B.1064.P18.0208191106
SCB2 Production BIOS Version 2.08
BIOS Build 1064
2 X Intel(R) Pentium(R) III CPU family      1400MHz
Testing system memory, memory size=2048MB
2048MB Extended Memory Passed
512K L2 Cache SRAM Passed
ATAPI CD-ROM SAMSUNG CD-ROM SN-124

```

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

```

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

```

```

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

```

```

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

```

```

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

```

SunOS - Intel Platform Edition Primary Boot Subsystem, vsn 2.0

Current Disk Partition Information

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

Please select the partition you wish to boot: * *

Solaris DCB

loading /solaris/boot.bin

SunOS Secondary Boot version 3.00

Solaris Intel Platform Edition Booting System

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

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

Press ESCape to interrupt autoboot in 2 seconds.

Initializing system

Please wait...

Warning: Resource Conflict - both devices are added

NON-ACPI device: ISY0050

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

ACPI device: ISY0050

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

<<< Current Boot Parameters >>>

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

Boot args:

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

<<< timeout in 5 seconds >>>

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

Size: 275683 + 22092 + 150244 Bytes

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

SunOS Release 5.9 Version Generic_112234-07 32-bit

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

Use is subject to license terms.

configuring IPv4 interfaces: e1000g2.

Hostname: phys-schost-1

Booting as part of a cluster

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

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

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

NOTICE: clcomm: Adapter e1000g3 constructed

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

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

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

NOTICE: clcomm: Adapter e1000g0 constructed

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

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

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

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

NOTICE: CMM: Cluster has reached quorum.

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

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

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

NOTICE: CMM: node reconfiguration #1 completed.

```

NOTICE: CMM: Node phys-schost-1: joined cluster.
WARNING: mod_installdrv: no major number for rsmrdt
ip: joining multicasts failed (18) on clprivnet0 - will use link layer
broadcasts for multicast
The system is coming up. Please wait.
checking ufs filesystems
/dev/rdisk/clt0d0s5: is clean.
NOTICE: clcomm: Path phys-schost-1:e1000g0 - phys-schost-2:e1000g0 online
NIS domain name is dev.eng.mycompany.com
starting rpc services: rpcbind keyserver ypbind done.
Setting netmask of e1000g2 to 192.168.255.0
Setting netmask of e1000g3 to 192.168.255.128
Setting netmask of e1000g0 to 192.168.255.128
Setting netmask of clprivnet0 to 192.168.255.0
Setting default IPv4 interface for multicast: add net 224.0/4: gateway phys-schost-1
syslog service starting.
obtaining access to all attached disks

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

phys-schost-1 console login:

```

클러스터의 단일 노드 종료 및 부트

전역 클러스터 노드 또는 영역 클러스터 노드를 종료할 수 있습니다. 이 절에서는 전역 클러스터 노드 및 영역 클러스터 노드를 종료하는 지침에 대해 설명합니다.

전역 클러스터 노드를 종료하려면 Oracle Solaris shutdown 명령과 함께 clnode evacuate 명령을 사용합니다. 전체 전역 클러스터를 종료하는 경우에만 cluster shutdown 명령을 사용합니다.

영역 클러스터 노드의 경우 전역 클러스터에서 clzonecluster halt 명령을 사용하여 단일 영역 클러스터 노드 또는 전체 영역 클러스터를 종료합니다. clnode evacuate 및 shutdown 명령을 사용하여 영역 클러스터 노드를 종료할 수도 있습니다.

자세한 내용은 [clnode\(1CL\)](#), [shutdown\(1M\)](#), [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

이 장의 절차에서 phys-schost#는 전역 클러스터 프롬프트를 반영합니다. clzonecluster 대화식 셸 프롬프트는 clzc:schost>입니다.

표 4 작업 맵: 노드 종료 및 부트

작업	도구	지침
노드 중지	전역 클러스터 노드의 경우 clnode evacuate 및 shutdown 명령을 사용합니다. 영역 클러스터 노드의 경우 clzonecluster halt 명령을 사용합니다.	노드 종료 방법 [80]
노드 시작 클러스터 멤버십을 얻으려면 클러스터 노드에 클러스터 상호 연결에 대하여 작동하는 연결이 있어야 합니다.	전역 클러스터 노드의 경우 boot 또는 b 명령을 사용합니다. 영역 클러스터 노드의 경우 clzonecluster boot 명령을 사용합니다.	노드 부트 방법 [83]
클러스터에서 한 노드를 중지하고 재시작(재부트) 클러스터 멤버십을 얻으려면 클러스터 노드에 클러스터 상호 연결에 대하여 작동하는 연결이 있어야 합니다.	전역 클러스터 노드의 경우 clnode evacuate 및 shutdown 명령을 사용한 다음 boot 또는 b를 사용합니다. 영역 클러스터 노드의 경우 clzonecluster reboot 명령을 사용합니다.	노드 재부트 방법 [87]
노드가 클러스터 구성원에 포함되지 않도록 노드 부트	전역 클러스터 노드의 경우 clnode evacuate 및 shutdown 명령을 사용한 다음 boot -x (SPARC) 또는 GRUB 메뉴 항목 편집(x86)을 사용합니다. 기본 전역 클러스터가 비클러스터 모드로 부트된 경우 영역 클러스터 노드는 자동으로 비클러스터 모드가 됩니다.	비클러스터 모드로 노드를 부트하는 방법 [90]

▼ 노드 종료 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.



주의 - 클러스터 콘솔에서 send brk 명령을 사용하여 전역 클러스터 또는 영역 클러스터의 노드를 종료하지 마십시오. 클러스터에서는 이 명령을 사용할 수 없습니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 전역 클러스터 노드를 비우고 모든 리소스 그룹 및 장치 그룹을 다음 선호 노드로 전환할 수도 있습니다. 또한 영역 클러스터 노드를 종료할 수 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터에서 Oracle RAC를 실행하는 경우 종료 중인 클러스터에서 모든 데이터베이스 인스턴스를 종료합니다.

종료 절차에 대한 내용은 Oracle RAC 제품 설명서를 참조하십시오.

2. 종료할 클러스터 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

3. 특정 영역 클러스터 구성원을 중지하려면 4 - 6단계를 건너뛰고 전역 클러스터 노드에서 다음 명령을 실행합니다.

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

특정 영역 클러스터 노드를 지정하면 해당 노드만 중지됩니다. 기본적으로 `halt` 명령은 모든 노드의 영역 클러스터를 중지합니다.

4. 종료할 노드의 모든 리소스 그룹, 리소스 및 장치 그룹을 다른 전역 클러스터 구성원으로 전환합니다.

종료할 전역 클러스터 노드에서 다음 명령을 입력합니다. `clnode evacuate` 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 리소스 그룹 및 장치 그룹을 전환합니다. 영역 클러스터 노드 내에서 `clnode evacuate`를 실행할 수도 있습니다.

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

`node` 전환할 리소스 그룹 및 장치 그룹이 있는 노드를 지정합니다.

5. 노드를 종료합니다.

종료할 전역 클러스터 노드에서 `shutdown` 명령을 실행합니다.

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

전역 클러스터 노드에 `ok` 프롬프트(SPARC 기반 시스템) 또는 `Press any key to continue` 메시지(x86 기반 시스템의 GRUB 메뉴)가 표시되는지 확인합니다.

6. 필요한 경우 노드의 전원을 끕니다.

예 22 SPARC: 전역 클러스터 노드 종료

다음 예에서는 `phys-schost-1` 노드가 종료될 때 표시되는 콘솔 출력을 보여 줍니다. `-g0` 옵션은 유예 시간을 0으로 설정하고, `-y` 옵션은 확인 질문에 자동으로 `yes` 응답을 제공합니다. 이 노드의 종료 메시지가 전역 클러스터에 있는 다른 노드의 콘솔에 나타납니다.

```
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
```

예 23 x86: 전역 클러스터 노드 종료

다음 예에서는 phys-schost-1 노드가 종료될 때 표시되는 콘솔 출력을 보여 줍니다. -g0 옵션은 유예 기간을 0으로 설정하고, -y 옵션은 확인 질문에 자동으로 yes 응답을 제공합니다. 이 노드의 종료 메시지가 전역 클러스터에 있는 다른 노드의 콘솔에 나타납니다.

```
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
```

예 24 영역 클러스터 노드 종료

다음 예에서는 clzonecluster halt를 사용하여 *sparse-sczone*이라는 영역 클러스터의 노드를 종료하는 방법을 보여 줍니다. 영역 클러스터 노드에서 clnode evacuate 및 shutdown 명령을 실행할 수도 있습니다.

```
phys-schost# clzonecluster status

=== Zone Clusters ===

--- Zone Cluster Status ---

Name          Node Name    Zone HostName  Status    Zone Status
----

```

```

sparse-sczone  schost-1  sczone-1  Online  Running
                schost-2  sczone-2  Online  Running
                schost-3  sczone-3  Online  Running
                schost-4  sczone-4  Online  Running

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

=== Zone Clusters ===

--- Zone Cluster Status ---

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

phys-schost#

```

참조 종료된 전역 클러스터 노드를 재시작하려면 [노드 부트 방법 \[83\]](#)을 참조하십시오.

▼ 노드 부트 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터 노드를 부트할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

전역 클러스터 또는 영역 클러스터에서 다른 활성 노드를 종료하거나 재부트하려면 부트하려는 노드에 대해 다중 사용자 서버 이정표가 온라인 상태가 될 때까지 기다리십시오. 그렇지 않으면 종료하거나 재부트하는 클러스터의 다른 노드에서 서비스를 인계할 수 없습니다.

주 - 쿼럼 구성에 따라 노드 시작 방법이 달라질 수 있습니다. 두 개의 노드로 구성된 클러스터에서는 클러스터의 총 쿼럼 계수가 3이 되도록 쿼럼 장치가 구성되어야 합니다. 즉, 각 노드에 대한 쿼럼 수가 하나씩 구성되고 쿼럼 장치에 대한 쿼럼 수 하나가 구성되어야 합니다. 이러한 경우에 첫번째 노드가 종료되면 두번째 노드가 계속 쿼럼 자격을 갖고 단일 클러스터 구성원으로 실행됩니다. 첫번째 노드가 다시 클러스터에 포함되어 클러스터 노드로 실행되려면 두번째 노드가 계속 실행되고 있어야 합니다. 또한 필요한 쿼럼 수(2)가 유지되어야 합니다.

게스트 도메인에서 Oracle Solaris Cluster를 실행 중인 경우 컨트롤 또는 I/O 도메인을 재부트하면 도메인 작동 중지를 포함하여 실행 중인 게스트 도메인에 영향을 미칠 수 있습니다. 컨트롤 또는 I/O 도메인을 재부트하기 전에 다른 노드로 작업 부하 균형을 조정하고 Oracle Solaris Cluster를 실행 중인 게스트 도메인을 중지해야 합니다.

컨트롤 또는 I/O 도메인이 재부트되면 게스트 도메인에서 하트비트를 받거나 보내지 않습니다. 이로 인해 정보 분리(split-brain) 및 클러스터 재구성이 발생합니다. 컨트롤 또는 I/O 도메인이 재부트된 후에는 게스트 도메인이 공유 장치에 액세스할 수 없습니다. 다른 클러스터 노드가 공유 장치로부터 이 게스트 도메인을 보호합니다. 컨트롤 또는 I/O 도메인이 재부트를 마치면 클러스터 재구성의 일부로 공유 디스크에서 보호가 해제되므로 게스트 도메인에서 I/O가 재개되고 공유 장치의 I/O로 인해 게스트 도메인 패닉이 발생합니다. 게스트가 중복을 위해 두 개의 I/O 도메인을 사용 중이고 한 번에 하나씩 I/O 도메인을 재부트할 경우 이 문제를 해결할 수 있습니다.

주 - 클러스터 구성원이 되려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.

1. 종료된 전역 클러스터 노드 또는 영역 클러스터 노드를 시작하려면 노드를 부트합니다.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot
```

- x86 기반 시스템에서는 다음 명령을 실행합니다.

GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다.

클러스터 구성요소가 활성화되면 부트된 노드의 콘솔에 메시지가 나타납니다.

- 한 개의 영역 클러스터가 있는 경우 부트할 노드를 지정할 수 있습니다.

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

2. 노드가 오류 없이 부트되고 온라인 상태인지 확인합니다.

- **cluster status** 명령을 실행하면 전역 클러스터 노드의 상태가 보고됩니다.

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

- 전역 클러스터의 한 노드에서 `clzonecluster status` 명령을 실행하면 모든 영역 클러스터 노드의 상태가 보고됩니다.

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

영역 클러스터 노드는 해당 노드를 호스트하는 노드가 클러스터 모드로 부트된 경우에만 클러스터 모드로 부트할 수 있습니다.

주 - 노드의 `/var` 파일 시스템이 꽉 차면 해당 노드에서 Oracle Solaris Cluster를 다시 시작하지 못할 수도 있습니다. 이런 문제가 발생하면 [꼭 찬 /var 파일 시스템을 복구하는 방법 \[92\]](#)을 참조하십시오.

예 25 SPARC: 전역 클러스터 노드 부트

다음 예에서는 `phys-schost-1` 노드를 전역 클러스터로 부트할 때 표시되는 콘솔 출력을 보여 줍니다.

```
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:
```

예 26 x86: 클러스터 노드 부트

다음 예에서는 `phys-schost-1` 노드를 클러스터로 부트할 때 표시되는 콘솔 출력을 보여 줍니다.

```
<<< 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 keyerv 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:
```

▼ 노드 재부트 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터 노드를 재부트할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.



주의 - 리소스에 대한 방법의 시간이 초과하여 종료할 수 없으면 리소스의 Failover_mode 등록 정보가 HARD로 설정된 경우에만 노드가 재부트됩니다. Failover_mode 등록 정보가 다른 값으로 설정된 경우에는 노드가 재부트되지 않습니다.

전역 클러스터 또는 영역 클러스터에서 다른 활성 노드를 종료하거나 재부트하려면 재부트하려는 노드에 대해 다중 사용자 서버 이정표가 온라인 상태가 될 때까지 기다리십시오. 그렇지 않으면 종료하거나 재부트하는 클러스터의 다른 노드에서 서비스를 인계할 수 없습니다.

1. 전역 클러스터 또는 영역 클러스터 노드에서 Oracle RAC를 실행 중인 경우에는 종료하려는 노드에서 데이터베이스 인스턴스를 모두 종료합니다.

종료 절차에 대한 내용은 Oracle RAC 제품 설명서를 참조하십시오.

2. 종료할 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다. 전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

3. `clnode evacuate` 및 `shutdown` 명령을 사용하여 전역 클러스터 노드를 종료합니다.

전역 클러스터의 한 노드에서 `clzonecluster halt` 명령을 실행하여 영역 클러스터를 종료합니다. (`clnode evacuate` 및 `shutdown` 명령도 영역 클러스터에서 작동합니다.)

전역 클러스터의 경우 종료할 노드에서 다음 명령을 입력합니다. `clnode evacuate` 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 장치 그룹을 전환합니다. 또한 이 명령은 지정된

노드의 전역 영역에서 다른 노드의 다음 우선 순위 전역 영역으로 모든 리소스 그룹을 전환합니다.

주 - 단일 노드를 종료하려면 `shutdown -g0 -y -i6` 명령을 사용합니다. 동시에 여러 노드를 종료하려면 `shutdown -g0 -y -i0` 명령을 사용하여 노드를 정지합니다. 모든 노드가 정지된 후 모든 노드에서 `boot` 명령을 사용하여 해당 노드를 클러스터로 다시 부트합니다.

- SPARC 기반 시스템에서는 단일 노드를 재부트하려면 다음 명령을 실행합니다.

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

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

- x86 기반 시스템에서는 단일 노드를 재부트하려면 다음 명령을 실행합니다.

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

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

GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다.

- 종료하고 재부트할 영역 클러스터 노드를 지정합니다.

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

주 - 클러스터 구성원이 되려면 노드가 현재 클러스터 상호 연결에 연결되어 있어야 합니다.

4. 노드가 오류 없이 부트되고 온라인 상태인지 확인합니다.

- 전역 클러스터 노드가 온라인 상태인지 확인합니다.

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

- 영역 클러스터 노드가 온라인 상태인지 확인합니다.

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

예 27 SPARC: 전역 클러스터 노드 재부트

다음 예에서는 `phys-schost-1` 노드가 재부트될 때 표시되는 콘솔 출력을 보여 줍니다. 종료 및 시작 알림과 같은 이 노드의 메시지가 전역 클러스터에 있는 다른 노드의 콘솔에 나타납니다.

```
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:

```

예 28 영역 클러스터 노드 재부트

다음 예에서는 영역 클러스터의 한 노드를 재부트하는 방법을 보여 줍니다.

```

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#
```

▼ 비클러스터 모드로 노드를 부트하는 방법

노드가 클러스터 멤버십에 포함되지 않는 비클러스터 모드로 전역 클러스터 노드를 부트할 수 있습니다. 비클러스터 모드는 클러스터 소프트웨어를 설치하거나 노드 업데이트와 같은 특정 관리 절차를 수행하는 경우에 유용합니다. 영역 클러스터 노드는 기본 전역 클러스터 노드의 상태와 다른 부트 상태가 될 수 없습니다. 전역 클러스터 노드가 비클러스터 모드로 부트된 경우 영역 클러스터 노드는 자동으로 비클러스터 모드가 됩니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 비클러스터 모드로 시작하려면 클러스터에서 **solaris.cluster.admin** RBAC 권한 부여를 제공하는 역할로 전환합니다.

전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.

2. 영역 클러스터 노드 또는 전역 클러스터 노드를 종료합니다.

clnode evacuate 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 장치 그룹을 전환합니다. 또한 이 명령은 지정된 노드의 전역 영역에서 다른 노드의 다음 우선 순위 전역 영역으로 모든 리소스 그룹을 전환합니다.

■ 특정 전역 클러스터 노드를 종료합니다.

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

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

■ 전역 클러스터 노드에서 특정 영역 클러스터 노드를 종료합니다.

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

영역 클러스터 내에서 **clnode evacuate** 및 **shutdown** 명령을 사용할 수도 있습니다.

3. 전역 클러스터 노드에 **ok** 프롬프트(Oracle Solaris 기반 시스템) 또는 **Press any key to continue** 메시지(x86 기반 시스템의 GRUB 메뉴)가 표시되는지 확인합니다.

4. 비클러스터 모드로 전역 클러스터 노드를 부트합니다.

■ SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot -xs
```

■ x86 기반 시스템에서는 다음 명령을 실행합니다.

- a. GRUB 메뉴에서 화살표 키를 사용하여 적절한 Oracle Solaris 항목을 선택하고 **e**를 입력하여 해당 명령을 편집합니다.

GRUB 메뉴가 나타납니다.

GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”](#)를 참조하십시오.

- b. 부트 매개 변수 화면에서 화살표 키를 사용하여 커널 항목을 선택하고 **e**를 입력하여 항목을 편집합니다.

GRUB 부트 매개 변수 화면이 나타납니다.

- c. 명령에 **-x**를 추가하여 시스템 부트를 비클러스터 모드로 지정합니다.

```
[ 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. Enter 키를 눌러 변경사항을 수락하고 부트 매개 변수 화면으로 복귀합니다.

화면에 편집된 명령이 표시됩니다.

- e. **b**를 입력하여 비클러스터 모드로 노드를 부트합니다.

주 - 커널 부트 매개 변수 명령에 대한 변경사항은 시스템을 재부트하면 사라집니다. 다음에 노드를 재부트하면 클러스터 모드로 부트됩니다. 대신 비클러스터 모드로 부트하려면, 이러한 단계를 다시 실행하여 **-x** 옵션을 커널 부트 매개 변수 명령에 추가합니다.

예 29 SPARC: 비클러스터 모드로 전역 클러스터 노드 부트

다음 예에서는 phys-schost-1 노드가 종료되고 비클러스터 모드로 다시 시작될 때 표시되는 콘솔 출력을 보여 줍니다. **-g0** 옵션은 유예 기간을 0으로 설정하고, **-y** 옵션은 확인 질문에 자동으로 yes 응답을 제공하며, **-i0** 옵션은 실행 레벨 0을 호출합니다. 이 노드의 종료 메시지가 전역 클러스터에 있는 다른 노드의 콘솔에 나타납니다.

```
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:
```

꽉 찬 /var 파일 시스템 복구

Oracle Solaris 소프트웨어와 Oracle Solaris Cluster 소프트웨어는 모두 /var/adm/messages 파일에 오류 메시지를 쓰므로 시간이 지나면 /var 파일 시스템이 꽉 찰 수 있습니다. 클러스터 노드의 /var 파일 시스템이 꽉 차면 Oracle Solaris Cluster가 다음 부트 시 해당 노드에서 시작하지 못할 수도 있습니다. 또한 노드에 로그인하지 못할 수도 있습니다.

▼ 꽉 찬 /var 파일 시스템을 복구하는 방법

노드에서 /var 파일 시스템이 꽉 찼다고 보고하고 계속 Oracle Solaris Cluster 서비스를 실행하는 경우 꽉 찬 파일 시스템을 지우려면 이 절차를 사용합니다. 자세한 내용은 [Oracle Solaris 11.3의 시스템 관리 문제 해결](#)의 “시스템 메시지 형식”을 참조하십시오.

1. 꽉 찬 /var 파일 시스템이 있는 클러스터 노드에서 root 역할로 전환합니다.
2. 꽉 찬 파일 시스템을 지웁니다.
예를 들어, 파일 시스템에서 반드시 필요한 파일이 아니면 삭제하십시오.

데이터 복제 접근 방식

이 장에서는 Oracle Solaris Cluster 소프트웨어에서 사용할 수 있는 데이터 복제 기술에 대해 설명합니다. 데이터 복제는 기본 저장 장치에서 백업 또는 보조 장치로 데이터를 복사하는 것으로 정의됩니다. 기본 장치가 실패할 경우, 보조 장치의 데이터를 사용할 수 있습니다. 데이터 복제를 사용하면 클러스터에 대해 고가용성 및 재해 허용 한계를 보장할 수 있습니다.

Oracle Solaris Cluster 소프트웨어에서는 다음과 같은 유형의 데이터 복제를 지원합니다.

- 클러스터 간 - 재해 복구에 Oracle Solaris Cluster Geographic Edition을 사용합니다.
- 클러스터 내 - 캠퍼스 클러스터 내에서 호스트 기반 미러링 대신 사용합니다.

데이터 복제를 수행하려면 복제할 객체와 이름이 같은 장치 그룹이 있어야 합니다. 장치는 한 번에 하나의 장치 그룹에만 속할 수 있으므로 장치가 포함된 Oracle Solaris Cluster 장치 그룹이 이미 있는 경우에는 새 장치 그룹에 해당 장치를 추가하기 전에 그룹을 삭제해야 합니다. Solaris Volume Manager, ZFS 또는 원시 디스크 장치 그룹 만들기 및 관리에 대한 지침은 “[장치 그룹 관리](#)” [113]를 참조하십시오.

클러스터에 가장 적합한 복제 접근 방식을 선택하려면 먼저 호스트 기반 데이터 복제와 저장소 기반 데이터 복제를 모두 이해해야 합니다. Oracle Solaris Cluster Geographic Edition 프레임워크를 사용하여 재해 복구를 위해 데이터 복제를 관리하는 것과 관련된 자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)를 참조하십시오.

이 장은 다음 절로 구성됩니다.

- “[데이터 복제 이해](#)” [93]
- “[캠퍼스 클러스터 내에서 저장소 기반 데이터 복제 사용](#)” [95]

데이터 복제 이해

Oracle Solaris Cluster 소프트웨어는 호스트 기반과 저장소 기반 데이터 복제를 지원합니다.

- 호스트 기반 데이터 복제는 소프트웨어를 사용하여 지리적으로 분산된 클러스터 간에 디스크 볼륨을 실시간으로 복제합니다. 원격 미러 복제를 사용하면 기본 클러스터의 마스터 볼륨에서 지리적으로 분산된 보조 클러스터의 마스터 볼륨으로 데이터를 복제할 수 있

습니다. 원격 미러 비트맵이 1차 디스크의 마스터 볼륨과 2차 디스크의 마스터 볼륨 사이의 차이를 추적합니다. 예를 들어 클러스터 간 복제 또는 클러스터와 클러스터에 없는 호스트 간 복제에 사용된 호스트 기반 복제 소프트웨어는 Oracle Solaris의 가용성 제품군 기능입니다.

호스트 기반 데이터 복제는 특별한 저장소 어레이가 아닌 호스트 리소스를 사용하므로 적은 비용이 드는 데이터 복제 솔루션입니다. Oracle Solaris OS가 실행되는 여러 호스트에서 공유 볼륨에 데이터를 쓸 수 있도록 구성된 데이터베이스, 응용 프로그램 또는 파일 시스템은 지원되지 않습니다(예: Oracle RAC).

- 두 클러스터 간의 호스트 기반 데이터 복제 사용에 대한 자세한 내용은 [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Solaris Availability Suite](#)를 참조하십시오.
- Geographic Edition 프레임워크를 사용하지 않는 호스트 기반 복제의 예를 보려면 부록 A, [부록 A. 배치 예제: Availability Suite 소프트웨어를 사용하여 클러스터 간 호스트 기반 데이터 복제 구성](#)을 참조하십시오.
- 저장소 기반 데이터 복제에서는 저장소 컨트롤러의 소프트웨어를 사용하여 클러스터 노드에서 저장 장치로 데이터 복제 작업을 이동합니다. 이 소프트웨어는 클러스터 요청에 응답하기 위해 전원을 처리하는 일부 노드를 해제합니다. 클러스터 내부에서 또는 클러스터 간에 데이터를 복제할 수 있는 저장소 기반 소프트웨어에는 EMC SRDF가 있습니다. 저장소 기반 데이터 복제는 특히 캠퍼스 클러스터 구성에서 중요하며 필요한 기반구조를 단순화할 수 있습니다.
- 캠퍼스 클러스터 환경에서 저장소 기반 데이터 복제를 사용하는 데 대한 자세한 내용은 [“캠퍼스 클러스터 내에서 저장소 기반 데이터 복제 사용” \[95\]](#)을 참조하십시오.
- 둘 이상의 클러스터 간 저장소 기반 복제와 해당 프로세스를 자동화하는 Oracle Solaris Cluster Geographic Edition 제품 사용에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)를 참조하십시오.

지원되는 데이터 복제 방법

Oracle Solaris Cluster 소프트웨어는 클러스터 간 또는 클러스터 내에서 다음과 같은 데이터 복제 방법을 지원합니다.

- **클러스터 간 복제** - Oracle Solaris Cluster Geographic Edition 소프트웨어를 사용한 재해 복구를 위해서는 호스트 기반 또는 저장소 기반 복제를 사용해서 클러스터 간 데이터 복제를 수행할 수 있습니다. 일반적으로 둘을 조합하기 보다는 호스트 기반 복제와 저장소 기반 복제 중 하나를 선택합니다. Geographic Edition 소프트웨어를 사용하여 두 가지 복제 유형을 모두 관리할 수 있습니다.

복제 유형	데이터 복제 제품
호스트 기반 복제	Oracle Solaris의 가용성 제품군 기능.
저장소 기반 복제	EMC SRDF(Symmetrix Remote Data Facility).

복제 유형	데이터 복제 제품
	Hitachi TrueCopy 및 Hitachi Universal Replicator(Oracle Solaris Cluster Geographic Edition 프레임워크 사용) Oracle ZFS 저장소 어플라이언스.

자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#) 의 “Data Replication”를 참조하십시오.

Oracle Solaris Cluster Geographic Edition 소프트웨어 없이 호스트 기반 복제를 사용하려면 [부록 A. 배치 예제: Availability Suite 소프트웨어를 사용하여 클러스터 간 호스트 기반 데이터 복제 구성](#)의 지침을 참조하십시오.

Oracle Solaris Cluster Geographic Edition 소프트웨어 없이 저장소 기반 복제를 사용하려면 해당 복제 소프트웨어 설명서를 참조하십시오.

- **캠퍼스 클러스터 내 복제** - 이 방법은 호스트 기반 미러링 대신 사용됩니다.

복제 유형	데이터 복제 제품
저장소 기반 복제	EMC SRDF(Symmetrix Remote Data Facility). Hitachi TrueCopy 및 Hitachi Universal Replicator(Oracle Solaris Cluster Geographic Edition 프레임워크 사용)

- **응용 프로그램 기반 복제** - Oracle Data Guard는 응용 프로그램 기반 복제 소프트웨어의 예입니다. 이 유형의 단일 인스턴스 Oracle 데이터베이스 또는 RAC 데이터베이스를 복제하기 위해 Geographic Edition 프레임워크에서 재해 복구용으로만 사용됩니다.

자세한 내용은 [Oracle Solaris Cluster Geographic Edition Data Replication Guide for Oracle Data Guard](#)를 참조하십시오.

캠퍼스 클러스터 내에서 저장소 기반 데이터 복제 사용

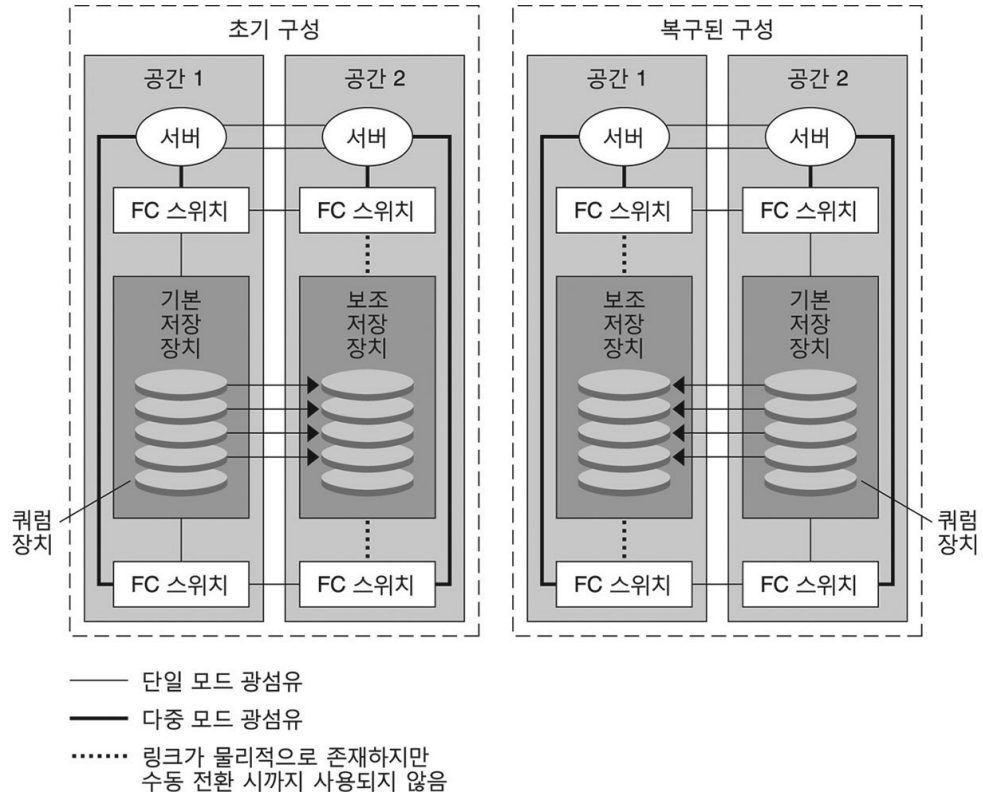
저장소 기반 데이터 복제에서는 저장 장치에 설치된 소프트웨어를 사용하여 캠퍼스 클러스터라는 클러스터 내부의 복제를 관리합니다. 그러한 소프트웨어는 특정 저장 장치에만 해당되며 재해 복구에는 사용되지 않습니다. 저장소 기반 데이터 복제를 구성할 때는 저장 장치와 함께 제공된 설명서를 참조하십시오.

사용하는 소프트웨어에 따라 저장소 기반 데이터 복제와 함께 자동 또는 수동 페일오버를 사용할 수 있습니다. Oracle Solaris Cluster에서는 복제의 수동 및 자동 페일오버를 둘 다 EMC SRDF 소프트웨어와 함께 사용합니다.

이 절에서는 캠퍼스 클러스터에서 사용되는 저장소 기반 데이터 복제에 대해 설명합니다. [그림 1](#)에서는 두 저장소 어레이 간에 데이터가 복제되는 샘플 2공간 구성을 보여 줍니다. 이 구성에서는 기본 저장소 어레이가 첫번째 공간에 포함되어 두 공간의 노드에 데이터를 제공합니다. 또한 기본 저장소 어레이는 보조 저장소 어레이에 복제할 데이터를 제공하기도 합니다.

주 - 그림 1에서는 쿼럼 장치가 복제되지 않은 볼륨에 있음을 보여 줍니다. 복제된 볼륨은 쿼럼 장치로 사용할 수 없습니다.

그림 1 저장소 기반 데이터 복제를 사용한 2룸 구성



Oracle Solaris Cluster에서 EMC SRDF와의 저장소 기반 동기식 복제가 지원됩니다. EMC SRDF에 대해서는 비동기식 복제가 지원되지 않습니다.

EMC SRDF의 Domino 모드나 Adaptive Copy 모드를 사용하지 마십시오. Domino 모드는 대상을 사용할 수 없을 때 로컬 및 대상 SRDF 볼륨을 호스트에 사용할 수 없게 합니다. Adaptive Copy 모드는 일반적으로 데이터 마이그레이션과 데이터 이동에 사용되며 재해 복구에는 사용하지 않는 것이 좋습니다.

원격 저장 장치와의 접속이 끊어지면 never 또는 async의 Fence_level을 지정하여 기본 클러스터에서 실행 중인 응용 프로그램이 차단되지 않았는지 확인합니다. data 또는 status의 Fence_level을 지정할 경우 업데이트를 원격 저장 장치로 복사할 수 없으면 기본 저장 장치가 업데이트를 거부합니다.

캠퍼스 클러스터 내에서 저장소 기반 데이터 복제를 사용할 경우 요구사항 및 제한 사항

데이터 무결성을 보장하려면 다중 경로 및 적절한 RAID 패키지를 사용합니다. 다음 목록에는 저장소 기반 데이터 복제를 사용하는 클러스터 구성을 구현하기 위한 고려 사항이 포함되어 있습니다.

- 자동 파일오버에 대해 클러스터를 구성하는 경우 동기식 복제를 사용합니다.
복제된 볼륨의 자동 파일오버에 대한 클러스터 구성 지침은 [“저장소 기반의 복제된 장치 구성 및 관리” \[101\]](#)를 참조하십시오. 캠퍼스 클러스터 설계 요구사항에 대한 자세한 내용은 [Oracle Solaris Cluster Hardware Administration Manual](#)의 [“Shared Data Storage”](#)를 참조하십시오.
- 특정 응용 프로그램 관련 데이터는 비동기식 데이터 복제에 적합하지 않을 수 있습니다. 응용 프로그램의 동작을 파악하여 저장 장치 간에 응용 프로그램 관련 데이터를 복제할 최적의 방법을 결정합니다.
- 노드 간 거리는 Oracle Solaris Cluster 광 섬유 채널 및 상호 연결 기반구조로 제한됩니다. 현재 제한 사항과 지원되는 기술에 대한 자세한 내용은 Oracle 서비스 공급자에게 문의하십시오.
- 복제된 볼륨을 쿼럼 장치로 구성하지 마십시오. 복제되지 않은 공유 볼륨에서 쿼럼 장치를 찾거나 쿼럼 서버를 사용합니다.
- 데이터의 기본 복사본만 클러스터 노드에 표시됩니다. 그렇지 않으면 볼륨 관리자가 데이터의 기본 복사본과 보조 복사본에 동시에 액세스를 시도할 수 있습니다. 데이터 복사본 표시 제어에 대한 자세한 내용은 저장소 어레이와 함께 제공된 설명서를 참조하십시오.
- EMC SRDF를 사용하여 사용자는 복제된 장치 그룹을 정의할 수 있습니다. 복제 장치 그룹마다 이름이 같은 Oracle Solaris Cluster 장치 그룹이 필요합니다.
- 동시 또는 계단식 RDF 장치에서 EMC SRDF를 사용한 세 사이트 또는 세 데이터 센터 구성의 경우 참가하는 모든 클러스터 노드의 Solutions Enabler SYMCLI 옵션 파일에서 다음 항목을 추가해야 합니다.

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

이 항목을 통해 클러스터 소프트웨어가 두 SRDF 동기 사이트 간에 응용 프로그램의 이동을 자동화할 수 있습니다. 항목에서 *rdf-group-number*는 호스트의 로컬 Symmetrix를 보조 사이트의 Symmetrix에 연결하는 RDF 그룹을 나타냅니다.

세 데이터 센터 구성에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)의 [“Three-Data-Center \(3DC\) Topologies”](#)를 참조하십시오.

- 클러스터 내부에서 복제할 때는 SRDF에 Oracle RAC(Oracle Real Application Clusters)가 지원되지 않습니다. 현재 기본 복제본이 아닌 복제본에 연결된 노드에는 쓰기 액세스가 없습니다. 복제된 장치에서는 클러스터의 모든 노드에서 직접 쓰기 액세스를 필요로 하는 확장 가능 응용 프로그램을 지원할 수 없습니다.
- Sun Cluster용 Solaris Volume Manager 다중 소유자 디스크 세트는 지원되지 않습니다.
- EMC SRDF에서 Domino 모드나 Adaptive Copy 모드를 사용하지 마십시오. 자세한 내용은 [“캠퍼스 클러스터 내에서 저장소 기반 데이터 복제 사용” \[95\]](#)를 참조하십시오.

캠퍼스 클러스터 내에서 저장소 기반 데이터 복제를 사용할 경우 수동 복구 문제

모든 캠퍼스 클러스터에서와 마찬가지로, 저장소 기반 데이터 복제를 사용하는 클러스터는 한 번 실패할 경우 일반적으로 개입이 필요하지 않습니다. 하지만 [그림 1](#)과 같이 수동 페일오버를 사용하는 경우 기본 저장 장치가 있는 공간이 없다면 2 노드 클러스터에 문제가 발생합니다. 남은 노드는 쿼럼 장치를 보유할 수 없으며 클러스터 구성원으로 부트할 수 없습니다. 이 경우 다음의 수동 개입이 클러스터에 필요합니다.

1. Oracle 서비스 공급자는 남은 노드가 클러스터 구성원으로 부트하도록 재구성해야 합니다.
2. 사용자 또는 해당 Oracle 서비스 공급자는 보조 저장 장치의 복제되지 않은 볼륨을 쿼럼 장치로 구성해야 합니다.
3. 사용자 또는 해당 Oracle 서비스 공급자는 남은 노드가 보조 저장 장치를 기본 저장소로 사용하도록 구성해야 합니다. 이 재구성에는 볼륨 관리자 볼륨 재구축, 데이터 복원 또는 저장소 볼륨과의 응용 프로그램 연관 변경이 포함될 수 있습니다.

저장소 기반 데이터 복제를 사용할 경우 최고 사례

저장소 기반 데이터 복제에 EMC SRDF 소프트웨어를 사용할 때는 정적 장치 대신 동적 장치를 사용하십시오. 정적 장치는 복제 기본 변경에 몇 분 가량 소요되며 페일오버 시간에 영향을 줄 수 있습니다.

전역 장치, 디스크 경로 모니터링 및 클러스터 파일 시스템 관리

이 장에서는 전역 장치, 디스크 경로 모니터링 및 클러스터 파일 시스템 관리에 대한 정보 및 절차를 소개합니다.

- “전역 장치 및 전역 이름 공간 관리 개요” [99]
- “저장소 기반의 복제된 장치 구성 및 관리” [101]
- “클러스터 파일 시스템 관리 개요” [113]
- “장치 그룹 관리” [113]
- “저장 장치에 대한 SCSI 프로토콜 설정 관리” [138]
- “클러스터 파일 시스템 관리” [143]
- “디스크 경로 모니터링 관리” [148]

이 장에 있는 관련 절차에 대한 자세한 내용은 표 7. “작업 맵: 장치 그룹 관리”를 참조하십시오.

전역 장치, 전역 이름 공간, 장치 그룹, 디스크 경로 모니터링 및 클러스터 파일 시스템과 관련된 개념 정보는 [Oracle Solaris Cluster 4.3 Concepts Guide](#)를 참조하십시오.

전역 장치 및 전역 이름 공간 관리 개요

Oracle Solaris Cluster 장치 그룹의 관리는 클러스터에 설치된 볼륨 관리자가 맡습니다. Solaris Volume Manager는 “클러스터 인식형”이므로 사용자가 Solaris Volume Manager `metaset` 명령을 사용하여 장치 그룹을 추가, 등록, 제거할 수 있습니다. 자세한 내용은 [metaset\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

Oracle Solaris Cluster 소프트웨어는 자동으로 클러스터의 각 디스크와 테이프 장치에 대한 원시 디스크 장치 그룹을 만듭니다. 그러나 사용자가 클러스터 장치 그룹을 전역 장치로 액세스할 때까지 클러스터 장치 그룹이 오프라인 상태를 유지합니다. 장치 그룹이나 볼륨 관리자 디스크 그룹을 관리하는 경우 사용자가 그룹의 기본 노드인 클러스터 노드에 있어야 합니다.

일반적으로 전역 장치 이름 공간은 관리할 필요가 없습니다. 전역 이름 공간은 설치 과정에서 자동으로 설정되고 Oracle Solaris OS 재부트 과정에서 자동으로 업데이트됩니다. 그

러나 전역 이름 공간을 업데이트해야 하는 경우에는 임의의 클러스터 노드에서 `cldevice populate` 명령을 실행할 수 있습니다. 이 명령을 실행하면 나중에 클러스터에 포함될 노드뿐 아니라 다른 모든 클러스터 노드 구성원에서 전역 이름 공간이 업데이트됩니다.

Solaris Volume Manager에 대한 전역 장치 사용 권한

전역 장치 권한을 변경하면 Solaris Volume Manager 및 디스크 장치에 대한 클러스터의 모든 노드에 변경사항이 자동으로 전파됩니다. 전역 장치에 대한 사용 권한을 변경하려면 클러스터의 모든 노드에서 직접 사용 권한을 변경해야 합니다. 예를 들어, 전역 장치 `/dev/global/dsk/d3s0`의 권한을 644로 변경하려는 경우 클러스터의 모든 노드에서 다음 명령을 실행해야 합니다.

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

전역 장치 동적 재구성

클러스터에서 디스크 및 테이프 장치에 대한 동적 재구성 작업을 완료할 때는 다음과 같은 사항을 고려해야 합니다.

- Oracle Solaris 동적 재구성 기능에 대해 문서화된 모든 요구사항, 절차 및 제한 사항은 Oracle Solaris Cluster 동적 재구성 지원에도 적용됩니다. 운영체제의 작동이 정지된 경우만은 예외입니다. 따라서 Oracle Solaris Cluster 소프트웨어와 함께 동적 재구성 기능을 사용하기 전에 Oracle Solaris 동적 재구성 기능에 대한 설명서를 검토하십시오. 특히 동적 재구성 연결 종료 작업 중 비네트워크 IO 장치에 영향을 주는 문제를 검토해야 합니다.
- Oracle Solaris Cluster에서는 기본 노드에서 활성 장치에 대한 동적 재구성 보드 제거 작업을 수행할 수 없습니다. 동적 재구성 작업은 기본 노드의 비활성 장치와 보조 노드의 모든 장치에 대해 수행할 수 있습니다.
- 동적 재구성 작업이 끝나면 작업 이전과 마찬가지로 클러스터 데이터 액세스가 계속됩니다.
- Oracle Solaris Cluster에서는 쿼럼 장치의 가용성에 영향을 주는 동적 재구성 작업을 수행할 수 없습니다. 자세한 내용은 [“쿼럼 장치 동적 재구성” \[158\]](#)을 참조하십시오.



주의 - 보조 노드에 대해 동적 재구성 작업을 수행할 때 현재 기본 노드에 장애가 발생하면 클러스터 가용성이 영향을 받습니다. 새로운 보조 노드가 제공될 때까지 기본 노드를 페일오버할 수 없습니다.

전역 장치에 대해 동적 재구성 작업을 수행하려면 다음 단계를 순서대로 완료하십시오.

표 5 작업 맵: 디스크 및 테이프 장치 동적 재구성

작업	지침
1. 현재 기본 노드에 대해 활성 장치 그룹에 영향을 주는 동적 재구성 작업을 수행해야 할 경우 장치에서 동적 재구성 제거 작업을 수행하기 전에 기본 및 보조 노드를 전환합니다.	장치 그룹에 대한 기본 노드를 전환하는 방법 [135]
2. 제거하려는 장치에 대해 동적 재구성 제거 작업을 수행합니다.	시스템과 함께 제공된 설명서를 참조하십시오.

저장소 기반의 복제된 장치 구성 및 관리

저장소 기반 복제를 사용하여 복제되는 장치를 포함하도록 Oracle Solaris Cluster 장치 그룹을 구성할 수 있습니다. Oracle Solaris Cluster 소프트웨어는 저장소 기반 복제를 위해 EMC Symmetrix Remote Data Facility 소프트웨어를 지원합니다.

EMC Symmetrix Remote Data Facility 소프트웨어를 사용하여 데이터를 복제하기 전에 저장소 기반 복제 설명서를 숙지하고 저장소 기반 복제 제품 및 최신 업데이트를 시스템에 설치해야 합니다. 저장소 기반 복제 소프트웨어 설치에 대한 자세한 내용은 제품 설명서를 참조하십시오.

저장소 기반 복제 소프트웨어는 장치 한 쌍을 복제본으로 구성하며 이때 장치 한 개는 기본 복제본, 다른 장치는 보조 복제본입니다. 모든 지정된 시간에 하나의 노드 세트에 연결된 장치가 기본 복제본이 됩니다. 다른 노드 세트에 연결된 장치는 보조 복제본이 됩니다.

Oracle Solaris Cluster 구성에서 복제본이 속한 Oracle Solaris Cluster 장치 그룹이 이동할 때마다 기본 복제본이 자동으로 이동합니다. 따라서 Oracle Solaris Cluster 구성에서 직접 기본 복제본을 이동하면 안 됩니다. 대신 연결된 Oracle Solaris Cluster 장치 그룹을 이동하여 인계해야 합니다.



주의 - 사용자가 만드는 Oracle Solaris Cluster 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

EMC Symmetrix Remote Data Facility 복제된 장치 관리

다음 표에는 EMC SRDF(Symmetrix Remote Data Facility) 저장소 기반의 복제된 장치를 설정 및 관리하기 위해 수행할 작업이 나열되어 있습니다.

표 6 작업 맵: EMC SRDF 저장소 기반의 복제된 장치 관리

작업	지침
저장 장치 및 노드에 SRDF 소프트웨어를 설치합니다.	EMC 저장 장치와 함께 제공된 설명서

작업	지침
EMC 복제 그룹 구성	EMC SRDF 복제 그룹을 구성하는 방법 [102]
DID 장치 구성	EMC SRDF를 사용하여 복제를 위해 DID 장치를 구성하는 방법 [104]
복제된 그룹 등록	장치 그룹 추가 및 등록 방법(Solaris Volume Manager) [120]
구성 확인	EMC SRDF 복제된 전역 장치 그룹 구성을 확인하는 방법 [105]
캠퍼스 클러스터의 기본 공간이 전체 실패 후 수동으로 데이터 복구	기본 공간이 전체 실패한 후 EMC SRDF 데이터를 복구하는 방법 [111]

▼ EMC SRDF 복제 그룹을 구성하는 방법

- 시작하기 전에
- EMC SRDF(Symmetrix Remote Data Facility) 복제 그룹을 구성하기 전에 EMC Solutions Enabler 소프트웨어를 모든 클러스터 노드에 설치해야 합니다. 먼저 클러스터의 공유 디스크에 EMC SRDF 장치 그룹을 구성합니다. EMC SRDF 장치 그룹 구성 방법에 대한 자세한 내용은 EMC SRDF 제품 설명서를 참조하십시오.
 - EMC SRDF를 사용하는 경우 정적 장치 대신 동적 장치를 사용하십시오. 정적 장치는 복제 기본 변경에 몇 분 가량 소요되며 페일오버 시간에 영향을 줄 수 있습니다.



주의 - 사용자가 만드는 Oracle Solaris Cluster 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

1. 저장소 어레이에 연결된 모든 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 동시 SRDF 또는 계단식 장치를 사용하는 세 사이트 또는 세 데이터 센터 구현의 경우 `SYMAPI_2SITE_CLUSTER_DG` 매개변수를 설정합니다.
모든 참여 클러스터 노드의 Solutions Enabler 옵션 파일에 다음 항목을 추가합니다.

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

`device-group` 장치 그룹의 이름을 지정합니다.

`rdf-group-number` 호스트의 로컬 Symmetrix를 보조 사이트의 Symmetrix에 연결하는 RDF 그룹을 지정합니다.

이 항목을 통해 클러스터 소프트웨어가 두 SRDF 동기 사이트 간에 응용 프로그램의 이동을 자동화할 수 있습니다.

세 데이터 센터 구성에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition Overview](#)의 “Three-Data-Center (3DC) Topologies”를 참조하십시오.

3. 복제된 데이터로 구성된 각 노드에서 Symmetrix 장치 구성을 검색합니다.
이 작업은 몇 분 가량 소요될 수 있습니다.

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

4. 아직 복제본 쌍을 만들지 않은 경우 지금 만듭니다.

`symrdf` 명령을 사용하여 복제본 쌍을 만듭니다. 복제본 쌍 만들기에 대한 지침은 SRDF 설명서를 참조하십시오.

주 - 세 사이트 또는 세 데이터 센터 구현을 위해 동시 RDF 장치를 사용 중인 경우 모든 `symrdf` 명령에 다음 매개변수를 추가합니다.

```
-rdg rdf-group-number
```

`symrdf` 명령에 RDF 그룹 번호를 지정하면 `symrdf` 작업이 올바른 RDF 그룹으로 지정됩니다.

5. 복제된 장치로 구성된 각 노드에서 데이터 복제가 올바르게 설정되었는지 확인합니다.

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

6. 장치 그룹 스왑을 수행합니다.

a. 기본 및 보조 복제본이 동기화되는지 확인합니다.

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

b. `symdg show` 명령을 사용하여 기본 복제본을 포함하는 노드와 보조 복제본을 포함하는 노드를 확인합니다.

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

RDF1 장치가 있는 노드는 기본 복제본을 포함하고 RDF2 장치 상태인 노드는 보조 복제본을 포함합니다.

c. 보조 복제본을 사용으로 설정합니다.

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

d. RDF1 및 RDF2 장치를 스왑합니다.

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

e. 복제 쌍을 사용으로 설정합니다.

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

f. 기본 노드 및 보조 복제본이 동기화되는지 확인합니다.

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

7. 원래 기본 복제본이 있던 노드에 대해 5단계를 모두 반복합니다.

다음 순서 EMC SRDF 복제된 장치에 대해 장치 그룹을 구성한 후에 복제된 장치에 사용되는 DID(장치 식별자) 드라이버를 구성해야 합니다.

▼ EMC SRDF를 사용하여 복제를 위해 DID 장치를 구성하는 방법

이 절차에서는 복제된 장치에 사용되는 DID(장치 식별자) 드라이버를 구성합니다. 지정된 DID 장치 인스턴스들이 서로의 복제본이고 지정된 장치 그룹에 속하는지 확인합니다.

시작하기 전에 `phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 구성된 RDF1 및 RDF2 장치에 해당하는 DID 장치를 확인합니다.

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

주 - 시스템에 전체 Oracle Solaris 장치 패치가 표시되지 않으면 환경 변수 `SYMCLI_FULL_PDEVNAME`을 1로 설정하고 `symdg -show` 명령을 다시 입력합니다.

3. Oracle Solaris 장치에 해당하는 DID 장치를 확인합니다.

```
# cldevice list -v
```

4. 일치하는 DID 장치 쌍마다 인스턴스를 하나의 복제된 DID 장치로 조합합니다. RDF2/secondary 측에서 다음 명령을 실행합니다.

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

주 - SRDF 데이터 복제 장치에는 `-t` 옵션이 지원되지 않습니다.

`-t replication-type`

복제 유형을 지정합니다. EMC SRDF일 경우 **SRDF**를 입력합니다.

`-g replication-device-group`

`symdg show` 명령에 표시된 대로 장치 그룹 이름을 지정합니다.

`-d destination-instance`

RDF1 장치에 해당하는 DID 인스턴스를 지정합니다.

source-instance

RDF2 장치에 해당하는 DID 인스턴스를 지정합니다.

주 - 잘못된 DID 장치를 조합할 경우 -b 옵션을 `scdidadm` 명령에 사용하여 두 DID 장치의 조합을 실행 취소합니다.

`scdidadm -b device`

-b *device* 인스턴스가 결합되었을 때 `destination_device`에 해당하는 DID 인스턴스입니다.

5. 복제 장치 그룹의 이름이 변경될 경우 다음과 같은 추가 단계를 수행합니다.
복제 장치 그룹(및 해당 전역 장치 그룹) 이름이 변경될 경우 먼저 `scdidadm -b` 명령을 사용하여 복제된 장치 정보를 업데이트하여 기존 정보를 제거해야 합니다. 마지막 단계에서는 `cldevice combine` 명령을 사용하여 업데이트된 장치를 새로 만듭니다.
6. DID 인스턴스가 조합되었는지 확인합니다.
`cldevice list -v device`
7. SRDF 복제가 설정되었는지 확인합니다.
`cldevice show device`
8. 모든 노드에서 조합된 모든 DID 인스턴스에 대한 DID 장치에 액세스할 수 있는지 확인합니다.
`cldevice list -v`

다음 순서 복제된 장치에 사용되는 DID(장치 식별자) 드라이버를 구성한 후 EMC SRDF 복제된 전역 장치 그룹 구성을 확인해야 합니다.

▼ EMC SRDF 복제된 전역 장치 그룹 구성을 확인하는 방법

시작하기 전에 전역 장치 그룹을 확인하기 전에 먼저 해당 그룹을 만듭니다. Solaris Volume Manager, ZFS 또는 원시 디스크에서 장치 그룹을 사용할 수 있습니다. 자세한 내용은 다음을 참조하십시오.

- [장치 그룹 추가 및 등록 방법\(Solaris Volume Manager\) \[120\]](#)
- [장치 그룹 추가 및 등록 방법\(원시 디스크\) \[122\]](#)
- [복제된 장치 그룹 추가 및 등록 방법\(ZFS\) \[123\]](#)



주의 - 사용자가 만든 Oracle Solaris Cluster 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 기본 장치 그룹이 기본 복제본이 포함된 노드와 같은 노드에 해당하는지 확인합니다.

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

2. 시험 삼아 스위치오버를 수행하여 장치 그룹이 올바르게 구성되고 복제본이 노드 간에 이동할 수 있게 합니다.

장치 그룹이 오프라인일 경우 온라인으로 전환합니다.

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

-n nodename 장치 그룹이 전환되는 노드입니다. 이 노드가 새 기본 노드가 됩니다.

3. 다음 명령의 출력을 비교하여 스위치오버가 성공적인지 확인합니다.

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

예: Oracle Solaris Cluster에 대한 SRDF 복제 그룹 구성

이 예에서는 클러스터에서 SRDF 복제를 설정하는 데 필요한 Oracle Solaris Cluster 특정 단계를 완료합니다. 이 예에서는 다음 작업을 이미 수행했다고 가정합니다.

- 어레이 간에 LUNS 복제 쌍 지정 완료
- 저장 장치 및 클러스터 노드에 SRDF 소프트웨어 설치

예 30 복제본 쌍 만들기

이 예에는 노드 2개가 symmetrix 하나에 연결되고 나머지 노드 2개가 두번째 symmetrix에 연결되는 4노드 클러스터가 포함됩니다. SRDF 장치를 dg1이라고 합니다.

모든 노드에서 다음 명령 실행

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

```
# symdev list pd
```

```
Symmetrix ID: 000187990182
```

Device Name	Directors	Device
-------------	-----------	--------

```

-----
Sym Physical          SA :P DA :IT Config      Attribute  Sts  Cap
-----
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

```

... RDF1 측의 모든 노드에서 다음 명령 실행

```

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

```

RDF2 측의 모든 노드에서 다음 명령 실행

```

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

```

예 31 데이터 복제 설정 확인

다음 명령은 클러스터의 노드 하나에서 수행됩니다.

```
# symdg show dg1
```

Group Name: dg1

```

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

```

```

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

```

Standard (STD) Devices (1):

```
{
```

```

-----
LdevName          PdevName          Sym      Cap
Dev  Att. Sts      (MB)
-----

```

```

DEV001            /dev/rdisk/c5t6006048000018790002353594D303637d0s2 0067  RW   4315
}

```

Device Group RDF Information

...

```
# symrdf -g dg1 establish
```

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

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

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

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

#

symrdf -g dg1 query

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

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

Legend for MODES:

```
M(mode 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
```

#

예 32 사용된 디스크에 해당하는 DID 표시

RDF1 및 RDF2 측에 같은 절차가 적용됩니다.

dymdg show dg 명령 출력의 PdevName 필드에서 확인할 수 있습니다.

RDF1 측에서 다음 명령 실행

```
# symdg show dg1
```

Group Name: dg1

Group Type : RDF1 (RDFA)

...

Standard (STD) Devices (1):

```
{
```

```
-----
LdevName      PdevName      Sym      Cap
Dev  Att. Sts  (MB)
-----
```

```
DEV001        /dev/rdisk/c5t6006048000018790002353594D303637d0s2 0067  RW  4315
}
```

Device Group RDF Information

...

해당하는 *DID* 가져오기

```
# cldevice list | grep c5t6006048000018790002353594D303637d0
```

```
217    pmoney1:/dev/rdisk/c5t6006048000018790002353594D303637d0 /dev/did/rdisk/d217
```

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

해당하는 *DID* 나열

```
# cldevice show d217
```

=== DID Device Instances ===

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

#

RDF2 측에서 다음 명령 실행

```
# symdg show dg1
```

Group Name: dg1

Group Type : RDF2 (RDFA)

...

Standard (STD) Devices (1):

```
{
```

```
-----
LdevName      PdevName      Sym      Cap
Dev  Att. Sts  (MB)
-----
```

```
DEV001        /dev/rdisk/c5t6006048000018799018253594D303637d0s2 0067  WD  4315
```

```

    }

    Device Group RDF Information
'''
    해당하는 DID 가져오기
# cldevice list | grep c5t6006048000018799018253594D303637d0
108      pmoney4:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108
108      pmoney3:/dev/rdisk/c5t6006048000018799018253594D303637d0 /dev/did/rdsk/d108

    해당하는 DID 나열
# cldevice show d108

=== DID Device Instances ===

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

#

```

예 33 DID 인스턴스 조합

RDF2 사이드에서 다음을 입력합니다.

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

예 34 조합된 DID 표시

클러스터의 아무 노드에서나 다음을 입력합니다.

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

```

=== DID Device Instances ===

DID Device Name:          /dev/did/rdsk/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

```

▼ 기본 공간이 전체 실패한 후 EMC SRDF 데이터를 복구하는 방법

이 절차에서는 캠퍼스 클러스터의 기본 공간이 전체 실패한 후 기본 공간이 보조 공간으로 페일오버한 다음 기본 공간이 다시 온라인 상태로 전환될 경우 데이터 복구를 수행합니다. 캠퍼스 클러스터의 기본 공간은 기본 노드이며 저장 사이트입니다. 공간의 전체 실패에는 공간에 있는 호스트와 저장소 둘 다의 실패가 포함됩니다. 기본 공간이 실패할 경우 Oracle Solaris Cluster에서는 보조 공간으로 자동 페일오버하고 보조 공간의 저장 장치를 읽고 쓸 수 있게 하며 해당 장치 그룹 및 리소스 그룹의 페일오버를 사용으로 설정합니다.

기본 공간이 다시 온라인 상태로 전환하면 보조 공간에 쓴 SRDF 장치 그룹에서 데이터를 수동으로 복구하고 데이터를 재동기화할 수 있습니다. 이 절차에서는 원래의 보조 공간(이 절차에서는 보조 공간에 phys-campus-2 사용)에서 원래의 기본 공간(phys-campus-1)으로 데이터를 동기화하여 SRDF 장치를 복구합니다. 또한 이 절차에서는 SRDF 장치 그룹 유형을 phys-campus-2에서 RDF1로, phys-campus-1에서 RDF2로 변경합니다.

시작하기 전에 수동 페일오버를 수행하려면 먼저 EMC 복제 그룹을 등록하고 EMC 복제 그룹 및 DID 장치를 구성해야 합니다. Solaris Volume Manager 장치 그룹 만들기에 대한 자세한 내용은 [장치 그룹 추가 및 등록 방법\(Solaris Volume Manager\) \[120\]](#)을 참조하십시오.

주 - 이 지침은 기본 공간이 완전히 페일오버하고 다시 온라인 상태로 전환한 후 수동으로 SRDF 데이터를 복구하는 데 사용할 수 있는 한 가지 방법을 보여 줍니다. 다른 방법은 EMC 설명서를 확인하십시오.

이 단계를 수행하려면 캠퍼스 클러스터의 기본 공간에 로그인합니다. 아래 절차에서 *dg1*은 SRDF 장치 그룹 이름입니다. 실패할 때 이 절차의 기본 공간은 phys-campus-1이고 보조 공간은 phys-campus-2입니다.

1. 캠퍼스 클러스터의 기본 공간에 로그인하고 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 기본 공간에서 `symrdf` 명령을 사용하여 RDF 장치의 복제 상태를 질의하고 해당 장치에 대한 정보를 봅니다.

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

팁 - split 상태의 장치 그룹은 동기화되지 않습니다.

3. RDF 쌍 상태가 split이고 장치 그룹 유형이 RDF1이면 SRDF 장치 그룹 페일오버가 강제 실행됩니다.

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

4. RDF 장치 상태를 봅니다.

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

5. **페일오버 후 페일오버한 RDF 장치의 데이터를 스왑할 수 있습니다.**

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

6. **RDF 장치에 대한 기타 정보와 상태를 확인합니다.**

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

7. **기본 공간에 SRDF 장치 그룹을 설정합니다.**

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

8. **장치 그룹이 동기화된 상태이고 장치 그룹 유형이 RDF2인지 확인합니다.**

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

예 35 기본 사이트 페일오버 후 수동으로 EMC SRDF 데이터 복구

이 예에서는 캠퍼스 클러스터의 기본 공간이 페일오버하고 보조 공간이 데이터를 인계 및 기록한 다음 기본 공간이 다시 온라인 상태로 전환된 후 수동으로 EMC SRDF 데이터를 복구하는 데 필요한 Oracle Solaris Cluster 특정 단계를 제공합니다. 예에서 SRDF 장치 그룹은 *dg1*이고 표준 논리 장치는 DEV001입니다. 실패할 때 기본 공간은 phys-campus-1이고 보조 공간은 phys-campus-2입니다. 캠퍼스 클러스터의 기본 공간 phys-campus-1에서 단계를 수행합니다.

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

```
phys-campus-1# syndg 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# syndg 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# syndg 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
```

클러스터 파일 시스템 관리 개요

클러스터 파일 시스템 관리에 특별한 Oracle Solaris Cluster 명령은 필요하지 않습니다. 다른 Oracle Solaris 파일 시스템을 관리하는 경우와 마찬가지로 `mount` 및 `newfs`와 같은 표준 Oracle Solaris 파일 시스템 명령을 사용하여 클러스터 파일 시스템을 관리합니다. 클러스터 파일 시스템을 마운트할 때는 `mount` 명령에 `-g` 옵션을 지정합니다. 클러스터 파일 시스템이 UFS를 사용하여 부트할 때 자동으로 마운트될 수도 있습니다. 클러스터 파일 시스템은 전역 클러스터의 노드에서만 볼 수 있습니다.

주 - 클러스터 파일 시스템이 파일을 읽을 때는 파일 시스템이 해당 파일에 대한 액세스 시간을 업데이트하지 않습니다.

클러스터 파일 시스템 제한

클러스터 파일 시스템 관리에 적용되는 제한 사항은 다음과 같습니다.

- `unlink` 명령은 비어 있지 않은 디렉토리에 지원되지 않습니다. 자세한 내용은 [unlink\(1M\)](#) 매뉴얼 페이지를 참조하십시오.
- `lockfs -d` 명령은 지원되지 않습니다. 해결 방법으로 `lockfs -n`을 사용합니다.
- 다시 마운트할 때 `directio` 마운트 옵션을 추가하여 클러스터 파일 시스템을 다시 마운트할 수 없습니다.

장치 그룹 관리

클러스터 요구사항이 변경됨에 따라 클러스터에서 장치 그룹을 추가, 제거 또는 수정해야 할 수도 있습니다. Oracle Solaris Cluster는 이러한 변경을 위해 사용할 수 있는 `clsetup`이라는 대화식 인터페이스를 제공합니다. `clsetup` 유틸리티는 `cluster` 명령을 생성합니다. 몇 가지 절차 뒤에 다음과 같은 생성된 명령의 예가 나옵니다. 다음 표에서는 장치 그룹 관리에 대한 작업을 나열하고 이 절의 해당 절차에 대한 링크를 제공합니다.

주 - Oracle Solaris Cluster Manager 브라우저를 사용해서 장치 그룹을 온라인 및 오프라인으로 전환할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

Oracle Solaris Cluster 소프트웨어는 자동으로 클러스터의 각 디스크와 테이프 장치에 대한 원시 디스크 장치 그룹을 만듭니다. 그러나 사용자가 클러스터 장치 그룹을 전역 장치로 액세스할 때까지 클러스터 장치 그룹이 오프라인 상태를 유지합니다.



주의 - 다른 노드가 활성 클러스터 멤버이고 그 중 하나 이상이 디스크 세트를 소유하고 있는 경우에는 클러스터 외부에서 부트되는 클러스터 노드에서 `metaset -s setname -f -t` 명령을 실행하지 마십시오.

표 7 작업 맵: 장치 그룹 관리

작업	지침
<code>cldevice populate</code> 명령을 사용하여 재구성 재부트 없이 전역 장치 이름 공간 업데이트	전역 장치 이름 공간 업데이트 방법 [115]
전역 장치 이름 공간에 사용되는 <code>lofi</code> 장치의 크기 변경	전역 장치 이름 공간에 사용되는 <code>lofi</code> 장치의 크기 변경 방법 [116]
기존 전역 장치 이름 공간 이동	전용 파티션에서 <code>lofi</code> 장치로 전역 장치 이름 공간 마이그레이션 방법 [117]
	<code>lofi</code> 장치에서 전용 파티션으로 전역 장치 이름 공간 마이그레이션 방법 [118]
<code>metaset</code> 명령을 사용하여 Solaris Volume Manager 디스크 세트를 추가하고 장치 그룹으로 등록	장치 그룹 추가 및 등록 방법(Solaris Volume Manager) [120]
<code>cldevicegroup</code> 명령을 사용하여 원시 디스크 장치 그룹 추가 및 등록	장치 그룹 추가 및 등록 방법(원시 디스크) [122]
<code>cldevicegroup</code> 명령을 사용하여 ZFS에 대한 명명된 장치 그룹 추가	복제된 장치 그룹 추가 및 등록 방법(ZFS) [123]
<code>metaset</code> 및 <code>metaclear</code> 명령을 사용하여 구성에서 Solaris Volume Manager 장치 그룹 제거	“장치 그룹 제거 및 등록 해제 방법(Solaris Volume Manager)” [126]
<code>cldevicegroup</code> , <code>metaset</code> , <code>clsetup</code> 명령을 사용하여 모든 장치 그룹에서 노드 제거	모든 장치 그룹에서 노드를 제거하는 방법 [126]
<code>metaset</code> 명령을 사용하여 Solaris Volume Manager 장치 그룹에서 노드 제거	장치 그룹에서 노드를 제거하는 방법(Solaris Volume Manager) [127]
<code>cldevicegroup</code> 명령을 사용하여 원시 디스크 장치 그룹에서 노드 제거	원시 디스크 장치 그룹에서 노드를 제거하는 방법 [129]
<code>cldevicegroup</code> 을 생성하기 위해 <code>clsetup</code> 명령을 사용하여 장치 그룹의 등록 정보 변경	장치 그룹 등록 정보를 변경하는 방법 [130]
<code>cldevicegroup show</code> 명령을 사용하여 장치 그룹 및 등록 정보 표시	장치 그룹 구성을 나열하는 방법 [134]
<code>cldevicegroup</code> 을 생성하기 위해 <code>clsetup</code> 을 사용하여 장치 그룹에 사용할 보조 노드 수 변경	장치 그룹에 원하는 보조 수를 설정하는 방법 [132]
<code>cldevicegroup switch</code> 명령을 사용하여 장치 그룹에 대한 기본 노드 전환	장치 그룹에 대한 기본 노드를 전환하는 방법 [135]
<code>metaset</code> 명령을 사용하여 장치 그룹을 유지 보수 상태로 전환	장치 그룹을 유지 보수 상태로 전환하는 방법 [136]

▼ 전역 장치 이름 공간 업데이트 방법

새 전역 장치를 추가하는 경우 `cldevice populate` 명령을 실행하여 전역 장치 이름 공간을 수동으로 업데이트합니다.

주 - `cldevice populate` 명령을 실행하는 노드가 현재 클러스터 구성원이 아니면 명령이 적용되지 않습니다. `/global/.devices/node@ nodeID` 파일 시스템이 마운트되지 않은 경우에도 명령이 적용되지 않습니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 클러스터의 각 노드에서 `devfsadm` 명령을 실행합니다.
클러스터의 모든 노드에서 동시에 이 명령을 실행할 수 있습니다. 자세한 내용은 [devfsadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.
3. 이름 공간을 재구성합니다.

```
# cldevice populate
```

4. 디스크 세트를 만들기 전에 각 노드에서 `cldevice populate` 명령이 완료되었는지 확인합니다.

`cldevice` 명령이 한 노드에서만 실행되는 경우에도 이 명령은 모든 노드에서 동일한 명령을 원격으로 호출합니다. `cldevice populate` 명령이 프로세스를 완료했는지 확인하려면 클러스터의 각 노드에서 다음 명령을 실행합니다.

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

예 36 전역 장치 이름 공간 업데이트

다음 예에서는 성공적으로 `cldevice populate` 명령을 실행한 경우 생성되는 출력을 보여 줍니다.

```
# 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
```

▼ 전역 장치 이름 공간에 사용되는 lofi 장치의 크기 변경 방법

전역 클러스터의 노드 하나 이상에서 전역 장치 이름 공간에 lofi 장치를 사용하는 경우 이 절차를 수행하여 장치 크기를 변경합니다.

1. 전역 장치 이름 공간에 대한 lofi 장치의 크기를 조정할 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 노드에서 서비스를 제외하고 노드를 비클러스터 모드로 재부트합니다.

이 절차를 수행하는 동안 이 노드에서 전역 장치가 처리되지 않도록 하기 위해 이 작업을 수행합니다. 지침은 [비클러스터 모드로 노드를 부트하는 방법 \[90\]](#)을 참조하십시오.

3. 전역 장치 파일 시스템을 마운트 해제하고 lofi 장치를 분리합니다.

전역 장치 파일 시스템이 로컬로 마운트됩니다.

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

lofi 장치가 분리되었는지 확인합니다.

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

장치가 분리되면 명령에서 출력을 반환하지 않습니다.

주 - `-m` 옵션을 사용하여 파일 시스템을 마운트하면 `mnttab` 파일에 항목이 추가되지 않습니다. `umount` 명령에서 다음과 같은 경고를 보고할 수 있습니다.

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

이 경고는 무시해도 됩니다.

4. `/.globaldevices` 파일을 삭제하고 필요한 크기로 다시 만듭니다.

다음 예에서는 200MB 크기의 새 `/.globaldevices` 파일 만들기를 보여 줍니다.

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

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

5. 전역 장치 이름 공간에 대한 새 파일 시스템을 만듭니다.

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

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

6. 클러스터 모드로 노드를 부트합니다.

이제 전역 장치가 새 파일 시스템에 채워집니다.

```
phys-schost# reboot
```

7. 해당 노드에서 실행하려는 모든 서비스를 노드로 마이그레이션합니다.

전역 장치 이름 공간 마이그레이션

전용 분할 영역에 전역 장치 이름 공간을 만들지 않고 lofi(loopback file interface) 장치에 이름 공간을 만들 수 있습니다.

주 - 루트 파일 시스템에 대한 ZFS가 지원되지만 한 가지 중요한 예외사항이 있습니다. 전역 장치 파일 시스템에 대해 부트 디스크의 전용 분할 영역을 사용하는 경우 파일 시스템으로 UFS만 사용해야 합니다. 전역 장치 이름 공간에는 UFS 파일 시스템에서 실행되는 PxFS (Proxy File System)가 필요합니다.

그러나 전역 장치 이름 공간에 대한 UFS 파일 시스템은 루트(/) 파일 시스템 및 다른 루트 파일 시스템(예: /var 또는 /home)에 대한 ZFS 파일 시스템과 공존할 수 있습니다. 또는 lofi 장치를 대신 사용하여 전역 장치 이름 공간을 호스팅하는 경우 루트 파일 시스템에 대한 ZFS 사용에 제한이 없습니다.

다음 절차에서는 기존의 전역 장치 이름 공간을 전용 파티션에서 lofi 장치로 이동하거나 그 반대로 이동하는 방법에 대해 설명합니다.

- [전용 파티션에서 lofi 장치로 전역 장치 이름 공간 마이그레이션 방법 \[117\]](#)
- [lofi 장치에서 전용 파티션으로 전역 장치 이름 공간 마이그레이션 방법 \[118\]](#)

▼ 전용 파티션에서 lofi 장치로 전역 장치 이름 공간 마이그레이션 방법

1. 이름 공간 위치를 변경하려는 전역 클러스터 노드에서 root 역할로 전환합니다.
2. 노드에서 서비스를 제외하고 노드를 비클러스터 모드로 재부트합니다.
이 절차를 수행하는 동안 이 노드에서 전역 장치가 처리되지 않도록 하기 위해 이 작업을 수행합니다. 지침은 [비클러스터 모드로 노드를 부트하는 방법 \[90\]](#)을 참조하십시오.
3. /.globaldevices라는 파일이 노드에 있는지 확인합니다.
파일이 있으면 삭제합니다.
4. lofi 장치를 만듭니다.

```
# 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. /etc/vfstab 파일에서 전역 장치 이름 공간 항목을 주석 처리합니다.
이 항목에는 /global/.devices/node@nodeID로 시작하는 마운트 경로가 있습니다.
6. 전역 장치 파티션 /global/.devices/node@nodeID를 마운트 해제합니다.

7. **globaldevices** 및 **scmountdev** SMF 서비스를 사용 안함으로 설정했다가 다시 사용으로 설정합니다.

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

이제 **lofi** 장치가 **/.globaldevices**에 만들어져 전역 장치 파일 시스템으로 마운트됩니다.

8. 전역 장치 이름 공간을 파티션에서 **lofi** 장치로 마이그레이션할 다른 노드에 대해 이러한 단계를 반복합니다.

9. 한 노드에서 전역 장치 이름 공간을 채웁니다.

```
# cldevice populate
```

클러스터에 대해 향후 작업을 수행하기 전에 각 노드에서 명령 처리가 완료되었는지 확인합니다.

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

이제 전역 장치 이름 공간이 **lofi** 장치에 상주합니다.

10. 해당 노드에서 실행하려는 모든 서비스를 노드로 마이그레이션합니다.

▼ **lofi** 장치에서 전용 파티션으로 전역 장치 이름 공간 마이그레이션 방법

1. 이름 공간 위치를 변경하려는 전역 클러스터 노드에서 **root** 역할로 전환합니다.
2. 노드에서 서비스를 제외하고 노드를 비클러스터 모드로 재부트합니다.
이 절차를 수행하는 동안 이 노드에서 전역 장치가 처리되지 않도록 하기 위해 이 작업을 수행합니다. 지침은 [비클러스터 모드로 노드를 부트하는 방법 \[90\]](#)을 참조하십시오.
3. 노드의 로컬 디스크에서 다음 요구사항을 충족하는 새 파티션을 만듭니다.
 - 크기 512MB 이상
 - UFS 파일 시스템 사용
4. 전역 장치 파일 시스템으로 마운트할 새 파티션에 대한 항목을 **/etc/vfstab** 파일에 추가합니다.

- 현재 노드의 노드 ID를 확인합니다.

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

- 다음 형식을 사용하여 `/etc/vfstab` 파일에서 새 항목을 만듭니다.

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

예를 들어 사용하도록 선택한 분할 영역이 `/dev/did/rdisk/d5s3`인 경우 `/etc/vfstab` 파일에 추가할 새 항목은 다음과 같습니다.

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

- 전역 장치 파티션 `/global/.devices/node@nodeID`를 마운트 해제합니다.
- `/.globaldevices` 파일과 관련된 `lofi` 장치를 제거합니다.

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

```
# rm /.globaldevices
```
- `globaldevices` 및 `scmountdev` SMF 서비스를 사용 안함으로 설정했다가 다시 사용으로 설정합니다.

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

이제 파티션이 전역 장치 이름 공간 파일 시스템으로 마운트됩니다.

- 전역 장치 이름 공간을 `lofi` 장치에서 파티션으로 마이그레이션할 다른 노드에 대해 이러한 단계를 반복합니다.
- 클러스터 모드로 부트하고 전역 장치 이름 공간을 채웁니다.
 - 클러스터의 한 노드에서 전역 장치 이름 공간을 채웁니다.

```
# cldevice populate
```

- 노드에 대해 향후 작업을 수행하기 전에 클러스터의 모든 노드에서 프로세스가 완료되었는지 확인합니다.

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

이제 전역 장치 이름 공간이 전용 파티션에 상주합니다.

- 해당 노드에서 실행하려는 모든 서비스를 노드로 마이그레이션합니다.

장치 그룹 추가 및 등록

Solaris Volume Manager, ZFS 또는 원시 디스크에 대한 장치 그룹을 추가하고 등록할 수 있습니다.

▼ 장치 그룹 추가 및 등록 방법(Solaris Volume Manager)

`metaset` 명령을 사용하여 Solaris Volume Manager 디스크 세트를 만들고 해당 디스크 세트를 Oracle Solaris Cluster 장치 그룹으로 등록합니다. 디스크 세트를 등록하면 디스크 세트에 지정한 이름이 자동으로 장치 그룹에 할당됩니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.



주의 - 사용자가 만든 Oracle Solaris Cluster 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

1. 디스크 세트를 만들려는 디스크에 연결된 노드 중 하나에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. Solaris Volume Manager 디스크 세트를 추가하고 Oracle Solaris Cluster가 있는 장치 그룹으로 등록합니다.
복수 소유자 디스크 그룹을 만들려면 `-M` 옵션을 사용합니다.

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

`-s diskset` 만들 디스크 세트를 지정합니다.

`-a -h nodelist` 디스크 세트를 마스터할 수 있는 노드 목록을 추가합니다.

`-M` 디스크 그룹의 소유자를 여러 명으로 지정합니다.

주 - `metaset` 명령을 실행하여 클러스터에 Solaris Volume Manager 장치 그룹을 설치하면 해당 장치 그룹에 포함된 노드 수에 관계없이 기본적으로 보조 노드 수가 하나가 됩니다. 장치 그룹이 만들어진 후 `clsetup` 유틸리티를 사용하여 원하는 보조 노드 수를 변경할 수 있습니다. 디스크 장애 조치에 대한 자세한 내용은 [장치 그룹에 원하는 보조 수를 설정하는 방법 \[132\]](#)을 참조하십시오.

3. 복제된 디스크 그룹을 구성하는 경우, 장치 그룹에 대한 복제 등록 정보를 설정합니다.

```
# cldevicegroup sync devicegroup
```

4. 장치 그룹이 추가되었는지 확인합니다.

장치 그룹 이름은 metaset로 지정한 디스크 세트 이름과 일치합니다.

```
# cldevicegroup list
```

5. DID 매핑을 나열하십시오.

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

- 디스크 세트를 마스터하거나 마스터할 수도 있는 클러스터 노드가 공유하는 드라이브를 선택하십시오.
- 디스크 세트에 드라이브를 추가할 때 /dev/did/rdisk/dN 형식의 전체 DID 장치 이름을 사용합니다.

다음 예에서 DID 장치 /dev/did/rdisk/d3에 대한 항목은 드라이브가 phys-schost-1 및 phys-schost-2에 의해 공유됨을 나타냅니다.

```
=== 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. 디스크 세트에 드라이브를 추가합니다.

전체 DID 경로 이름을 사용합니다.

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

-s setname 디스크 세트 이름을 장치 그룹 이름과 동일하게 지정합니다.

-a 디스크 세트에 드라이브를 추가합니다.

주 - 디스크 세트에 드라이브를 추가할 때 하위 레벨 장치 이름(cNtXdY)을 사용하지 마십시오. 하위 수준 장치 이름은 로컬 이름이므로 클러스터 전체에서 고유하지 않기 때문에 이 이름을 사용하면 메타 세트가 전환되지 않을 수도 있습니다.

7. 디스크 세트와 드라이브의 상태를 확인합니다.

```
# metaset -s setname
```

예 37 Solaris Volume Manager 장치 그룹 추가

다음 예에서는 /dev/did/rdisk/d1 및 /dev/did/rdisk/d2 디스크 드라이브가 포함된 디스크 세트와 장치 그룹을 만드는 방법을 보여주고 장치 그룹이 생성되었는지 확인합니다.

```
# 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
```

▼ 장치 그룹 추가 및 등록 방법(원시 디스크)

Oracle Solaris Cluster 소프트웨어는 다른 볼륨 관리자 사용과 함께 원시 디스크 장치 그룹의 사용을 지원합니다. 처음에 Oracle Solaris Cluster를 구성하면 클러스터에 있는 각 원시 장치에 대해 장치 그룹이 자동으로 구성됩니다. Oracle Solaris Cluster 소프트웨어와 함께 사용하기 위해 자동으로 작성된 장치 그룹을 이 절차를 사용하여 재구성합니다.

다음과 같은 이유로 인해 원시 디스크 유형의 새 장치 그룹을 만듭니다.

- 장치 그룹에 둘 이상의 DID를 추가하려고 합니다.
- 장치 그룹의 이름을 변경해야 합니다.
- cldevicegroup 명령의 -v 옵션을 사용하지 않고 장치 그룹 목록을 만들려고 합니다.



주의 - 복제된 장치에 장치 그룹을 만드는 경우 만든 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

1. 사용하려는 장치를 식별하고 사전 정의된 장치 그룹의 구성을 해제합니다.

다음 명령은 dN 및 dX 장치에 대해 사전 정의된 장치 그룹을 제거합니다.

```
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. 원하는 장치를 포함하는 새 원시 디스크 장치 그룹을 작성합니다.

다음 명령은 dN 및 dX 장치를 포함하는 전역 장치 그룹 *raw-disk-dg*를 만듭니다.

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

3. 사용자가 만든 원시 디스크 장치 그룹을 확인합니다.

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

▼ 복제된 장치 그룹 추가 및 등록 방법(ZFS)

이 절차에 따라 HAStoragePlus에서 관리하는 복제된 ZFS 장치 그룹을 만듭니다.

HAStoragePlus를 사용하지 않는 ZFS 저장소 풀(zpool)을 만들려면 대신 [HAStoragePlus 없이 로컬 ZFS 저장소 풀을 구성하는 방법 \[124\]](#)으로 이동하십시오.

시작하기 전에 ZFS를 복제하려면 이름이 지정된 장치 그룹을 만들고 zpool에 속하는 디스크를 나열해야 합니다. 장치는 한 번에 하나의 장치 그룹에만 속할 수 있으므로 장치가 포함된 Oracle Solaris Cluster 장치 그룹이 이미 있는 경우에는 새 ZFS 장치 그룹에 해당 장치를 추가하기 전에 그룹을 삭제해야 합니다.

사용자가 만드는 Oracle Solaris Cluster 장치 그룹(Solaris Volume Manager 또는 원시 디스크)의 이름은 복제된 장치 그룹의 이름과 같아야 합니다.

1. zpool의 장치에 해당하는 기본 장치 그룹을 삭제합니다.

예를 들어 /dev/did/dsk/d2 및 /dev/did/dsk/d13 장치 두 개를 포함하는 mypool이라는 zpool이 있는 경우 d2 및 d13이라는 기본 장치 그룹 두 개를 삭제해야 합니다.

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

2. [1단계](#)에서 제거한 장치 그룹의 DID에 해당하는 DID를 가진 명명된 장치 그룹을 만듭니다.

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

이 작업을 수행하면 mypool(zpool과 같은 이름)이라는 장치 그룹이 만들어져 /dev/did/dsk/d2 및 /dev/did/dsk/d13 원시 장치를 관리합니다.

3. 해당 장치를 포함하는 zpool을 만듭니다.

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

4. 리소스 그룹을 만들어 노드 목록에 전역 영역만 있는 복제된 장치(장치 그룹에 포함)의 마이그레이션을 관리합니다.

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

5. [4단계](#)에서 만든 리소스 그룹에 hasp-rs 리소스를 만들고 globaldevicepaths 등록 정보를 원시 디스크 유형의 장치 그룹으로 설정합니다.

이 장치는 [2단계](#)에서 만들었습니다.

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

6. 이 리소스 그룹에서 rg_affinities 등록 정보의 +++ 값을 [4단계](#)에서 만든 리소스 그룹으로 설정합니다.

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

7. 3단계에서 만든 zpool에 대한 HAStoragePlus 리소스(hasp-rs)를 4단계 또는 6단계에서 만든 리소스 그룹에 만듭니다.

resource_dependencies 등록 정보를 5단계에서 만든 hasp-rs 리소스로 설정합니다.

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

8. 장치 그룹 이름이 필요한 위치에 새 리소스 그룹 이름을 사용합니다.

▼ HAStoragePlus 없이 로컬 ZFS 저장소 풀을 구성하는 방법

이 절차에서는 HAStoragePlus 리소스를 구성하지 않고 로컬 장치에서 ZFS 저장소 풀(zpool)을 구성하는 방법에 대해 설명합니다.

주 - HAStoragePlus 리소스를 사용하는 로컬 zpool을 구성하려면 대신 [복제된 장치 그룹 추가 및 등록 방법\(ZFS\) \[123\]](#)으로 이동하십시오.

1. DID 매핑을 나열하고 사용할 로컬 장치를 식별합니다.

새 zpool을 사용할 클러스터 노드만 나열하는 장치를 선택합니다. cNtXdY 장치 이름과 /dev/did/rdisk/dN DID 장치 이름을 모두 확인하십시오.

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

다음 예제에서 DID 장치 /dev/did/rdisk/d1 및 /dev/did/rdisk/d2에 대한 항목은 해당 드라이브가 phys-schost-1에서만 사용됨을 보여 줍니다. 이 절차의 예제에서는 장치 이름이 c0t6d0인 DID 장치 /dev/did/rdisk/d2가 클러스터 노드 phys-schost-1에 대해 사용 및 구성됩니다.

```
=== 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
...
```

2. zpool에 대해 선택하는 DID 장치의 장치 그룹 이름을 확인합니다.

다음 예제 출력은 dsk/d2가 DID 장치 /dev/did/rdsk/d2s2에 대한 장치 그룹 이름임을 보여줍니다. 장치 그룹 이름이 DID 장치 이름의 일부인 관계가 자주 발생하지만 항상 그런 것은 아닙니다. 이 장치 그룹에는 해당 노드 목록에 phys-schost-1 노드 하나만 포함됩니다.

```
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/rdsk/d2s2
...
```

3. DID 장치에 대해 localonly 등록 정보를 설정합니다.

2단계에서 식별한 장치 그룹 이름을 지정합니다. 장치에 대한 보호를 사용 안함으로 설정하려는 경우 명령에 default_fencing=nofencing도 포함시킵니다.

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

cldevicegroup 등록 정보에 대한 자세한 내용은 [cldevicegroup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

4. 장치 설정을 확인합니다.

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

5. zpool을 만듭니다.

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

6. (옵션) ZFS 데이터 세트를 만듭니다.

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

7. 새 zpool을 확인합니다.

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

장치 그룹 유지 보수

장치 그룹에 대해 다양한 관리 작업을 수행할 수 있습니다. 또한 이러한 작업 중 일부는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 수행할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

장치 그룹 제거 및 등록 해제 방법(Solaris Volume Manager)

장치 그룹은 Oracle Solaris Cluster에 등록된 Solaris Volume Manager 디스크 세트입니다. Solaris Volume Manager 장치 그룹을 제거하려면 `metaclear` 및 `metaset` 명령을 사용합니다. 이 명령은 동일한 이름의 장치 그룹을 제거하고 Oracle Solaris Cluster 장치 그룹에서 디스크 그룹의 등록을 해제합니다.

디스크 세트 제거 단계는 [Solaris Volume Manager Administration Guide](#)를 참조하십시오.

▼ 모든 장치 그룹에서 노드를 제거하는 방법

잠재적 기본값에 해당 노드를 나열하는 모든 장치 그룹에서 클러스터 노드를 제거하려면 이 절차를 사용합니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 제거할 노드(모든 장치 그룹의 잠재적 기본 노드)에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 장치 그룹 또는 제거할 노드의 그룹이 구성원인지 판별합니다.

각 장치 그룹에 대한 `Device group node list`에서 해당 노드 이름을 찾습니다.

```
# cldevicegroup list -v
```

3. 2단계에서 식별된 장치 그룹이 `svm` 장치 그룹 유형인 경우 해당 유형의 각 장치 그룹에 대해 [장치 그룹에서 노드를 제거하는 방법\(Solaris Volume Manager\) \[127\]](#)의 단계를 수행합니다.

4. 제거할 노드가 속한 원시 디스크 장치 그룹이 구성원인지 확인합니다.

```
# cldevicegroup list -v
```

5. 4단계에 나열된 장치 그룹이 `Disk` 또는 `Local_Disk` 장치 그룹 유형인 경우 각 장치 그룹에 대해 [원시 디스크 장치 그룹에서 노드를 제거하는 방법 \[129\]](#)의 단계를 수행합니다.

6. 모든 장치 그룹의 잠재적인 기본 노드 목록에서 노드가 제거되었는지 확인합니다.

해당 노드가 장치 그룹의 잠재적인 기본 노드로서 목록에 포함되어 있지 않으면 명령을 실행해도 아무것도 반환되지 않습니다.

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

▼ 장치 그룹에서 노드를 제거하는 방법(Solaris Volume Manager)

Solaris Volume Manager 장치 그룹의 잠재적인 기본값 목록에서 클러스터 노드를 제거하려면 이 절차를 따릅니다. 제거할 노드가 있는 각 장치 그룹에 대해 `metaset` 명령을 반복합니다.



주의 - 다른 노드가 활성 클러스터 멤버이고 그 중 하나 이상이 디스크 세트를 소유하고 있는 경우에는 클러스터 외부에서 부트되는 클러스터 노드에서 `metaset -s setname -f -t` 명령을 실행하지 마십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 노드가 여전히 장치 그룹의 멤버인지와 장치 그룹이 Solaris Volume Manager 장치 그룹인지를 확인합니다.

장치 그룹 유형 `sds/svm`은 Solaris Volume Manager 장치 그룹을 나타냅니다.

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

2. 어느 노드가 현재 장치 그룹의 기본 노드인지 확인합니다.

```
# cldevicegroup status devicegroup
```

3. 현재 수정할 장치 그룹이 있는 노드에서 `root` 역할로 전환합니다.

4. 장치 그룹에서 노드의 호스트 이름을 삭제합니다.

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

`-s setname` 장치 그룹 이름을 지정합니다.

`-d` `-h`를 사용하여 확인한 노드를 장치 그룹에서 삭제합니다.

`-h nodelist` 제거할 노드의 노드 이름을 지정합니다.

주 - 업데이트를 완료하는 데 몇 분이 걸릴 수 있습니다.

명령이 실패하면 명령에 `-f`(강제 실행) 옵션을 추가합니다.

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

5. 잠재적인 기본값으로서 노드가 제거되는 각 장치 그룹에 대해 4단계를 반복합니다.

6. 노드가 장치 그룹에서 제거되었는지 확인합니다.

장치 그룹 이름은 metaset로 지정한 디스크 세트 이름과 일치합니다.

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

예 38 장치 그룹에서 노드 제거(Solaris Volume Manager)

다음 예에서는 장치 그룹 구성에서 호스트 이름 phys-schost-2를 제거하는 방법을 보여 줍니다. 이 예에서는 지정된 장치 그룹의 잠재적인 기본 노드인 phys-schost-2를 제거합니다. cldevicegroup show 명령을 실행하여 노드가 제거되었는지 확인합니다. 제거된 노드가 더 이상 화면의 텍스트에 표시되지 않는지 확인하십시오.

노드에 대한 Solaris Volume Manager 장치 그룹 확인

```
# 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
```

어느 노드가 현재 장치 그룹의 기본 노드인지 확인합니다

```
# 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

현재 장치 그룹을 소유하는 노드에서 루트 역할로 전환

장치 그룹에서 호스트 이름 제거

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

노드 제거 확인

```
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

dg-schost-1

phys-schost-1 -

Online

▼ 원시 디스크 장치 그룹에서 노드를 제거하는 방법

원시 디스크 장치 그룹의 잠재적인 기본 노드 목록에서 클러스터 노드를 제거하려면 다음 절차를 수행합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 제거할 노드가 아닌 클러스터의 한 노드에서 `solaris.cluster.read` 및 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 제거할 노드와 연결된 장치 그룹을 식별하고, 원시 디스크 장치 그룹을 판별합니다.

```
# cldevicegroup show -n nodename -t rawdisk +
```

3. 각 Local_Disk 원시 디스크 장치 그룹의 `localonly` 등록 정보를 사용 불가능하게 합니다.

```
# cldevicegroup set -p localonly=false devicegroup
```

`localonly` 등록 정보에 대한 자세한 내용은 `cldevicegroup(1CL)` 매뉴얼 페이지를 참조하십시오.

4. 제거할 노드에 연결된 모든 원시 디스크 장치 그룹의 `localonly` 등록 정보를 사용 불가능하게 했는지 확인합니다.

Disk 장치 그룹 유형은 해당 원시 디스크 장치 그룹에 대해 `localonly` 등록 정보가 사용 불가능하게 되었음을 나타냅니다.

```
# cldevicegroup show -n nodename -t rawdisk -v +
```

5. 2단계에서 식별된 모든 원시 디스크 장치 그룹에서 노드를 제거합니다.

제거할 노드가 연결된 각 원시 디스크 장치 그룹에 대하여 이 단계를 완료해야 합니다.

```
# cldevicegroup remove-node -n nodename devicegroup
```

예 39 원시 장치 그룹에서 노드 제거

이 예에서는 원시 디스크 장치 그룹에서 노드(phys-schost-2)를 제거하는 방법을 보여 줍니다. 모든 명령이 클러스터의 다른 노드(phys-schost-1)에서 실행됩니다.

제거하려는 노드에 연결된 장치 그룹 식별 및

원시 디스크 장치 그룹 확인

```
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
노드의 각 로컬 디스크에 대해 localonly 플래그를 사용 안함으로 설정
phys-schost-1# cldevicegroup set -p localonly=false dsk/d4
```

localonly 플래그가 사용 안함으로 설정되었는지 확인

```
phys-schost-1# cldevicegroup show -n
phys-schost-2 -t rawdisk +
(dsk/d4) Device group type:      Disk
(dsk/d8) Device group type:      Local_Disk
```

모든 원시 디스크 장치 그룹에서 노드 제거

```
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
```

▼ 장치 그룹 등록 정보를 변경하는 방법

장치 그룹의 기본 소유권을 설정하는 방법은 *preferenced*라는 소유권 기본 설정 속성의 설정을 기반으로 합니다. 이 속성이 설정되지 않은 경우에는 다른 노드가 소유하지 않은 장치 그룹의 디스크에 처음으로 액세스를 시도하는 노드가 해당 그룹을 소유하게 됩니다. 그러나 이 속성이 설정되면 노드가 소유권을 얻기 위해 시도하는 순서를 지정해야 합니다.

preferenced 속성을 사용 불가하게 하면 *failback* 속성도 자동으로 사용 불가하게 됩니다. 그러나 *preferenced* 속성을 사용 가능 또는 다시 사용 가능하게 하려는 경우 *failback* 속성을 사용 가능하게 하거나 사용 불가하게 할 수 있습니다.

preferenced 속성이 사용 가능 또는 다시 사용 가능한 경우 기본 소유권 설정 목록에서 노드 순서를 재설정해야 합니다.

이 절차에서는 `clsetup` 유틸리티를 사용하여 Solaris Volume Manager 장치 그룹에 대한 `preferenced` 속성 및 `failback` 속성을 설정 또는 설정 해제합니다.

시작하기 전에 이 절차를 수행하려면 속성 값을 변경할 장치 그룹의 이름이 필요합니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.read` 및 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. `clsetup` 유틸리티를 시작합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.

3. 장치 그룹 작업을 하려면 장치 그룹 및 볼륨 옵션에 대한 번호를 입력합니다.

장치 그룹 메뉴가 표시됩니다.

4. 장치 그룹의 키 등록 정보를 변경하려면 Solaris Volume Manager 장치 그룹의 키 등록 정보 변경 옵션에 대한 번호를 입력합니다.

주요 등록 정보 변경 메뉴가 표시됩니다.

5. 장치 그룹 등록 정보를 변경하려면 기본 설정 또는 페일백 등록 정보를 변경할 옵션 번호를 입력합니다.

지침에 따라 장치 그룹에 대한 `preferenced` 및 `failback` 옵션을 설정합니다.

6. 장치 그룹 속성이 변경되었는지 확인하십시오.

다음 명령을 실행하여 장치 그룹 정보가 표시되는지 확인합니다.

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

예 40 장치 그룹의 등록 정보 변경

다음 예에서는 장치 그룹(`dg-schost-1`)에 대한 속성 값을 설정할 때 `clsetup`에 의해 생성되는 `cldevicegroup` 명령을 보여 줍니다.

```
# cldevicegroup set -p preferred=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

```

▼ 장치 그룹에 원하는 보조 수를 설정하는 방법

`numsecondaries` 등록 정보는 기본 노드가 실패할 경우 그룹을 마스터할 수 있는 장치 그룹 내 노드 수를 지정합니다. 장치 서비스에 대한 기본 보조 노드 수는 1입니다. 값은 1부터 장치 그룹에서 작동하는 기본 이외의 공급자 노드 수까지, 정수로 설정할 수 있습니다.

이 설정은 클러스터의 성능과 가용성 사이에 균형을 맞추는 데 중요한 역할을 하는 값입니다. 예를 들어, 보조 노드 수를 증가시키면 클러스터에서 동시에 여러 번 장애가 발생할 경우에도 장치 그룹이 작동할 확률이 높아집니다. 또한 보조 노드 수를 늘리면 정상 작동 중에 주기적으로 성능이 저하됩니다. 일반적으로 보조 노드 수가 적을수록 성능은 좋아지지만 가용성은 떨어집니다. 그러나 보조 노드 수가 많다고 해서 문제가 발생하는 파일 시스템이나 장치 그룹의 가용성이 항상 높아지는 것은 아닙니다. 자세한 내용은 [Oracle Solaris Cluster 4.3 Concepts Guide](#)의 3장, “Key Concepts for System Administrators and Application Developers”을 참조하십시오.

`numsecondaries` 등록 정보를 변경하면 실제 보조 노드 수와 원하는 개수가 맞지 않을 경우에 보조 노드가 장치 그룹에 추가되거나 장치 그룹에서 제거됩니다.

이 절차에서는 장치 그룹의 모든 유형에 대해 `numsecondaries` 등록 정보를 설정하기 위해 `clsetup` 유틸리티를 사용합니다. 장치 그룹을 구성할 때의 장치 그룹 옵션에 대한 자세한 내용은 [cldevicegroup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.read` 및 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. `clsetup` 유틸리티를 시작합니다.

`clsetup`

주 메뉴가 표시됩니다.
3. 장치 그룹에서 작업하려면 장치 그룹 및 볼륨 메뉴 항목을 선택합니다.
장치 그룹 메뉴가 표시됩니다.

4. 장치 그룹의 키 등록 정보를 변경하려면 장치 그룹의 키 등록 정보 변경 메뉴 항목을 선택합니다.

주요 등록 정보 변경 메뉴가 표시됩니다.

5. 보조 노드 수를 변경하려면 `numsecondaries` 등록 정보를 변경하는 옵션에 대한 번호를 입력합니다.

지침에 따라 장치 그룹에 대해 구성할 원하는 보조 노드 수를 입력합니다. 그러면 해당하는 `cldevicegroup` 명령이 실행되고, 로그가 인쇄되며, 유틸리티가 이전 메뉴로 돌아갑니다.

6. 장치 그룹 구성을 검증합니다.

```
# 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
```

주 - 장치 그룹 구성 변경에는 기존 볼륨의 그룹, 소유자, 권한 변경뿐 아니라 볼륨 추가나 제거도 포함됩니다. 구성을 변경한 후에 다시 등록하면 전역 이름 공간이 올바른 상태가 됩니다. [전역 장치 이름 공간 업데이트 방법 \[115\]](#)을 참조하십시오.

7. 장치 그룹 속성이 변경되었는지 확인합니다.

다음 명령을 실행하여 표시되는 장치 그룹 정보를 확인합니다.

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

예 41 필요한 보조 노드 수 변경(Solaris Volume Manager)

다음 예에서는 장치 그룹(`dg-schost-1`)에 대한 보조 노드 수를 구성할 때 `clsetup`에 의해 생성되는 `cldevicegroup` 명령을 보여 줍니다. 이 예에서는 이전에 디스크 그룹 및 볼륨을 만들었다고 가정합니다.

```
# 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
```

예 42 원하는 보조 노드의 수를 기본값으로 설정

다음은 null 문자열 값을 사용하여 보조 노드의 기본 개수를 구성하는 예입니다. 기본값이 변경될 경우에도 장치 그룹이 기본값을 사용하도록 구성됩니다.

```
# 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
```

▼ 장치 그룹 구성을 나열하는 방법

구성을 나열하기 위해 root 역할로 전환할 필요는 없습니다. 그러나, `solaris.cluster.read` 인증이 필요합니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

● 다음 방법 중 하나를 사용합니다.

Oracle Solaris Cluster Manager 브라우저 인터페이스

자세한 내용은 [13장. Oracle Solaris Cluster Manager 브라우저 인터페이스 사용](#)을 참조하십시오.

`cldevicegroup show`

클러스터의 모든 장치 그룹에 대한 구성을 나열하려면 `cldevicegroup show` 명령을 사용합니다.

`cldevicegroup show devicegroup`

단일 장치 그룹의 구성을 나열하려면 `cldevicegroup show devicegroup` 명령을 사용합니다.

`cldevicegroup status devicegroup`

단일 장치 그룹의 상태를 판별하려면 `cldevicegroup status devicegroup` 명령을 사용합니다.

```
cldevicegroup status +
```

클러스터의 모든 장치 그룹의 상태를 판별하려면 `cldevicegroup status +` 명령을 사용합니다.

자세한 내용을 보려면 이러한 명령과 함께 `-v` 옵션을 사용합니다.

예 43 모든 장치 그룹의 상태 표시

```
# 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
```

예 44 특정 장치 그룹의 구성 표시

```
# 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
```

▼ 장치 그룹에 대한 기본 노드를 전환하는 방법

다음 절차를 수행하면 비활성 장치 그룹을 시작(온라인으로 전환)할 수도 있습니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 비활성 장치 그룹을 온라인으로 전환할 수도 있습니다. 자세한 내용은 Oracle Solaris Cluster Manager 온라인 도움말을 참조하십시오. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문 형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. `cldevicegroup switch`를 사용하여 기본 장치 그룹을 전환합니다.

```
# cldevicegroup switch -n nodename devicegroup
```

`-n nodename` 전환할 대상 노드의 이름을 지정합니다. 이 노드가 새 기본 노드가 됩니다.

`devicegroup` 전환할 장치 그룹을 지정합니다.

3. 장치 그룹이 새로운 기본 노드로 전환되었는지 확인합니다.

장치 그룹이 올바르게 등록되면 다음 명령을 사용할 때 새 장치 그룹에 대한 정보가 표시됩니다.

```
# cldevice status devicegroup
```

예 45 장치 그룹에 대한 기본 노드 전환

다음 예는 장치 그룹에 대한 기본 노드를 전환하는 방법과 변경을 확인하는 방법입니다.

```
# 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

▼ 장치 그룹을 유지 보수 상태로 전환하는 방법

장치 그룹을 유지 보수 상태로 전환하면 장치 중 하나에 액세스할 때 해당 장치 그룹이 자동으로 온라인으로 전환되지 않습니다. 복구 절차를 완료하기 위해 모든 I/O 작업을 중단해야 하는 경우에는 복구가 완료될 때까지 장치 그룹을 유지 보수 상태로 유지해야 합니다. 장치 그룹을 유지 보수 상태로 전환하면 특정 노드에서 디스크 세트가 복구되는 동안 다른 노드에서 장치 그룹이 온라인으로 전환되지 않으므로 데이터 손실이 방지됩니다.

손상된 디스크 세트 복원 방법에 대한 지침은 “[손상된 디스크 세트 복원](#)” [255]을 참조하십시오.

주 - 디스크 그룹을 유지 보수 상태로 만들려면 먼저 장치에 대한 모든 액세스를 중단하고 관련 파일 시스템의 마운트를 모두 해제해야 합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 활성 장치 그룹을 오프라인으로 전환할 수도 있습니다. 자세한 내용은 Oracle Solaris Cluster Manager 온라인 도움말을 참조하십시오. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법](#) [278]을 참조하십시오.

1. 장치 그룹을 유지 보수 상태로 만듭니다.

- a. 장치 그룹이 활성화되어 있으면 비활성화합니다.

```
# cldevicegroup disable devicegroup
```

- b. 장치 그룹을 오프라인으로 전환합니다.

```
# cldevicegroup offline devicegroup
```

2. 복구 절차를 수행하기 위해 디스크 세트에 대한 소유권이 필요한 경우 수동으로 해당 디스크 세트를 가져옵니다.

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



주의 - Solaris Volume Manager 디스크 세트에 대한 소유권을 얻는 경우 장치 그룹이 유지 보수 상태에 있을 때 반드시 `metaset -C take` 명령을 사용해야 합니다. `metaset -t` 명령을 사용하면 소유권을 얻는 과정에서 장치 그룹이 온라인 상태로 전환됩니다.

3. 수행해야 할 복구 절차를 완료합니다.

4. 디스크 세트에 대한 소유권을 해제합니다.



주의 - 장치 그룹을 유지 보수 상태에서 해제하기 전에 디스크 세트에 대한 소유권을 해제해야 합니다. 소유권 해제가 실패하면 데이터 손실이 일어날 수 있습니다.

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

5. 장치 그룹을 온라인으로 전환합니다.

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

예 46 장치 그룹을 유지 보수 상태로 만들기

이 예에서는 장치 그룹 dg-schost-1을 유지 보수 상태로 전환하고, 유지 보수 상태에서 장치 그룹을 제거하는 방법을 보여 줍니다.

```
[장치 그룹을 유지 보수 상태로 만듭니다.]
# cldevicegroup disable dg-schost-1
# cldevicegroup offline dg-schost-1
[필요한 경우 수동으로 디스크 세트를 가져옵니다.]
For Solaris Volume Manager:
# metaset -C take -f -s dg-schost-1
[필요한 모든 복구 절차를 완료합니다.]
[소유권을 해제합니다.]
Solaris Volume Manager:
# metaset -C release -s dg-schost-1
[장치 그룹을 온라인으로 전환합니다.]
# cldevicegroup online dg-schost-1
# cldevicegroup enable dg-schost-1
```

저장 장치에 대한 SCSI 프로토콜 설정 관리

Oracle Solaris Cluster 소프트웨어를 설치하면 모든 저장 장치에 SCSI 예약이 자동으로 할당됩니다. 다음 절차에 따라 장치 설정을 확인하고, 필요한 경우 장치 설정을 대체할 수 있습니다.

- 모든 저장 장치에 대한 기본 전역 SCSI 프로토콜 설정을 표시하는 방법 [138]
- 단일 저장 장치의 SCSI 프로토콜을 표시하는 방법 [139]
- 모든 저장 장치에 대한 기본 전역 보호(fencing) 프로토콜 설정을 변경하는 방법 [140]
- 단일 저장 장치에 대한 보호(fencing) 프로토콜을 변경하는 방법 [141]

▼ 모든 저장 장치에 대한 기본 전역 SCSI 프로토콜 설정을 표시하는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 모든 노드에 현재 기본 전역 SCSI 프로토콜 설정을 표시합니다.

```
# cluster show -t global
```

자세한 내용은 `cluster(1CL)` 매뉴얼 페이지를 참조하십시오.

예 47 모든 저장 장치에 대한 기본 전역 SCSI 프로토콜 설정 표시

다음 예에서는 클러스터의 모든 저장 장치에 대한 SCSI 프로토콜 설정을 표시합니다.

```
# 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
```

▼ 단일 저장 장치의 SCSI 프로토콜을 표시하는 방법

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 모든 노드에 저장 장치의 SCSI 프로토콜을 표시합니다.

```
# cldevice show device
```

`device` 장치 경로 이름 또는 장치 이름.

자세한 내용은 `cldevice(1CL)` 매뉴얼 페이지를 참조하십시오.

예 48 단일 장치의 SCSI 프로토콜 표시

다음 예에서는 장치 /dev/rdisk/c4t8d0에 대한 SCSI 프로토콜을 표시합니다.

```
# cldevice show /dev/rdisk/c4t8d0

=== DID Device Instances ===

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

▼ 모든 저장 장치에 대한 기본 전역 보호(fencing) 프로토콜 설정을 변경하는 방법

클러스터에 연결된 모든 저장 장치에 대해 전역적으로 보호(fencing)를 설정하거나 해제할 수 있습니다. 단일 저장 장치의 기본 보호(fencing) 설정은 장치의 기본 보호(fencing)가 pathcount, prefer3 또는 nofencing으로 설정된 경우 전역 설정으로 대체됩니다. 저장 장치의 기본 보호(fencing) 설정이 global인 경우 저장 장치는 전역 설정을 사용하게 됩니다. 예를 들어, 저장 장치에 기본 설정 pathcount가 있는 경우 이 절차를 사용하여 전역 SCSI 프로토콜 설정을 prefer3으로 변경해도 설정이 변경되지 않습니다. 단일 장치의 기본 설정을 변경하려면 [단일 저장 장치에 대한 보호\(fencing\) 프로토콜을 변경하는 방법 \[141\]](#) 절차를 사용해야 합니다.



주의 - 잘못된 상황에서 보호(fencing)를 해제하면 응용 프로그램 장애 조치 중에 데이터가 손상될 수 있습니다. 보호(fencing)를 해제하려는 경우 이러한 데이터 손상 가능성에 주의해야 합니다. 공유 저장 장치에서 SCSI 프로토콜을 지원하지 않는 경우 또는 클러스터 외부 호스트에서 클러스터의 저장소에 액세스할 수 있게 하려는 경우 보호(fencing)를 해제할 수 있습니다.

쿼럼 장치에 대한 기본 보호(fencing) 설정을 변경하려면 장치의 구성을 해제하고, 보호(fencing) 설정을 변경한 다음 쿼럼 장치를 재구성해야 합니다. 쿼럼 장치가 포함된 장치에 대해 정기적으로 보호(fencing)를 해제하고 다시 설정하려는 경우 쿼럼 작업이 중단되지 않도록 쿼럼 서버 서비스를 통해 쿼럼을 구성하는 것이 좋습니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 퀴럼 장치가 아닌 모든 저장 장치에 대해 보호(fencing) 프로토콜을 설정합니다.

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

```
-p global_fencing
```

모든 공유 장치에 대하여 현재 기본 경계 알고리즘을 설정합니다.

```
prefer3
```

경로가 세 개 이상 있는 장치에 대해 SCSI-3 프로토콜을 사용합니다.

```
pathcount
```

공유 장치에 첨부된 DID 경로의 수로 경계 프로토콜을 판별합니다. pathcount 설정은 퀴럼 장치에 사용됩니다.

```
nofencing
```

모든 저장 장치에 대해 보호(fencing) 상태를 설정하여 보호(fencing)를 해제합니다.

```
nofencing-noscrub
```

장치를 스크러빙하면 장치에서 지속적인 SCSI 예약 정보가 모두 지워지고 클러스터 외부 시스템에서 저장소에 액세스할 수 있습니다. SCSI 예약에 심각한 문제가 있는 저장 장치에 대해서만 nofencing-noscrub 옵션을 사용합니다.

예 49 모든 저장 장치에 대한 기본 전역 보호(fencing) 프로토콜 설정 지정

다음 예에서는 클러스터의 모든 저장 장치에 대한 보호(fencing) 프로토콜을 SCSI-3으로 설정합니다.

```
# cluster set -p global_fencing=prefer3
```

▼ 단일 저장 장치에 대한 보호(fencing) 프로토콜을 변경하는 방법

단일 저장 장치에 대한 보호(fencing) 프로토콜을 변경할 수도 있습니다.

주 - 퀴럼 장치에 대한 기본 보호(fencing) 설정을 변경하려면 장치의 구성을 해제하고, 보호(fencing) 설정을 변경한 다음 퀴럼 장치를 재구성해야 합니다. 퀴럼 장치가 포함된 장치에 대해 정기적으로 보호(fencing)를 해제하고 다시 설정하려는 경우 퀴럼 작업이 중단되지 않도록 퀴럼 서버 서비스를 통해 퀴럼을 구성하는 것이 좋습니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.



주의 - 잘못된 상황에서 보호(fencing)를 해제하면 응용 프로그램 장애 조치 중에 데이터가 손상될 수 있습니다. 보호(fencing)를 해제하려는 경우 이러한 데이터 손상 가능성에 주의해야 합니다. 공유 저장 장치에서 SCSI 프로토콜을 지원하지 않는 경우 또는 클러스터 외부 호스트에서 클러스터의 저장소에 액세스할 수 있게 하려는 경우 보호(fencing)를 해제할 수 있습니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. **저장 장치의 보호(fencing) 프로토콜을 설정합니다.**

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

-p default_fencing

장치의 default_fencing 등록 정보를 수정합니다.

pathcount

공유 장치에 첨부된 DID 경로의 수로 경계 프로토콜을 판별합니다.

scsi3

SCSI-3 프로토콜을 사용합니다.

global

기본 경계 설정으로 전역을 사용합니다. global 설정은 쿼럼이 아닌 장치에 사용됩니다.

nofencing

지정된 DID 인스턴스에 대해 보호(fencing) 상태를 설정하여 보호(fencing)를 해제합니다.

nofencing-noscrub

장치를 스크러빙하면 장치에서 지속적인 SCSI 예약 정보가 모두 지워지고 클러스터 외부 시스템에서 저장 장치에 액세스할 수 있습니다. SCSI 예약에 심각한 문제가 있는 저장 장치에 대해서만 nofencing-noscrub 옵션을 사용합니다.

device

장치 경로의 이름 또는 장치 이름을 지정합니다.

자세한 내용은 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

예 50 단일 장치의 보호(fencing) 프로토콜 설정

다음 예에서는 장치 번호로 지정된 장치 d5를 SCSI-3 프로토콜로 설정합니다.

```
# cldevice set -p default_fencing=prefer3 d5
```

또한 d11 장치에 대해 기본 보호(fencing)를 해제합니다.

```
#cldevice set -p default_fencing=nofencing d11
```

클러스터 파일 시스템 관리

클러스터 파일 시스템은 클러스터의 모든 노드에서 읽고 액세스할 수 있는 전역적으로 사용 가능한 파일 시스템입니다.

표 8 작업 맵: 클러스터 파일 시스템 관리

작업	지침
초기 Oracle Solaris Cluster 설치 후 클러스터 파일 시스템 추가	클러스터 파일 시스템을 추가하는 방법 [143]
클러스터 파일 시스템 제거	클러스터 파일 시스템을 제거하는 방법 [146]
노드 간 일관성을 위해 클러스터의 전역 마운트 지점 확인	클러스터의 전역 마운트를 확인하는 방법 [148]

▼ 클러스터 파일 시스템을 추가하는 방법

처음 Oracle Solaris Cluster를 설치한 후에 만드는 각 클러스터 파일 시스템에 대하여 이 작업을 수행하십시오.



주의 - 정확한 디스크 장치 이름을 지정해야 합니다. 클러스터 파일 시스템을 만들면 디스크에 있는 데이터가 모두 삭제됩니다. 잘못된 장치 이름을 지정하면 지우려고 하지 않은 데이터가 삭제됩니다.

다음의 필수 조건은 추가적인 클러스터 파일 시스템을 추가하기 전에 완료되어야 함을 확인하십시오.

- 클러스터에서 root 역할 권한은 노드에서 설정합니다.
- 클러스터에 볼륨 관리자 소프트웨어를 설치하고 구성합니다.
- 클러스터 파일 시스템을 만들 장치 그룹(예: Solaris Volume Manager 장치 그룹) 또는 블록 디스크 슬라이스가 나타납니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터 파일 시스템을 영역 클러스터에 추가할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

Oracle Solaris Cluster Manager를 사용하여 데이터 서비스를 설치한 경우에 클러스터 파일 시스템을 만들 충분한 공유 디스크가 있었으면 이미 하나 이상의 클러스터 파일 시스템이 있습니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터 노드에서 **root** 역할을 수행합니다.

팁 - 더 빨리 파일 시스템을 만들려면 파일 시스템을 만드는 현재 기본 전역 장치에서 **root** 역할을 수행합니다.

2. **newfs** 명령을 사용하여 UFS 파일 시스템을 만듭니다.



주의 - 파일 시스템을 만들면 디스크의 데이터가 모두 삭제됩니다. 정확한 디스크 장치 이름을 지정해야 합니다. 잘못된 장치 이름을 지정하면 삭제하지 않으려는 데이터가 지워질 수 있습니다.

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

다음 표에서는 *raw-disk-device* 인수 이름의 예를 보여 줍니다. 이름 지정 규칙은 볼륨 관리자마다 다릅니다.

볼륨 관리자	샘플 디스크 장치 이름	설명
Solaris Volume Manager	<code>/dev/md/nfs/rdisk/d1</code>	nfs 디스크 세트에 있는 원시 디스크 장치 d1
없음	<code>/dev/global/rdisk/d1s3</code>	원시 디스크 장치 d1s3

3. 클러스터의 각 노드에서 클러스터 파일 시스템의 마운트 지점 디렉토리를 만듭니다.

해당 노드에서 클러스터 파일 시스템에 액세스하지 않는 경우에도 각 노드에 마운트 지점이 필요합니다.

팁 - 쉽게 관리하려면 `/global/device-group/` 디렉토리에 마운트 지점을 만듭니다. 이 위치를 사용하면 로컬 파일 시스템에서 전역으로 사용 가능한 클러스터 파일 시스템을 쉽게 구별할 수 있습니다.

`phys-schost# mkdir -p /global/device-group/mount-point/`

device-group

장치를 포함하는 장치 그룹의 이름에 해당되는 디렉토리 이름

mount-point

클러스터 파일 시스템을 마운트할 디렉토리의 이름

4. 클러스터의 각 노드에서 `/etc/vfstab` 파일에 마운트 지점에 대한 항목을 추가합니다.
자세한 내용은 `vfstab(4)` 매뉴얼 페이지를 참조하십시오.
 - a. 각 항목에서 사용하는 파일 시스템 유형에 대한 필수 마운트 옵션을 지정합니다.
 - b. 클러스터 파일 시스템을 자동으로 마운트하려면 `mount at boot` 필드를 `yes`로 설정합니다.
 - c. 각 클러스터 파일 시스템에 대해 `/etc/vfstab` 항목의 정보가 각 노드에서 동일한지 확인합니다.
 - d. 각 노드의 `/etc/vfstab` 파일에 있는 장치 항목 순서가 동일한지 확인합니다.
 - e. 파일 시스템의 부트 순서 종속성을 확인하십시오.
예를 들어 `phys-schost-1`은 디스크 장치 `d0`을 `/global/oracle/`에 마운트하고 `phys-schost-2`는 디스크 장치 `d1`을 `/global/oracle/logs/`에 마운트하는 시나리오를 가정합니다. 이 구성에서는 `phys-schost-1`이 부트되어 `/global/oracle/`을 마운트한 후에만 `phys-schost-2`가 부트되어 `/global/oracle/logs/`를 마운트할 수 있습니다.
5. 클러스터의 노드에서 구성 검사 유틸리티를 실행합니다.

```
phys-schost# cluster check -k vfstab
```

구성 검사 유틸리티에서는 마운트 지점이 있는지 확인합니다. 또한 `/etc/vfstab` 파일 항목이 클러스터의 모든 노드에서 올바른지 확인합니다. 오류가 발생하지 않으면 아무 출력도 반환되지 않습니다.

자세한 내용은 `cluster(1CL)` 매뉴얼 페이지를 참조하십시오.

6. 클러스터의 노드에서 클러스터 파일 시스템을 마운트합니다.

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

7. 클러스터의 각 노드에서 클러스터 파일 시스템이 마운트되는지 확인합니다.
`df` 명령이나 `mount` 명령을 사용하여 마운트된 파일 시스템을 나열할 수 있습니다. 자세한 내용은 `df(1M)` 매뉴얼 페이지 또는 `mount(1M)` 매뉴얼 페이지를 참조하십시오.

▼ 클러스터 파일 시스템을 제거하는 방법

클러스터 파일 시스템을 마운트 해제하여 제거합니다. 또한, 데이터를 제거하거나 삭제하려면 시스템에서 주요 디스크 장치(또는 메타 장치나 볼륨)를 제거하십시오.

주 - cluster shutdown 명령을 실행하여 전체 클러스터를 중지하면 시스템이 종료될 때 클러스터 파일 시스템이 자동으로 마운트 해제됩니다. 단일 노드를 중지하기 위해 shutdown 명령을 실행하면 클러스터 파일 시스템이 마운트 해제되지 않습니다. 그러나 디스크에 연결된 노드가 현재 종료되는 노드 하나뿐인 경우에는 해당 디스크에 있는 클러스터 파일 시스템에 액세스하려고 하면 오류가 발생합니다.

다음의 필수 조건은 클러스터 파일 시스템을 마운트 해제하기 전에 완료되어야 함을 확인하십시오.

- 클러스터에서 root 역할 권한은 노드에서 설정합니다.
- 파일 시스템은 사용 중이 아닙니다. 사용자가 디렉토리에서 작업 중이거나 프로그램이 파일 시스템에서 열려 있다면 해당 파일 시스템이 사용 중인 것으로 간주됩니다. 사용자 또는 프로그램이 클러스터의 어느 노드에서나 실행할 수 있습니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터 파일 시스템을 제거할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 노드에서 root 역할을 수행합니다.
2. 마운트할 클러스터 파일 시스템을 결정합니다.

```
# mount -v
```

3. 각 노드에서 클러스터 파일 시스템을 사용하는 모든 프로세스를 표시하십시오. 그러면 중지시킬 프로세스를 알 수 있습니다.

```
# fuser -c [ -u ] mountpoint
```

-c 파일 시스템의 마운트 지점인 파일과 마운트된 파일 시스템 내의 모든 파일에 대하여 보고합니다.

-u (옵션) 각 프로세스 ID에 대한 사용자 로그인 이름을 표시합니다.

mountpoint 프로세스를 중지시킬 클러스터 파일 시스템의 이름을 지정합니다.

4. 각 노드에서 클러스터 파일 시스템에 대한 모든 프로세스를 중지시킵니다.
원하는 방법을 사용하여 프로세스를 중지시키십시오. 필요한 경우 다음 명령을 사용하여 클러스터 파일 시스템과 관련된 프로세스를 강제로 종료하십시오.

```
# fuser -c -k mountpoint
```

클러스터 파일 시스템을 사용하는 각 프로세스에 SIGKILL 명령이 전달됩니다.

5. 각 노드에서 파일 시스템을 사용하는 프로세스가 없는지 확인합니다.

```
# fuser -c mountpoint
```

6. 한 노드에서만 파일 시스템을 마운트 해제합니다.

```
# umount mountpoint
```

mountpoint 마운트를 해제할 클러스터 파일 시스템의 이름을 지정합니다. 이것은 클러스터 파일 시스템이 마운트되는 디렉토리 이름 또는 파일 시스템의 장치 이름 경로일 수 있습니다.

7. (옵션) **/etc/vfstab** 파일을 편집하여 제거되는 클러스터 파일 시스템에 대한 항목을 삭제합니다.

/etc/vfstab 파일에 이 클러스터 파일 시스템에 대한 항목이 있는 각 클러스터 노드에서 이 단계를 수행합니다.

8. (옵션) 디스크 장치 **group/metadevice/volume/plex**를 제거합니다.

자세한 내용은 볼륨 관리자 설명서를 참조하십시오.

예 51 클러스터 파일 시스템 제거

다음 예에서는 Solaris Volume Manager 메타 장치 또는 볼륨 /dev/md/oracle/rdisk/d1에 마운트된 UFS 클러스터 파일 시스템을 제거합니다.

```
# 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

    각 노드에서 강조 표시된 항목 제거
# 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
```

저장 후 종료

클러스터 파일 시스템의 데이터를 제거하려면 하부 장치를 제거하십시오. 자세한 내용은 볼륨 관리자 설명서를 참조하십시오.

▼ 클러스터의 전역 마운트를 확인하는 방법

`cluster(1CL)` 유틸리티는 `/etc/vfstab` 파일에서 클러스터 파일 시스템에 대한 항목의 구문을 확인합니다. 오류가 발생하지 않으면 아무 것도 반환되지 않습니다.

주 - 클러스터 파일 시스템 제거와 같이 장치나 볼륨 관리 구성요소에 영향을 준 클러스터 구성을 변경한 후 `cluster check` 명령을 실행합니다.

1. 클러스터 노드에서 `root` 역할을 수행합니다.
2. 클러스터 전역 마운트를 확인합니다.

```
# cluster check -k vfstab
```

디스크 경로 모니터링 관리

디스크 경로 모니터링(DPM) 관리 명령을 사용하면 보조 디스크 경로 오류에 대한 알림을 받을 수 있습니다. 디스크 경로 모니터링과 관련된 관리 작업을 수행하려면 이 절의 절차를 수행하십시오.

다음 추가 정보를 사용할 수 있습니다.

항목	정보
디스크 경로 모니터링 데몬에 대한 개념 정보	Oracle Solaris Cluster 4.3 Concepts Guide 의 3 장, "Key Concepts for System Administrators and Application Developers"
<code>cldevice</code> 명령 옵션 및 관련 명령에 대한 설명	cldevice(1CL) 매뉴얼 페이지
<code>scdpmd</code> 데몬 조정	scdpmd.conf(4) 매뉴얼 페이지
<code>syslogd</code> 데몬에서 보고하는 기록된 오류	syslogd(1M) 매뉴얼 페이지

주 - `cldevice` 명령을 사용하여 노드에 I/O 장치를 추가할 때 모니터링 목록에 디스크 경로가 자동으로 추가됩니다. Oracle Solaris Cluster 명령을 사용하여 노드에서 장치를 제거할 경우에도 디스크 경로가 자동으로 모니터 해제됩니다.

표 9 작업 맵: 디스크 경로 모니터링 관리

작업	지침
디스크 경로를 모니터링합니다.	디스크 경로를 모니터링하는 방법 [149]

작업	지침
디스크 경로를 모니터링 해제합니다.	디스크 경로의 모니터를 해제하는 방법 [150]
노드에 대한 오류 디스크 경로의 상태를 인쇄합니다.	오류 디스크 경로를 인쇄하는 방법 [151]
파일에서 디스크 경로를 모니터링합니다.	파일의 디스크 경로를 모니터링하는 방법 [152]
모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 활성화하거나 비활성화합니다.	모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 활성화하는 방법 [154] 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 비활성화하는 방법 [155]
잘못된 디스크 경로 상태를 확인합니다. 모니터링된 DID 장치를 부트 시 사용할 수 없고 DID 인스턴스가 DID 드라이버로 업로드되지 않으면 잘못된 디스크 경로 상태가 보고될 수 있습니다.	디스크 경로 상태 오류를 해결하는 방법 [152]

cldevice 명령을 실행하는 다음 절의 절차에는 디스크 경로 인수가 포함됩니다. 디스크 경로 인수는 노드 이름 및 디스크 이름으로 구성됩니다. 노드 이름은 필수 항목이 아니며 노드 이름을 지정하지 않은 경우 기본적으로 all로 설정됩니다.

▼ 디스크 경로를 모니터링하는 방법

클러스터의 디스크 경로를 모니터링하려면 이 작업을 수행하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 디스크 경로의 모니터링을 사용으로 설정할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 디스크 경로를 모니터링합니다.

```
# cldevice monitor -n node disk
```

3. 디스크 경로가 모니터링되는지 확인합니다.

```
# cldevice status device
```

예 52 단일 노드의 디스크 경로 모니터링

다음 예에서는 단일 노드의 `schost-1:/dev/did/rdisk/d1` 디스크 경로를 모니터링합니다. `schost-1` 노드의 DPM 데몬에서만 `/dev/did/dsk/d1` 디스크에 대한 경로를 모니터링합니다.

```
# cldevice monitor -n schost-1 /dev/did/dsk/d1
# cldevice status d1
```

Device Instance	Node	Status

/dev/did/rdisk/d1	phys-schost-1	Ok

예 53 모든 노드의 디스크 경로 모니터링

다음 예에서는 모든 노드의 `schost-1:/dev/did/dsk/d1` 디스크 경로를 모니터링합니다. `/dev/did/dsk/d1`이 유효한 경로인 모든 노드에서 DPM이 시작됩니다.

```
# cldevice monitor /dev/did/dsk/d1
# cldevice status /dev/did/dsk/d1
```

Device Instance	Node	Status

/dev/did/rdisk/d1	phys-schost-1	Ok

예 54 CCR의 디스크 구성 다시 읽기

다음 예에서는 데몬이 CCR의 디스크 구성을 다시 읽어서 모니터링된 디스크 경로를 상태와 함께 인쇄합니다.

```
# 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

▼ 디스크 경로의 모니터를 해제하는 방법

디스크 경로의 모니터를 해제하려면 다음 절차를 수행합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 디스크 경로의 모니터링을 사용 안함으로 설정할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 모니터링을 해제할 디스크 경로의 상태를 확인합니다.

```
# cldevice status device
```

3. 각 노드에서 해당하는 디스크 경로의 모니터링을 해제합니다.

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

예 55 디스크 경로 모니터링 취소

다음 예에서는 schost-2: /dev/did/rdisk/d1 디스크 경로에 대한 모니터링을 해제하고 디스크 경로를 전체 클러스터 상태와 함께 인쇄합니다.

```
# 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

▼ 오류 디스크 경로를 인쇄하는 방법

클러스터의 오류 디스크 경로를 인쇄하려면 다음 절차를 사용하십시오.

1. 클러스터 노드에서 `root` 역할을 수행합니다.
2. 클러스터에서 오류가 발생한 디스크 경로를 인쇄합니다.

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

예 56 오류 디스크 경로 인쇄

다음 예에서는 전체 클러스터에서 오류가 발생한 디스크 경로를 인쇄합니다.

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

Device Instance	Node	Status
-----	----	-----
dev/did/dsk/d4	phys-schost-1	fail

▼ 디스크 경로 상태 오류를 해결하는 방법

다음 이벤트가 발생하는 경우 DPM이 온라인 상태로 전환되면 오류가 있는 경로의 상태가 업데이트되지 않을 수 있습니다.

- 모니터된 경로 오류로 인해 노드가 재부트됩니다.
- 모니터된 DID 경로 아래에 있는 장치는 재부트된 노드가 온라인 상태로 전환될 때까지 온라인으로 전환되지 않습니다.

모니터된 DID 장치를 부트 시 사용할 수 없고 따라서 DID 인스턴스가 DID 드라이버로 업로드되지 않으므로 잘못된 디스크 경로 상태가 보고됩니다. 이 상황이 발생하면 DID 정보를 수동으로 업데이트해야 합니다.

1. 한 노드에서 전역 장치 이름 공간을 업데이트합니다.

```
# cldevice populate
```

2. 각 노드에서 다음 단계로 진행하기 전에 명령 처리가 완료되었는지 확인합니다.

이 명령이 하나의 노드에서 실행되더라도 모든 노드에서 원격으로 실행됩니다. 명령 처리가 완료되었는지 판별하려면 클러스터의 각 노드에서 다음 명령을 실행하십시오.

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

3. DPM 폴링 시간 프레임 내에서 오류가 있는 디스크 경로 상태가 이제 정상인지 확인합니다.

```
# cldevice status disk-device
```

Device Instance	Node	Status
-----	----	-----
dev/did/dsk/dN	phys-schost-1	Ok

▼ 파일의 디스크 경로를 모니터링하는 방법

파일의 디스크 경로를 모니터링 또는 모니터링 해제하려면 다음 절차를 수행하십시오.

파일을 사용하여 클러스터 구성을 변경하려면 맨 먼저 현재 구성을 내보내야 합니다. 이 내보내기 작업에서는 변경할 구성 항목을 설정하기 위해 수정할 수 있는 XML 파일을 만듭니다. 이 절차의 지침은 전체 프로세스를 설명합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 장치 구성을 XML 파일로 내보냅니다.

```
# cldevice export -o configurationfile
```

-o XML 파일의 파일 이름을 지정합니다.
configurationfile

3. 구성 파일을 수정하여 장치 경로를 모니터링합니다.

모니터할 장치 경로를 찾고 `monitored` 속성을 `true`로 설정합니다.

4. 장치 경로를 모니터링합니다.

```
# cldevice monitor -i configurationfile
```

-i 수정된 XML 파일의 이름을 지정합니다.
configurationfile

5. 이제 장치 경로가 모니터링되는지 확인합니다.

```
# cldevice status
```

예 57 파일을 사용한 디스크 경로 모니터

다음 예에서는 XML 파일을 사용하여 노드 `phys-schost-2` 및 장치 `d3` 사이의 장치 경로를 모니터링합니다. `deviceconfig` XML 파일은 `phys-schost-2`와 `d3` 간의 경로가 현재 모니터링되고 있지 않음을 나타냅니다.

현재 클러스터 구성 내보내기

```
# 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="clt8d0">
<devicePath nodeRef="phys-schost-1" monitored="true"/>
<devicePath nodeRef="phys-schost-2" monitored="false"/>
</device>
</deviceList>
```

```
</cluster>
```

모니터링 속성을 *true*로 설정하여 경로 모니터

```
<?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>
```

파일을 읽은 후 모니터링 설정

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

장치가 모니터링되는지 확인

```
# cldevice status
```

참조 클러스터 구성 내보내기 및 결과로 생성된 XML 파일을 사용하여 클러스터 구성 설정에 대한 자세한 내용은 [cluster\(1CL\)](#) 및 [clconfiguration\(5CL\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 활성화하는 방법

이 기능을 활성화한 상태에서는 다음 조건이 충족될 경우 노드가 자동으로 재부트됩니다.

- 노드의 모니터링된 모든 디스크 경로가 실패합니다.
- 최소 하나의 디스크가 클러스터의 다른 노드에서 액세스할 수 있습니다.

노드를 재부트하면 해당 노드에서 마스터되는 모든 리소스 그룹 및 장치 그룹이 다른 노드에서 재시작됩니다.

노드가 자동으로 재부트된 후에도 노드의 모니터링된 모든 공유 디스크 경로에 계속 액세스할 수 없는 경우 노드가 다시 자동으로 재부트되지 않습니다. 그러나, 노드가 재부트된 후 디스크 경로가 사용 가능한 상태로 되었다가 실패한 경우 노드가 다시 자동으로 재부트됩니다.

`reboot_on_path_failure` 등록 정보를 활성화하면 노드 재부트가 필요한지 여부를 확인할 때 로컬 디스크 경로의 상태가 고려되지 않습니다. 모니터링된 공유 디스크만 영향을 받습니다.

주 - 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 `reboot_on_path_failure` 노드 등록 정보를 편집할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 클러스터의 모든 노드에 대해 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 활성화합니다.

```
# clnode set -p reboot_on_path_failure=enabled +
```

▼ 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 비활성화하는 방법

이 기능을 비활성화한 경우 노드의 모니터링된 모든 공유 디스크 경로가 실패하면 노드가 자동으로 재부트되지 않습니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 클러스터의 모든 노드에 대해 모니터링된 모든 공유 디스크 경로가 실패할 경우 노드의 자동 재부트를 비활성화합니다.

```
# clnode set -p reboot_on_path_failure=disabled +
```


쿼럼 관리

이 장에서는 Oracle Solaris Cluster 및 Oracle Solaris Cluster 쿼럼 서버에서 쿼럼 장치를 관리하는 절차에 대해 설명합니다. 쿼럼 개념에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Concepts Guide](#)의 “Quorum and Quorum Devices”를 참조하십시오.

- “쿼럼 장치 관리” [157]
- “Oracle Solaris Cluster 쿼럼 서버 관리” [175]

쿼럼 장치 관리

쿼럼 장치는 공유된 저장 장치 또는 쿼럼 서버로 2개 이상의 노드에서 공유되고 쿼럼을 설정하는 데 사용하는 투표를 구성합니다. 이 절에서는 쿼럼 장치를 관리하는 절차에 대해 설명합니다.

`clquorum` 명령을 사용하여 모든 쿼럼 장치 관리 절차를 수행할 수 있습니다. 더불어, `clsetup` 대화식 유틸리티 또는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 일부 절차를 수행할 수 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster 4.3 시스템 관리 설명서](#)의 “Oracle Solaris Cluster Manager에 액세스하는 방법”을 참조하십시오.

가능한 경우 이 절에서는 `clsetup` 유틸리티를 사용하여 쿼럼 절차에 대해 설명합니다. Oracle Solaris Cluster Manager 온라인 도움말에서는 Oracle Solaris Cluster Manager를 사용하여 쿼럼 절차를 수행하는 방법에 대해 설명합니다. 자세한 내용은 `clquorum(1CL)` 및 `clsetup(1CL)` 매뉴얼 페이지를 참조하십시오.

쿼럼 장치로 작업할 때 다음의 안내 사항을 숙지하십시오.

- 모든 쿼럼 명령은 전역 클러스터 노드에서 실행되어야 합니다.
- 모든 쿼럼 관련 `clquorum` 명령이 중단되거나 실패하면 클러스터 구성 데이터베이스에서 쿼럼 구성 정보가 일치하지 않을 수 있습니다. 이러한 불일치가 발생하는 경우, 위의 명령을 재실행하거나 `clquorum reset` 명령을 실행하여 쿼럼 구성을 재설정합니다.
- 클러스터의 고가용성을 위해, 쿼럼 장치에 의해 분배된 총 투표수가 노드에 의해 분배된 총 투표수보다 적은지 확인하십시오. 그렇지 않은 경우 모든 쿼럼 장치를 사용할 수 없으면 모든 노드가 작동하고 있더라도 노드가 클러스터를 형성할 수 없습니다.

- 현재 쿼럼 장치로서 구성된 디스크를 Oracle Solaris ZFS 저장소 풀에 추가하지 마십시오. 구성된 쿼럼 장치를 ZFS 저장소 풀에 추가하면 디스크는 EFI 디스크로 레이블이 다시 지정되고 쿼럼 구성 정보가 손실되어 디스크는 클러스터에 더 이상의 쿼럼 투표를 제공하지 않습니다. 일단 디스크가 저장소 풀에 들어가면 해당 디스크를 쿼럼 장치로 구성할 수 있습니다. 또는 디스크 구성을 취소하고 저장소 풀에 추가한 후 디스크를 쿼럼 장치로서 다시 구성할 수 있습니다.

주 - `clsetup` 명령은 다른 Oracle Solaris Cluster 명령과의 대화식 인터페이스입니다. `clsetup`을 실행하면 `clquorum` 명령과 같은 적절한 특정 명령이 생성됩니다. 생성된 이러한 명령은 절차 끝의 예에 나와 있습니다.

쿼럼 구성을 보려면 `clquorum show`를 사용합니다. 클러스터에서 `clquorum list` 명령은 쿼럼 장치의 이름을 표시합니다. `clquorum status` 명령은 상태 및 투표 수 정보를 제공합니다.

이 절에 있는 예에서는 대부분 세 개의 노드로 구성된 클러스터를 기준으로 설명합니다.

표 10 작업 목록: 쿼럼 관리

작업	지침
<code>clsetup</code> 유틸리티를 사용하여 클러스터에 쿼럼 장치 추가	“쿼럼 장치 추가” [159]
<code>clsetup</code> 유틸리티를 사용하여 클러스터에서 쿼럼 장치 제거(<code>clquorum</code> 생성)	쿼럼 장치를 제거하는 방법 [165]
<code>clsetup</code> 유틸리티를 사용하여 클러스터에서 마지막 쿼럼 장치 제거(<code>clquorum</code> 생성)	클러스터에서 마지막 쿼럼 장치를 제거하는 방법 [166]
추가 및 제거 절차를 사용하여 클러스터에서 쿼럼 장치 교체	쿼럼 장치를 교체하는 방법 [168]
추가 및 제거 절차를 사용하여 쿼럼 장치 목록 수정	쿼럼 장치 노드 목록을 수정하는 방법 [169]
<code>clsetup</code> 유틸리티를 사용하여 쿼럼 장치를 유지 보수 상태로 전환(<code>clquorum</code> 생성)	쿼럼 장치를 유지 보수 상태로 전환하는 방법 [170]
유지 보수 상태에 있으면 쿼럼 장치가 쿼럼을 구성하는데 포함되지 않습니다.	
<code>clsetup</code> 유틸리티를 사용하여 쿼럼 구성을 기본 상태로 재설정(<code>clquorum</code> 생성)	쿼럼 장치를 유지 보수 상태에서 해제하는 방법 [171]
<code>clquorum</code> 명령을 사용하여 쿼럼 장치 및 투표 수 나열	쿼럼 구성을 나열하는 방법 [173]

쿼럼 장치 동적 재구성

클러스터의 쿼럼 장치에 대해 동적 재구성 작업을 완료할 때는 몇 가지 사항을 고려해야 합니다.

- Oracle Solaris 동적 재구성 기능에 대해 문서화된 모든 요구사항, 절차 및 제한 사항은 운영체제의 작동이 정지된 경우를 제외하고는 Oracle Solaris Cluster 동적 재구성 지원에도 적용됩니다. 따라서 Oracle Solaris Cluster 소프트웨어와 함께 동적 재구성 기능을

사용하기 전에 Oracle Solaris 동적 재구성 기능에 대한 설명서를 검토하십시오. 특히 동적 재구성 연결 종료 작업 중 비네트워크 IO 장치에 영향을 주는 문제를 검토해야 합니다.

- Oracle Solaris Cluster에서는 쿼럼 장치를 위해 구성된 인터페이스가 있을 경우 동적 재구성 보드 제거 작업을 수행할 수 없습니다.
- 동적 재구성 작업이 활성 장치에 영향을 주는 경우에는 Oracle Solaris Cluster가 작업을 거부하고 작업의 영향을 받는 장치를 확인합니다.

쿼럼 장치를 제거하려면 표시된 순서대로 다음 단계를 완료해야 합니다.

표 11 작업 맵: 쿼럼 장치 동적 재구성

작업	지침
1. 제거되는 쿼럼 장치를 교체할 새 쿼럼 장치 활성화	“쿼럼 장치 추가” [159]
2. 제거할 쿼럼 장치 비활성화	쿼럼 장치를 제거하는 방법 [165]
3. 제거하려는 장치에 대해 동적 재구성 제거 작업을 수행합니다.	

쿼럼 장치 추가

이 절에서는 쿼럼 장치를 추가하는 절차를 제공합니다. 새 쿼럼 장치를 추가하기 전에 클러스터의 모든 노드가 온라인 상태인지 확인합니다. 클러스터에 필요한 쿼럼 투표 수 결정, 권장되는 쿼럼 구성 및 실패 보호에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Concepts Guide](#)의 “Quorum and Quorum Devices”를 참조하십시오.



주의 - 현재 쿼럼 장치로서 구성된 디스크를 Solaris ZFS 저장소 풀에 추가하지 마십시오. 구성된 쿼럼 장치를 Solaris ZFS 저장소 풀에 추가하면 디스크는 EFI 디스크로 레이블이 다시 지정되고 쿼럼 구성 정보가 손실되어 디스크는 클러스터에 더 이상의 쿼럼 투표를 제공하지 않습니다. 일단 디스크가 저장소 풀에 들어가면 해당 디스크를 쿼럼 장치로 구성할 수 없습니다. 디스크 구성을 취소하고 저장소 풀에 추가한 후 디스크를 쿼럼 장치로 다시 구성할 수도 있습니다.

Oracle Solaris Cluster 소프트웨어는 다음 유형의 쿼럼 장치를 지원합니다.

- 다음 위치의 공유 LUN:
 - 공유 SCSI 디스크
 - SATA(Serial Attached Technology Attachment) 저장소
 - Oracle ZFS Storage Appliance
 - 지원되는 NAS 장치
- Oracle Solaris Cluster 쿼럼 서버

이 장치를 추가하는 절차는 다음 절에서 설명합니다.

- [공유 디스크 쿼럼 장치를 추가하는 방법 \[160\]](#)
- [쿼럼 서버 쿼럼 장치를 추가하는 방법 \[163\]](#)

주 - 복제된 디스크를 퀵럼 장치로 구성할 수 없습니다. 복제된 디스크를 퀵럼 장치로 추가하려고 하면 다음 오류 메시지가 수신되고 명령이 오류 코드로 종료됩니다.

Disk-name is a replicated device. Replicated devices cannot be configured as quorum devices.

공유 디스크 퀵럼 장치는 Oracle Solaris Cluster 소프트웨어에서 지원하는 연결된 모든 저장 장치입니다. 공유 디스크는 클러스터의 두 개 이상 노드에 연결됩니다. 보호(fencing)를 설정하면 SCSI-2 또는 SCSI-3(기본값은 SCSI-2임)을 사용하는 퀵럼 장치로 이중 포트 디스크를 구성할 수 있습니다. 보호(fencing)가 설정되고 공유 장치가 세 개 이상의 노드에 연결되어 있으면 SCSI-3 프로토콜(세 개 이상의 노드에 대한 기본 프로토콜)을 사용하는 퀵럼 장치로 공유 디스크를 구성할 수 있습니다. SCSI 대체 플래그를 사용하여 Oracle Solaris Cluster 소프트웨어가 이중 포트 공유 디스크에 대해 SCSI-3 프로토콜을 사용하도록 할 수 있습니다.

공유 디스크에 대해 보호(fencing)를 해제하면 소프트웨어 퀵럼 프로토콜을 사용하는 퀵럼 장치로 디스크를 구성할 수 있습니다. 이것은 디스크에서 SCSI-2 또는 SCSI-3 프로토콜을 지원하는지 여부에 관계없이 적용됩니다. 소프트웨어 퀵럼은 Oracle에서 제공하는 프로토콜로, SCSI PGR(Persistent Group Reservations)의 형식을 에뮬레이트합니다.



주의 - SCSI를 지원하지 않는 디스크(예: SATA)를 사용하는 경우 SCSI 보호(fencing)를 해제해야 합니다.

퀵럼 장치에 대해 사용자 데이터를 포함하거나 장치 그룹의 멤버인 디스크를 사용할 수 있습니다. `cluster show` 명령의 출력에서 공유 디스크에 대한 `access-mode` 값을 확인하여 공유 디스크가 있는 퀵럼 부속 시스템에서 사용하는 프로토콜을 확인합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 퀵럼 서버 장치 또는 공유 디스크 퀵럼 장치를 만들 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

다음 절차에서 사용되는 명령에 대한 내용은 `clsetup(1CL)` 및 `clquorum(1CL)` 매뉴얼 페이지를 참조하십시오.

▼ 공유 디스크 퀵럼 장치를 추가하는 방법

Oracle Solaris Cluster 소프트웨어는 공유 디스크(SCSI 및 SATA 모두) 장치를 퀵럼 장치로 지원합니다. SATA 장치는 SCSI 예약을 지원하지 않으며, SCSI 예약 보호(fencing) 플래그를 사용 불가하게 하고 소프트웨어 퀵럼 프로토콜을 사용하여 이러한 디스크를 퀵럼 장치로 구성해야 합니다.

이 절차를 완료하려면 노드에서 공유되는 디스크 드라이브를 해당 장치 ID(DID)로 식별합니다. `cldevice show` 명령을 사용하여 DID 이름 목록을 표시합니다. 자세한 내용은 [cldevice\(1CL\)](#) 매뉴얼 페이지를 참조하십시오. 새 퀵럼 장치를 추가하기 전에 클러스터의 모든 노드가 온라인 상태인지 확인합니다.

SCSI 또는 SATA 장치를 구성하려면 이 절차를 사용합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. `clsetup` 유틸리티를 시작합니다.

```
# clsetup
```

`clsetup` 주 메뉴가 표시됩니다.

3. 퀴럼 옵션에 대한 번호를 입력합니다.
퀴럼 메뉴가 표시됩니다.
4. 퀴럼 장치를 추가하는 옵션에 대한 번호를 입력한 다음, `clsetup` 유틸리티에서 퀴럼 장치 추가에 대한 확인을 요청하면 `yes`를 입력합니다.
`clsetup` 유틸리티에서 추가할 퀴럼 장치 유형을 묻습니다.
5. 공유 디스크 퀴럼 장치 옵션에 대한 번호를 입력합니다.
`clsetup` 유틸리티에서 사용할 전역 장치를 묻습니다.
6. 사용 중인 전역 장치를 입력합니다.
`clsetup` 유틸리티에서 새 퀴럼 장치를 지정된 전역 장치에 추가할 것을 확인하도록 요청합니다.
7. 계속해서 새 퀴럼 장치를 추가하려면 `yes`를 입력합니다.
새 퀴럼 장치가 성공적으로 추가되면 `clsetup` 유틸리티에서 추가된 장치를 보여주는 메시지를 표시합니다.
8. 퀴럼 장치가 추가되었는지 확인합니다.

```
# clquorum list -v
```

▼ Oracle ZFS Storage Appliance NAS 퀴럼 장치를 추가하는 방법

새 퀴럼 장치를 추가하기 전에 클러스터의 모든 노드가 온라인 상태인지 확인합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 Oracle ZFS Storage Appliance NAS 장치를 추가할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. iSCSI 장치 설치에 대한 지침은 Oracle ZFS Storage Appliance와 함께 제공된 설치 설명서 또는 어플라이언스의 온라인 도움말을 참조하십시오.
2. 각 클러스터 노드에서 iSCSI LUN을 검색하고 정적 구성에 대한 iSCSI 액세스 목록을 설정합니다.

```
# 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. 하나의 클러스터 노드에서 iSCSI LUN용 DID를 구성합니다.

```
# cldevice populate
```

4. 방금 iSCSI를 사용하여 클러스터로 구성한 NAS 장치 LUN을 나타내는 DID 장치를 식별합니다.

cldevice show 명령을 사용하여 DID 이름 목록을 표시합니다. 자세한 내용은 [cldevice\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

5. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
6. `clquorum` 명령을 사용하여 [4단계](#)에서 식별된 DID 장치를 사용하는 NAS 장치를 쿼럼 장치로 추가합니다.

```
# clquorum add d20
```

클러스터에는 scsi-2, scsi-3 또는 소프트웨어 쿼럼 프로토콜 중 사용할 프로토콜을 결정하는 기본 규칙이 있습니다. 자세한 내용은 [clquorum\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 쿼럼 서버 쿼럼 장치를 추가하는 방법

시작하기 전에 Oracle Solaris Cluster 쿼럼 서버를 쿼럼 장치로 추가하려면 먼저 Oracle Solaris Cluster 쿼럼 서버 소프트웨어를 호스트 시스템에 설치하고 쿼럼 서버를 시작하여 실행 중이어야 합니다. 쿼럼 서버 설치에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 쿼럼 서버 소프트웨어를 설치하고 구성하는 방법”을 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 쿼럼 서버 장치를 만들 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 모든 Oracle Solaris Cluster 노드가 온라인 상태이고 Oracle Solaris Cluster 쿼럼 서버와 통신할 수 있는지 확인합니다.
 - a. 클러스터 노드와 바로 연결되는 네트워크 스위치가 다음 조건 중 하나를 충족하는지 확인합니다.
 - 이 스위치는 RSTP(Rapid Spanning Tree Protocol)를 지원합니다.
 - 스위치에 고속 포트 모드가 활성화되어 있습니다.

이 기능 중 하나는 클러스터 노드와 쿼럼 서버 사이의 즉각적인 통신을 확인하는 데 필요합니다. 스위치에 의해 이 통신이 두드러지게 지연되는 경우 클러스터는 이러한 통신 장애를 쿼럼장치의 손실로 해석합니다.
 - b. 공용 네트워크가 CIDR(Classless Inter-Domain Routing)이라고도 하는 가변 길이 서브넷 기능을 사용하는 경우 각 노드에서 다음 파일을 수정합니다.

RFC 791에 정의된 Classful 서브넷을 사용하는 경우에는 본 단계를 수행할 필요가 없습니다.

 - i. 클러스터가 사용하는 각 공용 서브넷의 항목을 `/etc/inet/netmasks` 파일에 추가합니다.

다음은 공용 네트워크 IP 주소 및 넷마스크를 포함하는 항목의 예입니다.

```
10.11.30.0 255.255.255.0
```

- ii. 각 `/etc/hostname.adapter` 파일의 호스트 이름 항목에 `netmask + broadcast +`를 추가합니다.

`nodename netmask + broadcast +`

- c. 클러스터의 각 노드에서 `/etc/inet/hosts` 파일 또는 `/etc/inet/ipnodes` 파일에 쿼럼 서버 호스트 이름을 추가합니다.

다음과 같이 호스트 이름 대 주소 매핑을 파일에 추가합니다.

`ipaddress qshost1`

`ipaddress` 쿼럼 서버를 실행 중인 컴퓨터 IP 주소

`qshost1` 쿼럼 서버가 실행 중인 컴퓨터의 호스트 이름

- d. 이름 지정 서비스를 사용하는 경우, 쿼럼 서버 호스트의 이름 대 주소 매핑을 이름 서비스 데이터베이스에 추가합니다.

3. **clsetup** 유틸리티를 시작합니다.

`# clsetup`

`clsetup` 주 메뉴가 표시됩니다.

4. 쿼럼 옵션에 대한 번호를 입력합니다.

쿼럼 메뉴가 표시됩니다.

5. 쿼럼 장치를 추가하는 옵션에 대한 번호를 입력합니다.

그런 다음, 쿼럼 장치 추가를 확인하기 위해 **yes**를 입력합니다.

`clsetup` 유틸리티에서 추가할 쿼럼 장치 유형을 묻습니다.

6. 쿼럼 서버 쿼럼 장치에 대한 옵션 번호를 입력한 다음 **yes**를 입력하여 쿼럼 서버 쿼럼 장치를 추가할 것인지 확인합니다.

`clsetup` 유틸리티에서 새 쿼럼 장치의 이름을 입력할 것을 요청합니다.

7. 추가할 쿼럼 장치의 이름을 입력합니다.

임의로 선택한 이름을 쿼럼 장치 이름으로 사용할 수 있습니다. 이름은 이후의 관리 명령을 처리할 때만 사용됩니다.

`clsetup` 유틸리티에서 쿼럼 서버 호스트의 이름을 입력할 것을 요청합니다.

8. 쿼럼 서버의 호스트 이름을 입력합니다.

이 이름은 네트워크에서 쿼럼 서버가 실행되는 시스템의 IP 주소 또는 시스템에 있는 시스템의 호스트 이름을 지정합니다.

호스트의 IPv4 또는 IPv6 구성에 따라, 시스템의 IP 주소를 `/etc/hosts` 파일, `/etc/inet/` `ipnodes` 파일 또는 두 파일 모두에 지정해야 합니다.

주 - 지정하는 시스템은 모든 클러스터 노드로 연결할 수 있어야 하고 쿼럼 서버를 실행해야 합니다.

`clsetup` 유틸리티에서 쿼럼 서버의 포트 번호를 입력할 것을 요청합니다.

9. 클러스터 노드와 통신할 수 있도록 쿼럼 서버에서 사용하는 포트 번호를 입력하십시오.

`clsetup` 유틸리티에서 새 쿼럼 장치를 추가할 것을 확인하도록 요청합니다.

10. 계속해서 새 쿼럼 장치를 추가하려면 **yes**를 입력합니다.

새 쿼럼 장치가 성공적으로 추가되면 `clsetup` 유틸리티에서 추가된 장치를 보여주는 메시지를 표시합니다.

11. 쿼럼 장치가 추가되었는지 확인합니다.

```
# clquorum list -v
```

쿼럼 장치 제거 또는 교체

이 절에서는 쿼럼 장치를 제거하거나 교체하기 위한 절차를 다음과 같이 설명합니다.

- 쿼럼 장치를 제거하는 방법 [165]
- 클러스터에서 마지막 쿼럼 장치를 제거하는 방법 [166]
- 쿼럼 장치를 교체하는 방법 [168]

▼ 쿼럼 장치를 제거하는 방법

쿼럼 장치를 제거하면 더 이상 쿼럼을 구성하는 데 포함되지 않습니다. 2 노드 클러스터에도 하나 이상의 쿼럼 장치가 구성되어야 합니다. 클러스터에 있는 마지막 쿼럼 장치의 경우 `clquorum(1CL)` 명령을 실행해도 구성에서 장치가 제거되지 않습니다. 노드를 제거하는 경우 해당 노드에 연결된 모든 쿼럼 장치를 제거합니다.

주 - 제거할 장치가 클러스터에 있는 마지막 쿼럼 장치인 경우 [클러스터에서 마지막 쿼럼 장치를 제거하는 방법 \[166\]](#)의 절차를 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 쿼럼 장치를 제거할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 제거할 쿼럼 장치를 결정합니다.

```
# clquorum list -v
```

3. `clsetup` 유틸리티를 실행합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.

4. 쿼럼 옵션에 대한 번호를 입력합니다.

5. 쿼럼 장치를 제거하는 옵션에 대한 번호를 입력합니다.

제거 프로세스 동안 표시되는 질문에 응답하십시오.

6. `clsetup`을 종료합니다.

7. 쿼럼 장치가 제거되었는지 확인합니다.

```
# clquorum list -v
```

일반 오류 쿼럼 서버 쿼럼 장치를 제거하는 동안 클러스터와 쿼럼 서버 호스트 간의 통신이 끊어지는 경우 더 이상 사용되지 않는 쿼럼 서버 호스트의 구성 정보를 정리해야 합니다. 이 정리를 수행하는 지침은 [“오래된 쿼럼 서버 클러스터 정보 정리” \[178\]](#)를 참조하십시오.

▼ 클러스터에서 마지막 쿼럼 장치를 제거하는 방법

이 절차에서는 `clquorum force` 옵션 `-f`를 사용하여 2 노드 클러스터의 마지막 쿼럼 장치를 제거합니다. 일반적으로 먼저 실패한 장치를 제거한 다음 교체용 쿼럼 장치를 추가해야 합니다. 이 장치가 2 노드 클러스터의 마지막 쿼럼 장치가 아닌 경우 [쿼럼 장치를 제거하는 방법 \[165\]](#)의 단계를 수행합니다.

쿼럼 장치를 추가하려면 오류가 발생한 쿼럼 장치에 접근하는 노드 재구성이 필요하며 시스템이 패닉 상태가 될 수 있습니다. Force 옵션을 사용하면 시스템에서 패닉 상태를 발생시키지 않고 오류가 발생한 쿼럼 장치를 제거할 수 있습니다. `clquorum` 명령을 사용하면 구성에서 장치를 제거할 수 있습니다. 자세한 내용은 [clquorum\(1CL\)](#) 매뉴얼 페이지를 참조하십시오. 실패한 쿼럼 장치를 제거한 후 `clquorum add` 명령을 사용하여 새 장치를 추가할 수 있습니다. [“쿼럼 장치 추가” \[159\]](#)를 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. `clquorum` 명령을 사용하여 쿼럼 장치를 제거합니다.
쿼럼 장치가 실패한 경우 `-F(Force)` 옵션을 사용하여 오류가 발생한 장치를 제거합니다.

```
# clquorum remove -F qd1
```

주 - 제거할 노드를 유지 보수 상태로 전환한 다음 `clquorum remove quorum` 명령을 사용하여 쿼럼 장치를 제거할 수도 있습니다. `clsetup` 클러스터 관리 메뉴 옵션은 클러스터가 설치 모드에 있는 동안 사용할 수 없습니다. 자세한 내용은 [노드를 유지 보수 상태로 전환하는 방법 \[227\]](#) 및 `clsetup(1CL)` 매뉴얼 페이지를 참조하십시오.

3. 쿼럼 장치가 제거되었는지 확인합니다.

```
# clquorum list -v
```
4. 마지막 쿼럼 장치를 제거하려는 이유에 따라 다음 단계 중 하나를 진행합니다.

■ 제거된 쿼럼 장치를 교체하는 경우 다음 하위 단계를 완료합니다.

- a. 새 쿼럼 장치를 추가합니다.
새 쿼럼 장치 추가에 대한 지침은 [“쿼럼 장치 추가” \[159\]](#)를 참조하십시오.
- b. 설치 모드에서 클러스터를 제거합니다.

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

■ 현재 클러스터를 단일 노드 클러스터로 축소하는 경우 설치 모드에서 클러스터를 제거합니다.

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

예 58 마지막 쿼럼 장치 제거

이 예에서는 클러스터를 유지 보수 모드로 전환하고 클러스터 구성에서 마지막 남은 쿼럼 장치를 제거하는 방법을 보여 줍니다.

설치 모드로 클러스터 전환

```
# cluster set -p installmode=enabled

쿼럼 장치 제거
# clquorum remove d3

쿼럼 장치가 제거되었는지 확인합니다
# clquorum list -v
Quorum      Type
-----
scphyshost-1 node
scphyshost-2 node
scphyshost-3 node
```

▼ 쿼럼 장치를 교체하는 방법

이 절차를 사용하여 기존의 쿼럼 장치를 다른 쿼럼 장치로 교체합니다. 쿼럼 장치는 NAS 장치를 다른 NAS 장치로 교체하는 것처럼 유사한 장치 유형으로 교체하거나, NAS 장치를 공유 디스크로 교체하는 것처럼 다른 장치로 교체할 수 있습니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 새 쿼럼 장치를 구성합니다.

이전 장치를 대신하도록 하려면 먼저 구성에 새 쿼럼 장치를 추가해야 합니다. 새 쿼럼 장치를 클러스터에 추가하려면 [“쿼럼 장치 추가” \[159\]](#)를 참조하십시오.

2. 교체할 쿼럼 장치를 제거합니다.

구성에서 기존 쿼럼 장치를 제거하려면 [쿼럼 장치를 제거하는 방법 \[165\]](#)을 참조하십시오.

3. 쿼럼 장치가 오류가 있는 디스크인 경우에는 디스크를 교체합니다.

디스크 외장 장치는 하드웨어 설명서의 하드웨어 절차를 참조하십시오. [Oracle Solaris Cluster Hardware Administration Manual](#)도 참조하십시오.

쿼럼 장치 유지 보수

이 절에서는 쿼럼 장치를 유지 보수하기 위한 절차를 다음과 같이 설명합니다.

- [쿼럼 장치 노드 목록을 수정하는 방법 \[169\]](#)
- [쿼럼 장치를 유지 보수 상태로 전환하는 방법 \[170\]](#)

- 쿼럼 장치를 유지 보수 상태에서 해제하는 방법 [171]
- 쿼럼 구성을 나열하는 방법 [173]
- 쿼럼 장치를 복구하는 방법 [174]
- “쿼럼 기본 시간 초과 변경” [174]

▼ 쿼럼 장치 노드 목록을 수정하는 방법

clsetup 유틸리티를 사용하여 기존 쿼럼 장치의 노드 목록에 노드를 추가하거나 목록에서 노드를 제거할 수 있습니다. 쿼럼 장치의 노드 목록을 변경하려면 쿼럼 장치를 제거하고 제거한 쿼럼 장치에 대한 노드의 물리적 연결을 수정한 후에 쿼럼 장치를 다시 클러스터 구성에 추가해야 합니다. 쿼럼 장치를 추가하면 clquorum 명령이 디스크에 연결된 모든 노드에 대해 노드-디스크 경로를 자동으로 구성합니다. 자세한 내용은 [clquorum\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 변경하는 쿼럼 장치의 이름을 확인합니다.

`clquorum list -v`
3. `clsetup` 유틸리티를 시작합니다.

`clsetup`

주 메뉴가 표시됩니다.
4. 쿼럼 옵션에 대한 번호를 입력합니다.

쿼럼 메뉴가 표시됩니다.
5. 쿼럼 장치를 제거하는 옵션에 대한 번호를 입력합니다.

지침을 따릅니다. 제거할 디스크의 이름을 묻는 메시지가 표시됩니다.
6. 쿼럼 장치와의 노드 연결을 추가하거나 삭제합니다.
7. 쿼럼 장치를 추가하는 옵션에 대한 번호를 입력합니다.

지침을 따릅니다. 쿼럼 장치로 사용할 디스크의 이름을 묻는 메시지가 표시됩니다.
8. 쿼럼 장치가 추가되었는지 확인합니다.

```
# clquorum list -v
```

▼ 쿼럼 장치를 유지 보수 상태로 전환하는 방법

clquorum 명령을 사용하여 쿼럼 장치를 유지 보수 상태로 전환할 수 있습니다. 자세한 내용은 [clquorum\(1CL\)](#) 매뉴얼 페이지를 참조하십시오. 현재 clsetup 유틸리티에는 이 기능이 없습니다.

오랫동안 쿼럼 장치의 서비스를 중단하는 경우 쿼럼 장치를 유지 보수 상태로 전환합니다. 이 방법으로 쿼럼 장치의 쿼럼 투표 수는 0으로 설정되며 해당 장치가 서비스되는 중에는 쿼럼 수에 포함되지 않습니다. 쿼럼 장치의 구성 정보는 유지 보수 상태에 있는 동안에도 보존됩니다.

주 - 두 개의 노드로 구성된 클러스터에도 하나 이상의 쿼럼 장치가 구성되어야 합니다. 유지 보수 상태로 전환할 장치가 2 노드 클러스터의 마지막 쿼럼 장치인 경우에는 clquorum을 실행해도 장치가 유지 보수 상태로 전환되지 않습니다.

클러스터 노드를 유지 보수 상태로 전환하려면 [노드를 유지 보수 상태로 전환하는 방법 \[227\]](#)을 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 쿼럼 장치를 사용 안 함으로 설정하고 유지 보수 상태로 설정할 수 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

클러스터가 설치 모드인 경우 Reset Quorum Devices(쿼럼 장치 재설정)를 눌러 설치 모드를 종료합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 쿼럼 장치를 유지 보수 상태로 만듭니다.

```
# clquorum disable device
```

device 변경할 디스크 장치의 DID 이름을 지정합니다(예: d4).

3. 쿼럼 장치가 현재 유지 보수 상태에 있는지 확인합니다.

유지 보수 상태로 만든 장치의 출력은 쿼럼 장치 투표가 0이 되어야 합니다.

```
# clquorum status device
```

예 59 쿼럼 장치를 유지 보수 상태로 만들기

다음 예는 쿼럼 장치를 유지 보수 상태로 만들고 결과를 확인하는 방법입니다.

```
# clquorum disable d20
# clquorum status d20

=== Cluster Quorum ===

--- Quorum Votes by Device ---

Device Name      Present    Possible    Status
-----
d20              1          1          Offline
```

참조 쿼럼 장치를 다시 사용 가능하게 하려면 [쿼럼 장치를 유지 보수 상태에서 해제하는 방법 \[171\]](#)을 참조하십시오.

노드를 유지 보수 상태로 전환하려면 [노드를 유지 보수 상태로 전환하는 방법 \[227\]](#)을 참조하십시오.

▼ 쿼럼 장치를 유지 보수 상태에서 해제하는 방법

쿼럼 장치가 유지 보수 상태이며 쿼럼 장치를 유지 보수 상태에서 해제하여 쿼럼 투표 수를 기본값으로 재설정하려는 경우 이 절차를 실행합니다.



주의 - globaldev 또는 node 옵션을 지정하지 않으면 쿼럼 계수가 전체 클러스터에 대해 재설정됩니다.

쿼럼 장치를 구성할 때 Oracle Solaris Cluster 소프트웨어는 $N-1$ 의 투표 수를 쿼럼 장치에 지정합니다. 여기서 N 은 쿼럼 장치에 연결된 투표 수입니다. 예를 들어, 투표 수가 0이 아닌 두 노드에 연결된 쿼럼 장치의 쿼럼 수는 1입니다($2 - 1$).

- 클러스터 노드와 관련 쿼럼 장치를 유지 보수 상태에서 해제하려면 [노드를 유지 보수 상태에서 해제하는 방법 \[229\]](#)을 참조하십시오.
- 쿼럼 투표 수에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Concepts Guide](#)의 “About Quorum Vote Counts”를 참조하십시오.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 쿼럼 장치를 사용으로 설정하고 유지 보수 상태를 종료할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 임의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 쿼럼 계수를 다시 설정합니다.

```
# clquorum enable device
```

device 재설정할 쿼럼 장치의 DID 이름을 지정합니다(예: d4).

3. 노드가 유지 보수 상태에 있었기 때문에 쿼럼 계수를 초기화하려면 노드를 재부트합니다.

4. 쿼럼 투표 수를 확인하십시오.

```
# clquorum show +
```

예 60 쿼럼 투표 수 재설정(쿼럼 장치)

다음 예에서는 쿼럼 장치에 대한 쿼럼 수를 다시 기본값으로 초기화하고 결과를 확인합니다.

```
# clquorum enable d20
```

```
# clquorum show +
```

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

```
Node Name: phys-schost-2
Node ID: 1
Quorum Vote Count: 1
Reservation Key: 0x43BAC41300000001
```

```
Node Name: phys-schost-3
Node ID: 2
Quorum Vote Count: 1
Reservation Key: 0x43BAC41300000002
```

```
=== Quorum Devices ===
```

```
Quorum Device Name: d3
Enabled: yes
Votes: 1
Global Name: /dev/did/rdisk/d20s2
Type: shared_disk
Access Mode: scsi3
Hosts (enabled): phys-schost-2, phys-schost-3
```

▼ 쿼럼 구성을 나열하는 방법

쿼럼 구성을 나열하기 위해 root 역할로 전환할 필요는 없습니다. `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할은 모두 가능합니다.

주 - 쿼럼 장치에 대한 노드 연결 수를 늘리거나 줄일 경우 쿼럼 투표 수가 자동으로 재계산되지 않습니다. 모든 쿼럼 장치를 제거한 다음 다시 구성에 추가하면 올바른 쿼럼 투표 수를 다시 설정할 수 있습니다. 2 노드 클러스터의 경우 원래 쿼럼 장치를 제거했다가 다시 추가하기 전에 새 쿼럼 장치를 임시로 추가합니다. 그런 다음 임시 쿼럼 장치를 제거합니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 쿼럼 구성을 확인할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

● `clquorum` 명령을 사용하여 쿼럼 구성을 나열합니다.

```
% clquorum show +
```

예 61 쿼럼 구성 표시

```
% 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

▼ 쿼럼 장치를 복구하는 방법

오작동하는 쿼럼 장치를 교체하려면 이 절차를 사용합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 교체할 쿼럼 장치를 제거합니다.

주 - 제거할 장치가 마지막 쿼럼 장치인 경우 먼저 다른 디스크를 새 쿼럼 장치로 추가하는 것이 좋습니다. 이 단계를 수행하면 교체 절차 중에 장애가 발생할 경우 쿼럼 장치가 유효한지를 확인할 수 있습니다. 새 쿼럼 장치를 추가하려면 “[쿼럼 장치 추가](#)” [159]를 참조하십시오.

쿼럼 장치로서 디스크 장치를 제거하려면 [쿼럼 장치를 제거하는 방법](#) [165]을 참조하십시오.

2. 디스크 장치를 교체합니다.

디스크 장치를 교체하려면 하드웨어 설명서에서 디스크 외장 장치에 대한 절차를 참조하십시오. [Oracle Solaris Cluster Hardware Administration Manual](#)도 참조하십시오.

3. 교체된 디스크를 새 쿼럼 장치로 추가합니다.

디스크를 새 쿼럼 장치로 추가하려면 “[쿼럼 장치 추가](#)” [159]를 참조하십시오.

주 - 1단계에서 다른 쿼럼 장치를 추가한 경우 이제 제거해도 됩니다. 쿼럼 장치를 제거하려면 [쿼럼 장치를 제거하는 방법](#) [165]을 참조하십시오.

쿼럼 기본 시간 초과 변경

클러스터 재구성 중 쿼럼 작업 완료에 대해 기본 25초의 시간 초과가 적용됩니다. [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “[쿼럼 장치를 구성하는 방법](#)”에 설명된 지침에 따라 쿼럼 시간 초과 값을 높일 수 있습니다. 또한 시간 초과 값을 높이는 대신 다른 쿼럼 장치로 전환할 수도 있습니다.

추가 문제 해결 정보는 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “[쿼럼 장치를 구성하는 방법](#)”에서 확인할 수 있습니다.

주 - Oracle Real Application Clusters(Oracle RAC)에 대한 기본 퀴럼 시간 초과 값인 25초는 변경하지 **마십시오**. 특정 정보 분리(split-brain) 시나리오에서는 시간 초과 기간이 길어지면 VIP 리소스 시간이 초과되어 Oracle RAC VIP 페일오버가 실패할 수 있습니다.

사용 중인 퀴럼 장치에서 기본 25초 시간 초과를 준수하지 않는 경우 다른 퀴럼 장치를 사용하십시오.

Oracle Solaris Cluster 퀴럼 서버 관리

Oracle Solaris Cluster 퀴럼 서버에서는 공유 저장 장치가 아닌 퀴럼 장치를 제공합니다. 이 절에서는 Oracle Solaris Cluster 퀴럼 서버를 관리하는 절차에 대해 다음과 같이 설명합니다.

- “퀴럼 서버 소프트웨어 시작 및 중지” [175]
- 퀴럼 서버를 시작하는 방법 [175]
- 퀴럼 서버를 중지하는 방법 [176]
- “퀴럼 서버에 대한 정보 표시” [177]
- “오래된 퀴럼 서버 클러스터 정보 정리” [178]

Oracle Solaris Cluster 퀴럼 서버 설치 및 구성에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 퀴럼 서버 소프트웨어를 설치하고 구성하는 방법”을 참조하십시오.

퀴럼 서버 소프트웨어 시작 및 중지

이 절차에서는 Oracle Solaris Cluster 소프트웨어를 시작하고 중지하는 방법에 대해 설명합니다.

퀴럼 서버 구성 파일 `/etc/scqsd/scqsd.conf`의 내용을 사용자정의하지 않은 경우 기본적으로 이 절차에서는 단일 기본 퀴럼 서버를 시작하고 중지합니다. 기본 퀴럼 서버는 포트 9000에 바인드되고 퀴럼 정보용으로 `/var/scqsd` 디렉토리를 사용합니다.

퀴럼 서버 소프트웨어 설치에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 퀴럼 서버 소프트웨어를 설치하고 구성하는 방법”을 참조하십시오. 퀴럼 시간 초과 값 변경에 대한 자세한 내용은 “퀴럼 기본 시간 초과 변경” [174]을 참조하십시오.

▼ 퀴럼 서버를 시작하는 방법

1. Oracle Solaris Cluster 소프트웨어를 시작하려는 호스트에서 **root** 역할로 전환합니다.

2. **clquorumserver start** 명령을 사용하여 소프트웨어를 시작합니다.

```
# clquorumserver start quorumserver
```

quorumserver 쿼럼 서버를 식별합니다. 쿼럼 서버가 수신하는 포트 번호를 사용할 수 있습니다. 구성 파일의 인스턴스 이름을 제공한 경우 해당 이름을 대신 사용할 수 있습니다.

단일 쿼럼 서버를 시작하려면 인스턴스 이름 또는 포트 번호를 입력합니다. 여러 개의 쿼럼 서버를 구성한 경우 모든 쿼럼 서버를 시작하려면 + 피연산자를 사용합니다.

예 62 구성된 모든 쿼럼 서버 시작

다음 예에서는 구성된 모든 쿼럼 서버를 시작합니다.

```
# clquorumserver start +
```

예 63 특정 쿼럼 서버 시작

다음 예에서는 포트 번호 2000에서 수신하는 쿼럼 서버를 시작합니다.

```
# clquorumserver start 2000
```

▼ 쿼럼 서버를 중지하는 방법

1. Oracle Solaris Cluster 소프트웨어를 시작하려는 호스트에서 **root** 역할로 전환합니다.

2. **clquorumserver stop** 명령을 사용하여 소프트웨어를 중지합니다.

```
# clquorumserver stop [-d] quorumserver
```

-d 다음에 시스템을 부트할 때 쿼럼 서버가 시작되는지 여부를 제어합니다. *-d* 옵션을 지정하면 다음에 시스템을 부트할 때 쿼럼 서버가 시작되지 않습니다.

quorumserver 쿼럼 서버를 식별합니다. 쿼럼 서버가 수신하는 포트 번호를 사용할 수 있습니다. 구성 파일에서 인스턴스 이름을 제공한 경우 해당 이름을 대신 사용할 수 있습니다.

단일 쿼럼 서버를 중지하려면 인스턴스 이름 또는 포트 번호를 입력합니다. 여러 개의 쿼럼 서버를 구성한 경우 모든 쿼럼 서버를 중지하려면 + 피연산자를 사용합니다.

예 64 구성된 모든 쿼럼 서버 중지

다음 예에서는 구성된 모든 쿼럼 서버를 중지합니다.

```
# clquorumserver stop 2000
```

6장. 쿼럼 관리 177

```
--- 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
```

예 67 여러 개의 쿼럼 서버 구성 표시

다음 예에서는 3개의 쿼럼 서버 qs1, qs2 및 qs3의 구성 정보를 표시합니다.

```
# clquorumserver show qs1 qs2 qs3
```

예 68 실행 중인 모든 쿼럼 서버 구성 표시

다음 예에서는 실행 중인 모든 쿼럼 서버의 구성 정보를 표시합니다.

```
# clquorumserver show +
```

오래된 쿼럼 서버 클러스터 정보 정리

유형이 quorumserver인 쿼럼 장치를 제거하려면 How to Remove a Quorum Device에서 설명한 대로 [쿼럼 장치를 제거하는 방법 \[165\]](#) 명령을 사용합니다. 일반 작업 시 이 명령은 쿼럼 서버 호스트에 대한 쿼럼 서버 정보도 제거합니다. 하지만 클러스터에서 쿼럼 서버 호스트와의 통신이 끊긴 경우 쿼럼 장치를 제거하면 이 정보가 정리되지 않습니다.

쿼럼 서버 클러스터 정보는 다음 환경에서 유효하지 않게 됩니다.

- clquorum remove 명령을 사용하여 먼저 클러스터 쿼럼 장치를 제거하지 않고 클러스터의 서비스를 해제한 경우
- 쿼럼 서버 호스트가 중지된 동안 클러스터에서 quorum_server 유형의 쿼럼 장치를 제거한 경우



주의 - 유형이 quorumserver인 쿼럼 장치를 클러스터에서 아직 제거하지 않은 경우 이 절차를 사용하여 유효한 쿼럼 서버를 정리하면 클러스터 쿼럼이 손상될 수 있습니다.

▼ 쿼럼 서버 구성 정보를 정리하는 방법

시작하기 전에 [쿼럼 장치를 제거하는 방법 \[165\]](#)에서 설명한 대로 클러스터에서 쿼럼 서버 쿼럼 장치를 제거합니다.



주의 - 클러스터에서 이 쿼럼 서버를 계속 사용 중인 경우 이 절차를 수행하면 클러스터 쿼럼이 손상될 수 있습니다.

1. 쿼럼 서버 호스트에서 **root** 역할로 전환합니다.
2. **clquorumserver clear** 명령을 사용하여 구성 파일을 정리합니다.

```
# clquorumserver clear -c clustername -I clusterID quorumserver [-y]
```

-c *clustername* 이전에 쿼럼 서버를 쿼럼 장치로 사용한 클러스터 이름입니다. 클러스터 노드에서 **cluster show**를 실행하여 클러스터 이름을 얻을 수 있습니다.

-I *clusterID* 클러스터 ID입니다. 클러스터 ID는 8자리의 16진수입니다. 클러스터 노드에서 **cluster show**를 실행하여 클러스터 ID를 얻을 수 있습니다.

quorumserver 하나 이상의 쿼럼 서버에 대한 식별자입니다. 쿼럼 서버는 포트 이름 또는 인스턴스 이름으로 식별할 수 있습니다. 포트 번호는 쿼럼 서버와 통신하기 위해 클러스터 노드에서 사용되며 인스턴스 이름은 쿼럼 서버 구성 파일 `/etc/scqsd/scqsd.conf`에 지정됩니다.

-y **clquorumserver clear** 명령을 강제로 실행하여 확인 메시지를 먼저 표시하지 않고 구성 파일에서 클러스터 정보를 정리합니다. 오래된 클러스터 정보를 쿼럼 서버에서 제거하려는 경우에만 이 옵션을 사용합니다.

3. (옵션) 기타 쿼럼 장치가 이 서버 인스턴스에서 구성되지 않은 경우에는 쿼럼 서버를 중지합니다.

예 69 쿼럼 서버 구성에서 오래된 클러스터 정보 정리

이 예에서는 포트 9000을 사용하는 쿼럼 서버에서 이름이 `sc-cluster`로 지정된 클러스터 정보를 제거합니다.

```
# 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
```


클러스터 상호 연결 및 공용 네트워크 관리

이 장에서는 Oracle Solaris Cluster 상호 연결 및 공용 네트워크를 관리하는 소프트웨어 절차를 설명합니다.

클러스터 상호 연결 및 공용 네트워크 관리는 하드웨어 및 소프트웨어 절차로 구성됩니다. 일반적으로 클러스터를 처음 설치하고 구성할 때 모든 PNM(공용 네트워크 관리) 객체를 포함하여 클러스터 상호 연결과 공용 네트워크를 구성합니다.

PNM 객체는 IPMP(Internet Protocol network multipathing) 그룹, 트렁크와 DLMP(Datalink Multipathing) 링크 통합 및 링크 통합에 의해 직접 지원되는 VNIC를 포함합니다. Multipathing은 Oracle Solaris 11 OS와 함께 자동으로 설치되며 사용자가 사용으로 설정해야 합니다. 클러스터 상호 연결 네트워크 구성을 나중에 변경해야 할 경우에는 이 장에 있는 소프트웨어 절차를 사용할 수 있습니다.

클러스터에서 IP Network Multipathing 그룹을 구성하는 방법에 대한 내용은 “공용 네트워크 관리” [195] 절을 참조하십시오. IPMP에 대한 자세한 내용은 [Oracle Solaris 11.3의 TCP/IP 네트워크, IPMP 및 IP 터널 관리의 3장](#), “IPMP 관리”를 참조하십시오. 링크 통합에 대한 자세한 내용은 [Oracle Solaris 11.3의 네트워크 데이터 링크 관리의 2장](#), “링크 집계를 사용하여 고가용성 구성”을 참조하십시오.

이 장에서는 다음의 항목에 대한 정보 및 절차를 설명합니다.

- “클러스터 상호 연결 관리” [181]
- “공용 네트워크 관리” [195]

이 장에 있는 관련 절차에 대한 자세한 내용은 표 12. “작업 목록: 클러스터 상호 연결 관리”를 참조하십시오.

클러스터 상호 연결 및 공용 네트워크에 대한 배경 및 개요 정보는 [Oracle Solaris Cluster 4.3 Concepts Guide](#)를 참조하십시오.

클러스터 상호 연결 관리

이 절에서는 클러스터 전송 어댑터 및 클러스터 전송 케이블과 같은 클러스터 상호 연결을 재구성하는 절차를 제공합니다. 이 절차를 수행하려면 Oracle Solaris Cluster 소프트웨어를 설치해야 합니다.

대부분의 경우, `clsetup` 유틸리티를 사용하여 클러스터 상호 연결에 대한 클러스터 전송을 관리할 수 있습니다. 자세한 내용은 [clsetup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오. 모든 클러스터 상호 연결 명령은 전역 클러스터 노드에서 실행되어야 합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 이러한 작업 중 일부를 수행할 수도 있습니다. 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

클러스터 소프트웨어의 설치 절차는 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)를 참조하십시오. 클러스터 하드웨어 구성요소의 서비스 절차는 [Oracle Solaris Cluster Hardware Administration Manual](#)을 참조하십시오.

주 - 기본 포트 이름이 필요할 경우에는 일반적으로 클러스터 상호 연결 절차에서 기본 포트 이름을 사용하도록 선택할 수 있습니다. 기본 포트 이름은 케이블 끝에 있는 어댑터를 호스트하는 내부 노드 ID 번호와 동일합니다.

표 12 작업 목록: 클러스터 상호 연결 관리

작업	지침
<code>clsetup(1CL)</code> 을 사용하여 클러스터 전송 관리	클러스터 구성 유틸리티에 액세스하는 방법 [32]
<code>clinterconnect status</code> 를 사용하여 클러스터 상호 연결 상태 확인	클러스터 상호 연결 상태를 확인하는 방법 [183]
<code>clsetup</code> 를 사용하여 클러스터 전송 케이블, 전송 어댑터 또는 스위치 추가	클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치를 추가하는 방법 [184]
<code>clsetup</code> 를 사용하여 클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치 제거	클러스터 전송 케이블, 전송 어댑터, 전송 스위치를 제거하는 방법 [186]
<code>clsetup</code> 을 사용하여 클러스터 전송 케이블을 사용 가능하게 설정	클러스터 전송 케이블을 활성화하는 방법 [188]
<code>clsetup</code> 을 사용하여 클러스터 전송 케이블을 사용 불가능하게 설정	클러스터 전송 케이블을 비활성화하는 방법 [189]
전송 어댑터의 인스턴스 번호 결정	전송 어댑터의 인스턴스 번호를 결정하는 방법 [191]
IP 주소 또는 기존 클러스터의 주소 범위 변경	기존 클러스터의 개인 네트워크 주소 또는 주소 범위를 변경하는 방법 [192]

클러스터 상호 연결 동적 재구성

클러스터 상호 연결에서 동적 재구성(Dynamic Reconfiguration, DR) 작업을 완료하는 경우 몇 가지 문제를 고려해야 합니다.

- Oracle Solaris 동적 재구성 기능에 대해 문서화된 모든 요구사항, 절차 및 제한 사항은 운영체제의 작동이 정지된 경우를 제외하고는 Oracle Solaris Cluster 동적 재구성 지원에도 적용됩니다. 따라서 Oracle Solaris Cluster 소프트웨어와 함께 동적 재구성 기능을 사용하기 전에 Oracle Solaris 동적 재구성 기능에 대한 설명서를 검토하십시오. 특히 동적 재구성 연결 종료 작업 중 비네트워크 IO 장치에 영향을 주는 문제를 검토해야 합니다.

- Oracle Solaris Cluster 소프트웨어에서는 활성 개인 상호 연결 인터페이스에서 수행되는 동적 재구성 보드 제거 작업을 수행할 수 없습니다.
- 활성 클러스터 상호 연결에서 동적 재구성을 수행하려면 클러스터에서 활성 어댑터를 완전히 제거해야 합니다. clsetup 메뉴 또는 적절한 명령을 사용합니다.



주의 - Oracle Solaris Cluster 소프트웨어에서는 각 클러스터 노드에서 다른 모든 클러스터 노드에 대해 하나 이상의 경로가 작동하고 있어야 합니다. 다른 클러스터 노드에 대한 마지막 경로를 지원하는 독립 상호 연결 인터페이스를 비활성화하면 안 됩니다.

공용 네트워크 인터페이스에 대해 동적 재구성 작업을 수행하는 경우에는 다음 절차를 순서대로 완료하십시오.

표 13 작업 맵: 공용 네트워크 인터페이스 동적 재구성

작업	지침
1. 현재 작동하는 상호 연결에서 인터페이스 비활성화 및 제거	“공용 네트워크 인터페이스 동적 재구성” [197]
2. 공용 네트워크 인터페이스에 대한 동적 재구성 작업 수행	

▼ 클러스터 상호 연결 상태를 확인하는 방법

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

이 절차를 수행하기 위해 root 역할로 로그인할 필요는 없습니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터 상호 연결의 상태를 확인할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 상호 연결의 상태를 확인합니다.

```
% clinterconnect status
```

2. 일반적인 상태 메시지는 다음 표를 참조하십시오.

상태 메시지	설명 및 가능한 조치
Path online	경로가 현재 정상적으로 작동하고 있습니다. 조치를 취할 필요가 없습니다.
Path waiting	경로가 현재 초기화되고 있습니다. 조치를 취할 필요가 없습니다.

상태 메시지	설명 및 가능한 조치
Faulted	경로가 작동하지 않습니다. 경로가 대기 상태와 온라인 상태 사이에 있을 경우에는 이것이 일시적인 상태일 수 있습니다. <code>clinterconnect status</code> 명령을 다시 실행해도 계속 이 상태가 지속되면 교정 조치를 수행합니다.

예 70 클러스터 상호 연결 상태 확인

다음은 작동하는 클러스터 상호 연결의 상태를 표시하는 예입니다.

```
% 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
```

▼ 클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치를 추가하는 방법

클러스터 개인 전송 요구사항에 대한 자세한 내용은 [Oracle Solaris Cluster Hardware Administration Manual](#)의 “Interconnect Requirements and Restrictions”을 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 케이블, 전송 어댑터 및 개인 어댑터를 클러스터에 추가할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 물리적 클러스터 전송 케이블이 설치되어 있는지 확인합니다.
클러스터 전송 케이블 설치 절차는 [Oracle Solaris Cluster Hardware Administration Manual](#)을 참조하십시오.
2. 클러스터 노드에서 `root` 역할을 수행합니다.
3. `clsetup` 유틸리티를 시작합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.

4. 클러스터 상호 연결 메뉴를 표시하는 옵션에 대한 번호를 입력합니다.
5. 전송 케이블을 추가하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오.
6. 노드에 전송 어댑터를 추가하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오.
클러스터 상호 연결에 다음 어댑터를 하나를 사용하려면 각 클러스터 노드에 있는 `/etc/system` 파일에 관련 항목을 추가합니다. 다음에 시스템을 재부트하면 이 항목이 적용됩니다.

어댑터	항목
nge	set nge:nge_taskq_disable=1
e1000g	set e1000g:e1000g_taskq_disable=1

7. 전송 스위치를 추가하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오.
8. 클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치가 추가되었는지 확인합니다.

```
# clinterconnect show node:adapter,adapternode
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

예 71 클러스터 전송 케이블, 전송 어댑터 또는 전송 스위치 추가 확인

다음 예에서는 노드에 전송 케이블, 전송 어댑터 또는 전송 스위치 추가를 확인하는 방법을 보여 줍니다. 예제에는 DLPI(Data Link Provider Interface) 전송 유형에 대한 설정이 포함됩니다.

```
# 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
```

```

Adapter Property (dlpi_heartbeat_quantum): 1000
Adapter Property (nw_bandwidth): 80
Adapter Property (bandwidth): 70
Adapter Property (ip_address): 172.16.0.129
Adapter Property (netmask): 255.255.255.128
Adapter Port Names: 0
Adapter Port State (0): Enabled

# clinterconnect show phys-schost-1:hub2

=== Transport Switches ===
Transport Switch: hub2
Switch State: Enabled
Switch Type: switch
Switch Port Names: 1 2
Switch Port State(1): Enabled
Switch Port State(2): Enabled

```

다음 순서 클러스터 전송 케이블의 상호 연결 상태를 확인하려면 [클러스터 상호 연결 상태를 확인하는 방법 \[183\]](#)을 참조하십시오.

▼ 클러스터 전송 케이블, 전송 어댑터, 전송 스위치를 제거하는 방법

다음 절차에 따라 노드 구성에서 클러스터 전송 케이블, 전송 어댑터, 전송 스위치를 제거합니다. 케이블이 비활성화되어도 케이블의 두 종점은 계속 구성되어 있습니다. 어댑터가 전송 케이블에서 종점으로 계속 사용되는 경우에는 제거할 수 없습니다.



주의 - 각 클러스터 노드에서 다른 모든 클러스터 노드에 대하여 하나 이상의 전송 경로가 작동하고 있어야 합니다. 어떤 노드도 두 노드 사이가 끊어지면 안됩니다. 케이블을 비활성화하기 전에 항상 노드의 클러스터 상호 연결 상태를 확인하십시오. 여분의 연결이 가능한지 확인한 후에 케이블 연결을 사용 불가능하게 합니다. 즉, 다른 연결을 사용할 수 있는지 먼저 확인해야 합니다. 노드에서 작동하는 마지막 케이블까지 비활성화하면 노드가 클러스터 구성원에서 제외됩니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터에서 케이블, 전송 어댑터 및 개인 어댑터를 제거할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 노드에서 root 역할을 수행합니다.
2. 남은 클러스터 전송 경로의 상태를 확인합니다.

```
# clinterconnect status
```



주의 - 2 노드 클러스터에서 노드 하나를 제거하려고 할 때 path faulted와 같은 오류 메시지가 나타나면 문제가 있는지 조사한 후에 이 절차를 계속하십시오. 이러한 문제가 발생하면 노드 경로를 사용하지 못할 수도 있습니다. 남은 작동 경로를 제거하면 노드가 클러스터 구성원에서 제외되어 클러스터가 재구성될 수도 있습니다.

3. clsetup 유틸리티를 시작합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.

4. 클러스터 상호 연결 메뉴에 액세스하는 옵션에 대한 번호를 입력합니다.
5. 전송 케이블을 사용하지 않는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오. 적용할 수 있는 노드 이름, 어댑터 이름 및 스위치 이름을 알아야 합니다.
6. 전송 케이블을 제거하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오. 적용할 수 있는 노드 이름, 어댑터 이름 및 스위치 이름을 알아야 합니다.

주 - 물리적인 케이블을 제거할 경우에는 포트와 대상 장치 사이의 케이블 연결을 끊으십시오.

7. 노드에서 전송 어댑터를 제거하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오. 적용할 수 있는 노드 이름, 어댑터 이름 및 스위치 이름을 알아야 합니다.
노드에서 물리적 어댑터를 제거하려는 경우 [Oracle Solaris Cluster Hardware Administration Manual](#)에서 하드웨어 서비스 절차를 참조하십시오.
8. 전송 스위치를 제거하는 옵션에 대한 번호를 입력합니다.
지시에 따라 요청하는 정보를 입력하십시오. 적용할 수 있는 노드 이름, 어댑터 이름 및 스위치 이름을 알아야 합니다.

주 - 전송 케이블에서 포트를 종점으로 사용하고 있으면 스위치를 제거할 수 없습니다.

9. 케이블, 어댑터 또는 스위치가 제거되었는지 확인합니다.

```
# clinterconnect show node:adapter,adapternode
```

```
# clinterconnect show node:adapter
# clinterconnect show node:switch
```

이 명령을 실행했을 때 해당 노드에서 제거된 전송 케이블이나 어댑터가 나타나면 안 됩니다.

예 72 전송 케이블, 전송 어댑터 또는 전송 스위치 제거 확인

다음 예제에서는 전송 케이블, 전송 어댑터 또는 전송 스위치 제거를 확인하는 방법을 보여 줍니다.

```
# 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
```

▼ 클러스터 전송 케이블을 활성화하는 방법

이 옵션은 기존 클러스터 전송 케이블을 사용으로 설정하는 데 사용됩니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 케이블을 사용으로 설정할 수도 있습니다. Private Interconnects(개인 상호 연결), Cables(케이블), 케이블 번호를 차례로 눌러서 강조 표시한 후 Enable(사용)을 누릅니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 노드에서 root 역할을 수행합니다.

2. clsetup 유틸리티를 시작합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.

3. 클러스터 상호 연결 메뉴에 액세스하는 옵션에 대한 번호를 입력합니다.

4. 전송 케이블을 사용으로 설정하는 옵션에 해당하는 번호를 입력합니다.

화면의 지시를 따르십시오. 식별하려는 케이블 종점 중 하나의 노드와 어댑터 이름을 모두 입력해야 합니다.

5. 케이블이 활성화되었는지 확인합니다.

```
# clinterconnect show node:adapter,adaptnode
```

다음과 같이 출력됩니다.

```
# 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
```

▼ 클러스터 전송 케이블을 비활성화하는 방법

클러스터 상호 연결 경로를 일시적으로 종료하기 위해 클러스터 전송 케이블을 사용 안함으로 설정할 필요가 있습니다. 일시적인 종료는 클러스터 상호 연결 문제를 해결하거나 클러스터 상호 연결 하드웨어를 교체할 때 사용합니다.

케이블이 비활성화되어도 케이블의 두 종점은 계속 구성되어 있습니다. 어댑터가 전송 케이블에서 종점으로 계속 사용되는 경우에는 제거할 수 없습니다.



주의 - 각 클러스터 노드에서 다른 모든 클러스터 노드에 대하여 하나 이상의 전송 경로가 작동하고 있어야 합니다. 어떤 노드도 두 노드 사이가 끊어지면 안됩니다. 케이블을 비활성화하기 전에 항상 노드의 클러스터 상호 연결 상태를 확인하십시오. 여분의 연결이 가능한지 확인한 후에 케이블 연결을 사용 불가하게 합니다. 즉, 다른 연결을 사용할 수 있는지 먼저 확인해야 합니다. 노드에서 작동하는 마지막 케이블까지 비활성화하면 노드가 클러스터 구성원에서 제외됩니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 케이블을 사용 안함으로 설정할 수도 있습니다. Private Interconnects(개인 상호 연결), Cables(케이블), 케이블 번호를 차례로 눌러서 강조 표시한 후 Disable(사용 안함)을 누릅니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 노드에서 **root** 역할을 수행합니다.
2. 케이블을 비활성화하기 전에 클러스터 상호 연결의 상태를 확인합니다.

```
# clinterconnect status
```



주의 - 2 노드 클러스터에서 노드 하나를 제거하려고 할 때 “path faulted”와 같은 오류 메시지가 나타나면 문제가 있는지 조사한 후에 이 절차를 계속하십시오. 이러한 문제가 발생하면 노드 경로를 사용하지 못할 수도 있습니다. 남은 작동 경로를 제거하면 노드가 클러스터 구성원에서 제외되어 클러스터가 재구성될 수도 있습니다.

3. **clsetup** 유틸리티를 시작합니다.

```
# clsetup
```

주 메뉴가 표시됩니다.
4. Cluster Interconnect(클러스터 상호 연결) 메뉴에 액세스하는 옵션에 해당하는 번호를 입력합니다.
5. 전송 케이블을 사용하지 않는 옵션에 대한 번호를 입력합니다.

프롬프트에 따라 요청하는 정보를 제공합니다. 이 클러스터 상호 연결의 모든 구성요소가 비활성화됩니다. 식별하려는 케이블 중점 중 하나의 노드와 어댑터 이름을 모두 입력해야 합니다.
6. 케이블이 비활성화되었는지 확인합니다.

```
# clinterconnect show node:adapter,adapternode
```

다음과 같이 출력됩니다.

```
# 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
```

▼ 전송 어댑터의 인스턴스 번호를 결정하는 방법

clsetup 명령을 통해 올바른 전송 어댑터를 추가 및 제거하려면 전송 어댑터의 인스턴스 번호를 결정해야 합니다. 어댑터 이름은 어댑터 유형과 어댑터의 인스턴스 번호로 이루어져 있습니다.

1. 슬롯 번호를 기준으로 어댑터 이름을 찾습니다.

다음 화면은 하나의 예이므로 사용자의 하드웨어 내용과 일치하지 않을 수 있습니다.

```
# 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. 어댑터의 경로를 사용하여 어댑터의 인스턴스 번호를 찾습니다.

다음 화면은 하나의 예이므로 사용자의 하드웨어 내용과 일치하지 않을 수 있습니다.

```
# grep sci /etc/path_to_inst
"/xyz@1f,400/pci11c8,0@2" 0 "ttt"
"/xyz@1f,4000.pci11c8,0@4 "ttt"
```

3. 어댑터 이름과 슬롯 번호를 사용하여 어댑터의 인스턴스 번호를 찾습니다.

다음 화면은 하나의 예이므로 사용자의 하드웨어 내용과 일치하지 않을 수 있습니다.

```
# prtconf
...
xyz, instance #0
xyz11c8,0, instance #0
xyz11c8,0, instance #1
...
```

▼ 기존 클러스터의 개인 네트워크 주소 또는 주소 범위를 변경하는 방법

개인 네트워크 주소, 네트워크 주소 범위 또는 이 둘을 모두 변경하려면 다음 절차를 따릅니다. 명령줄을 사용하여 이 작업을 수행하려면 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

시작하기 전에 root 역할에 대한 원격 셸([rsh\(1M\)](#)) 또는 보안 셸([ssh\(1\)](#)) 액세스가 모든 클러스터 노드에 대해 사용으로 설정되어 있는지 확인합니다.

1. 각 클러스터 노드에서 다음 보조 단계를 수행하여 모든 클러스터 노드를 비클러스터 모드로 재부팅합니다.

- a. 비클러스터 모드로 시작하려면 클러스터 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

- b. `clnode evacuate` 및 `cluster shutdown` 명령을 사용하여 노드를 종료합니다.

`clnode evacuate` 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 장치 그룹을 전환합니다. 또한 이 명령은 지정된 노드에서 다음 기본 노드로 모든 리소스 그룹을 전환합니다.

```
# clnode evacuate node
# cluster shutdown -g0 -y
```

2. 한 노드에서 `clsetup` 유틸리티를 시작합니다.

비클러스터 모드로 실행되는 경우 `clsetup` 유틸리티에 비클러스터 모드 작업을 위한 주 메뉴가 표시됩니다.

3. Change Network Addressing and Ranges for the Cluster Transport(클러스터 전송에 대한 네트워크 주소 지정 및 범위 변경) 메뉴 항목을 선택합니다.

`clsetup` 유틸리티에 현재의 개인 네트워크 구성이 표시되고 이 구성을 변경할지 묻는 메시지가 표시됩니다.

4. 개인 네트워크 IP 주소 또는 IP 주소 범위를 변경하려면 `yes`를 입력하고 Enter 키를 누릅니다.

`clsetup` 유틸리티에 기본 개인 네트워크 IP 주소(172.16.0.0)가 표시되고 이 기본값을 사용할지를 묻는 메시지가 표시됩니다.

5. 해당 개인 네트워크 IP 주소를 변경하거나 사용합니다.

■ 기본 개인 네트워크 IP 주소를 사용하고 IP 주소 범위 변경을 진행하려면 `yes`를 입력하고 Enter 키를 누릅니다.

■ 기본 개인 네트워크 IP 주소를 변경하려면:

- a. **clsetup** 유틸리티에서 기본 주소를 사용할지를 물으면 그에 대한 응답으로 **no**를 입력한 후 Enter 키를 누릅니다.
clsetup 유틸리티에 새 개인 네트워크 IP 주소를 묻는 메시지가 표시됩니다.
 - b. **새 IP 주소를 입력하고 Enter 키를 누릅니다.**
clsetup 유틸리티에 기본 넷마스크가 표시되고 이 기본 넷마스크를 사용할지를 묻는 메시지가 표시됩니다.
6. **기본 개인 네트워크 IP 주소 범위를 변경하거나 사용합니다.**
기본 넷마스크는 255.255.240.0입니다. 이 기본 IP 주소 범위는 클러스터에서 최대 64개의 노드, 12개의 영역 클러스터 및 10개의 개인 네트워크를 지원합니다.
 - **기본 IP 주소 범위를 사용하려면 yes를 입력하고 Enter 키를 누릅니다.**
 - **IP 주소 범위를 변경하려면:**
 - a. **clsetup** 유틸리티에서 기본 주소 범위를 사용할지를 물으면 그에 대한 응답으로 **no**를 입력한 후 Enter 키를 누릅니다.
기본 넷마스크의 사용을 거부할 경우 클러스터에 구성할 노드, 개인 네트워크 및 영역 클러스터 수를 묻는 메시지가 clsetup 유틸리티에서 표시됩니다.
 - b. **클러스터에 구성할 노드, 개인 네트워크 및 영역 클러스터의 수를 제공합니다.**
clsetup 유틸리티는 이러한 숫자로 계산하여 두 개의 넷마스크를 제안합니다.
 - 첫번째 넷마스크는 지정한 노드, 개인 네트워크 및 영역 클러스터 수를 지원하는 최소 넷마스크입니다.
 - 두번째 넷마스크는 지정한 노드, 개인 네트워크 및 영역 클러스터 수의 두 배를 지원하여 향후 확대될 경우에도 수용할 수 있도록 합니다.
 - c. **계산된 넷마스크 중 하나를 지정하거나 예상 노드, 개인 네트워크 및 영역 클러스터 수를 지원하는 다른 넷마스크를 지정합니다.**
7. **clsetup** 유틸리티에서 업데이트를 진행할지를 물으면 그에 대한 응답으로 **yes**를 입력합니다.
8. **모두 완료되면 clsetup 유틸리티를 종료합니다.**
9. **각 클러스터 노드에 대해 다음 보조 단계를 완료하여 각 클러스터 노드를 클러스터 모드로 재부트합니다.**
 - a. **노드를 부트합니다.**
 - SPARC 기반 시스템에서는 다음 명령을 실행합니다.

ok boot

- x86 기반 시스템에서는 다음 명령을 실행합니다.

GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다.

10. 노드가 오류 없이 부트되고 온라인 상태인지 확인합니다.

```
# cluster status -t node
```

클러스터 상호 연결 문제 해결

이 절에서는 클러스터 전송 어댑터 및 전송 케이블과 같은 클러스터 상호 연결을 사용 안함으로 설정했다가 사용으로 설정하기 위한 문제 해결 절차를 제공합니다.

클러스터 전송 어댑터를 관리하기 위해 ipadm 명령을 사용하지 마십시오. ipadm disable-if 명령을 사용하여 전송 어댑터가 사용 안함으로 설정된 경우 clinterconnect 명령을 사용하여 전송 경로를 사용 안함으로 설정했다가 사용으로 설정해야 합니다.

이 절차를 수행하려면 Oracle Solaris Cluster 소프트웨어가 설치되어야 합니다. 이러한 명령은 전역 클러스터 노드에서 실행되어야 합니다.

▼ 클러스터 상호 연결을 사용으로 설정하는 방법

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 클러스터 상호 연결을 사용으로 설정할 수도 있습니다. Private Interconnects(개인 상호 연결), Cables(케이블), 케이블 번호를 차례로 눌러서 강조 표시한 후 Enable(사용)을 누릅니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 클러스터 상호 연결의 상태를 확인합니다.

```
% clinterconnect status
```

```
=== Cluster Transport Paths===
Endpoint1      Endpoint2      Status
-----
pnode1:net1    pnode2:net1    waiting
pnode1:net5    pnode2:net5    Path online
```

2. 클러스터 상호 연결 경로를 사용 안함으로 설정합니다.

- a. 클러스터 상호 연결 경로를 확인합니다.

```
% clinterconnect show | egrep -ie "cable.*pnode1"
```

```
Transport Cable: pnode1:net5,switch2@1
```

```
Transport Cable: pnode1:net1,switch1@1
```

- b. 클러스터 상호 연결 경로를 사용 안함으로 설정합니다.

```
% clinterconnect disable pnode1:net1,switch1@1
```

3. 클러스터 상호 연결 경로를 사용으로 설정합니다.

```
% clinterconnect enable pnode1:net1,switch1@1
```

4. 클러스터 상호 연결이 사용으로 설정되었는지 확인합니다.

```
% clinterconnect status
```

```
=== Cluster Transport Paths===
```

Endpoint1	Endpoint2	Status
-----	-----	-----
pnode1:net1	pnode2:net1	Path online
pnode1:net5	pnode2:net5	Path online

공용 네트워크 관리

Oracle Solaris Cluster 소프트웨어는 공용 네트워크에 대해 IPMP, 링크 통합 및 VNIC를 Oracle Solaris 소프트웨어로 구현할 수 있도록 지원합니다. 기본 공용 네트워크 관리 방법은 클러스터 환경과 비클러스터 환경 모두에서 동일합니다.

Multipathing은 Oracle Solaris 11 OS를 설치할 때 자동으로 설치되며 사용자가 사용으로 설정해야 합니다. Multipathing 관리 방법은 해당 Oracle Solaris OS 설명서에서 설명합니다. 하지만 Oracle Solaris Cluster 환경에서 IPMP, 링크 통합 및 VNIC를 관리하기 전에 준수해야 하는 지침을 검토하십시오.

IPMP에 대한 자세한 내용은 [Oracle Solaris 11.3의 TCP/IP 네트워크, IPMP 및 IP 터널 관리의 3 장, "IPMP 관리"](#)를 참조하십시오. 링크 통합에 대한 자세한 내용은 [Oracle Solaris 11.3의 네트워크 데이터 링크 관리의 2 장, "링크 집계를 사용하여 고가용성 구성"](#)을 참조하십시오.

클러스터에서 IP Network Multipathing 그룹을 관리하는 방법

클러스터에서 IPMP 절차를 수행하기 전에 다음 지침을 확인하십시오.

- SUNW.SharedAddress 네트워크 리소스를 사용하는 확장 가능한 서비스 리소스(리소스 유형 등록 파일에 SCALABLE=TRUE 설정)를 구성할 때 PNM은 SUNW.SharedAddress에 구성된

것 외에도 클러스터 노드의 모든 IPMP 그룹에서 IPMP 그룹 상태를 모니터하도록 구성할 수 있습니다. 이와 같이 구성하면 클러스터 노드와 동일한 서브넷에서 공존하는 네트워크 클라이언트의 서비스 가용성을 최대화하기 위해 클러스터 노드의 IPMP 그룹 중 하나가 실패한 경우 서비스를 다시 시작하고 페일오버할 수 있습니다. 예를 들면 다음과 같습니다.

```
# echo ssm_monitor_all > /etc/cluster/pnm/pnm.conf
```

노드를 재부트합니다.

- 이더넷 어댑터에 대한 local-mac-address? 변수의 값은 true여야 합니다.
 - 클러스터에서 프로브 기반 IPMP 그룹이나 링크 기반 IPMP 그룹을 사용할 수 있습니다. 프로브 기반 IPMP 그룹은 대상 IP 주소를 테스트하고 가용성을 손상시킬 수 있는 더 많은 조건을 인식하여 최상의 보호를 제공합니다.
- iSCSI 저장소를 쿼럼 장치로 사용하는 경우 프로브 기반 IPMP 장치가 올바르게 구성되었는지 확인합니다. iSCSI 네트워크가 클러스터 노드와 iSCSI 저장 장치만 포함하는 개인 네트워크이고 iSCSI 네트워크에 다른 호스트가 없는 경우, 클러스터 노드 중 하나만 작동 중지되면 프로브 기반 IPMP 방식이 고장날 수 있습니다. IPMP가 프로브할 다른 호스트가 iSCSI 네트워크에 없기 때문에 문제가 발생하므로, 클러스터에 하나의 노드만 남을 때 IPMP는 이를 네트워크 장애로 취급합니다. IPMP와 iSCSI 네트워크 어댑터가 오프라인이 되면 남은 노드는 iSCSI 저장소와 쿼럼 장치에 액세스할 수 없게 됩니다. 이 문제를 해결하기 위해 iSCSI 네트워크에 라우터를 추가할 수 있습니다. 그러면 클러스터 밖의 다른 호스트가 프로브에 응답할 수 있어서 IPMP와 네트워크 어댑터가 오프라인이 되는 것을 막을 수 있습니다. 다른 방법으로, 프로브 기반 페일오버 대신 링크 기반 페일오버로 IPMP를 구성할 수 있습니다.
- 공용 네트워크 구성에 하나 이상의 non-link-local IPv6 공용 네트워크 인터페이스가 없는 한, scinstall 유틸리티는 동일한 서브넷을 사용하는 클러스터의 각 공용 네트워크 어댑터 세트에 대해 여러 어댑터 IPMP 그룹을 자동으로 구성합니다. 이러한 그룹은 전이성 프로브를 사용하는 링크 기반 그룹입니다. 프로브 기반 실패 감지가 필요한 경우 테스트 주소를 추가할 수 있습니다.
 - 동일한 복수 경로 그룹에 포함된 모든 어댑터의 테스트 IP 주소가 하나의 IP 서브넷에 속해야 합니다.
 - 테스트 IP 주소는 가용성이 높지 않기 때문에 일반 응용 프로그램에서 사용하면 안됩니다.
 - 복수 경로 그룹의 이름 지정에 대한 제한 사항은 없습니다. 그러나 리소스 그룹을 구성하는 경우 netiflist 이름 지정 규칙에서 다중 경로 이름 뒤에 노드 ID 번호나 노드 이름을 사용해야 합니다. 예를 들어, 다중 경로 그룹의 이름이 sc_ipmp0이면 netiflist의 이름 지정은 sc_ipmp0@1 또는 sc_ipmp0@phys-schost-1이 됩니다. 여기서 어댑터는 nodeID가 1인 phys-schost-1 노드에 있습니다.
 - **if_mpadm(1M)** 명령을 사용하여 제거할 어댑터의 IP 주소를 그룹의 대체 어댑터로 스위치오버하기 전에는 IP Network Multipathing 그룹의 어댑터를 구성 해제하거나(배관을 끊거나) 중단하지 마십시오.
 - Oracle Solaris Cluster HA IP 주소가 연결된 IPMP 그룹에서 네트워크 인터페이스의 연결을 해제하거나 제거하지 마십시오. 이 IP 주소는 논리 호스트 이름 리소스 또는 공유 주소 리소스에 속할 수 있습니다. 하지만 ifconfig 명령을 사용하여 활성 인터페이스의 연결을 해제하면 이제 Oracle Solaris Cluster가 이 이벤트를 인식합니다. IPMP 그룹을 프

로세스에서 사용할 수 없게 되는 경우 리소스 그룹을 다른 정상 노드로 페일오버합니다. Oracle Solaris Cluster는 IPMP 그룹이 유효하지만 HA IP 주소가 누락된 경우 동일한 노드에서 리소스 그룹을 다시 시작할 수 있습니다. IPv4 연결 끊김이나 IPv6 연결 끊김과 같은 여러 가지 이유로 IPMP 그룹을 사용할 수 없게 됩니다. 자세한 내용은 [if_mpadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

- 어댑터를 각 복수 경로 그룹에서 먼저 제거하지 않고 다른 서브넷에 다시 연결하지 마십시오.
- 복수 경로 그룹에 대한 모니터링을 실행할 경우에도 어댑터에 대하여 논리 어댑터 작동을 실행할 수 있습니다.
- 클러스터에 있는 각 노드에 적어도 하나의 공용 네트워크 연결을 유지해야 합니다. 공용 네트워크 연결이 없으면 클러스터에 액세스할 수 없습니다.
- 클러스터에 있는 IP Network Multipathing 그룹의 상태를 보려면 `ipmpstat -g` 명령을 사용합니다. 자세한 내용은 [Oracle Solaris 11.3의 TCP/IP 네트워크, IPMP 및 IP 터널 관리의 3 장, “IPMP 관리”](#)를 참조하십시오.

클러스터 소프트웨어의 설치 절차는 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)를 참조하십시오. 공용 네트워킹 하드웨어 구성요소의 서비스 절차는 [Oracle Solaris Cluster Hardware Administration Manual](#)을 참조하십시오.

공용 네트워크 인터페이스 동적 재구성

클러스터의 공용 네트워크 인터페이스에서 동적 재구성(DR) 작업을 완료하는 경우 몇 가지 문제를 고려해야 합니다.

- Oracle Solaris 동적 재구성 기능에 대해 문서화된 모든 요구사항, 절차 및 제한 사항은 운영체제의 작동이 정지된 경우를 제외하고는 Oracle Solaris Cluster 동적 재구성 지원에도 적용됩니다. 따라서 Oracle Solaris Cluster 소프트웨어와 함께 동적 재구성 기능을 사용하기 전에 Oracle Solaris 동적 재구성 기능에 대한 설명서를 검토하십시오. 특히 동적 재구성 연결 종료 작업 중 비네트워크 IO 장치에 영향을 주는 문제를 검토해야 합니다.
- 동적 재구성 보드 제거 작업은 공용 네트워크 인터페이스가 활성 상태가 아닌 경우에만 가능합니다. 활성 공용 네트워크 인터페이스를 제거하기 전에, `if_mpadm` 명령을 사용하여 제거할 어댑터의 IP 주소를 다중 경로 그룹의 다른 어댑터로 전환하십시오. 자세한 내용은 [if_mpadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.
- 활성 네트워크 인터페이스로 작동하는 공용 네트워크 인터페이스 카드를 제대로 사용 안함으로 설정하지 않고 제거하면 Oracle Solaris Cluster가 작업을 거부하고 작업의 영향을 받는 인터페이스를 확인합니다.



주의 - 2개의 어댑터가 있는 다중 경로 그룹의 경우 사용 안함으로 설정된 네트워크 어댑터에 대해 동적 재구성 제거 작업을 수행할 때 남은 네트워크 어댑터에 장애가 발생하면 가용성이 영향을 받습니다. 동적 재구성 작업을 수행하는 동안 남은 어댑터를 페일오버할 수 없습니다.

공용 네트워크 인터페이스에 대해 동적 재구성 작업을 수행하는 경우에는 다음 절차를 순서대로 완료하십시오.

표 14 작업 맵: 공용 네트워크 인터페이스 동적 재구성

작업	지침
1. <code>if_mpadm</code> 명령을 사용하여 제거할 어댑터의 IP 주소를 다중 경로 그룹의 다른 어댑터로 전환	if_mpadm(1M) 매뉴얼 페이지 <i>Oracle Solaris 11.3의 TCP/IP 네트워크, IPMP 및 IP 터널 관리</i> 의 “하나의 IPMP 그룹에서 다른 IPMP 그룹으로 인터페이스를 이동하는 방법”
2. <code>ipadm</code> 명령을 사용하여 다중 경로 그룹에서 어댑터 제거	ipadm(1M) 매뉴얼 페이지 <i>Oracle Solaris 11.3의 TCP/IP 네트워크, IPMP 및 IP 터널 관리</i> 의 “IPMP 그룹에서 인터페이스를 제거하는 방법”
3. 공용 네트워크 인터페이스에 대한 동적 재구성 작업 수행	

◆◆◆ 8 장 8

클러스터 노드 관리

이 장에서는 클러스터에 노드를 추가하는 방법 및 노드를 제거하는 방법에 대한 지침을 제공합니다.

- “클러스터 또는 영역 클러스터에 노드 추가” [199]
- “클러스터 노드 복원” [202]
- “클러스터에서 노드 제거” [206]

클러스터 유지 보수 작업에 대한 자세한 내용은 [9장. 클러스터 관리](#)를 참조하십시오.

클러스터 또는 영역 클러스터에 노드 추가

이 절에서는 전역 클러스터 또는 영역 클러스터에 노드를 추가하는 방법에 대해 설명합니다. 전역 클러스터 노드에서 특정 영역 클러스터의 노드를 호스트하지 않는 한 영역 클러스터를 호스트하는 전역 클러스터의 한 노드에 새 영역 클러스터 노드를 만들 수 있습니다.

주 - 추가 노드는 참여 중인 클러스터와 동일한 버전의 Oracle Solaris Cluster 소프트웨어를 실행해야 합니다.

각 영역 클러스터 노드에 대해 IP 주소 및 NIC를 지정하는 것은 선택사항입니다.

주 - 각 영역 클러스터 노드에 대해 IP 주소를 구성하지 않으면 다음과 같은 두 가지 상황이 발생합니다.

1. 특정 영역 클러스터에서 영역 클러스터에 사용할 NAS 장치를 구성할 수 없습니다. 클러스터에서는 NAS 장치와 통신할 때 영역 클러스터 노드의 IP 주소를 사용하므로 IP 주소가 없으면 클러스터에서 NAS 장치 보호(fencing)를 지원하지 못합니다.
 2. 클러스터 소프트웨어가 모든 NIC에서 논리 호스트 IP 주소를 활성화합니다.
-

원래 영역 클러스터 노드에 IP 주소나 NIC가 지정되어 있지 않으면 새 영역 클러스터 노드에 대해 해당 정보를 지정할 필요가 없습니다.

이 장에서 `phys-schost#`는 전역 클러스터 프롬프트를 반영합니다. `clzonecluster` 대화식 셸 프롬프트는 `clzc:schost>`입니다.

다음 표에는 기존 클러스터에 노드를 추가하기 위해 수행하는 작업이 나열되어 있습니다. 표시된 순서대로 작업을 수행합니다.

표 15 작업 맵: 기존 전역 또는 영역 클러스터에 노드 추가

작업	지침
노드에 호스트 어댑터를 설치하고 기존 클러스터 상호 연결이 새 노드를 지원할 수 있는지 확인	Oracle Solaris Cluster Hardware Administration Manual
공유 저장소 추가	Oracle Solaris Cluster Hardware Administration Manual 의 지침에 따라 공유 저장소를 수동으로 추가합니다. Oracle Solaris Cluster Manager를 통해서도 영역 클러스터에 공유 저장 장치를 추가할 수 있습니다. Oracle Solaris Cluster Manager에서 영역 클러스터에 대한 페이지로 탐색하고 Solaris Resources(Solaris 리소스) 탭을 누릅니다. Oracle Solaris Cluster Manager 로그인 지침은 Oracle Solaris Cluster Manager에 액세스하는 방법 [278] 을 참조하십시오.
인증된 노드 목록에 노드 추가	<code>claccess allow -h node-being-added</code>
새 클러스터 노드에 소프트웨어를 설치하고 구성	Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서 의 2 장, “전역 클러스터 노드에 소프트웨어 설치”
기존 클러스터에 새 노드 추가	기존 클러스터 또는 영역 클러스터에 노드를 추가하는 방법 [200]
클러스터가 Oracle Solaris Cluster Geographic Edition 파트너십에 구성되어 있는 경우 새 노드를 구성의 활성 참여자로 구성	Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide 의 “How to Add a New Node to a Cluster in a Partnership”

▼ 기존 클러스터 또는 영역 클러스터에 노드를 추가하는 방법

기존 전역 클러스터 또는 영역 클러스터에 Oracle Solaris 호스트나 가상 시스템을 추가하기 전에 개인 클러스터 상호 연결에 대해 작동하는 물리적 연결을 포함하여 필요한 모든 하드웨어가 해당 노드에 올바르게 설치 및 구성되었는지 확인합니다.

하드웨어 설치 정보는 [Oracle Solaris Cluster Hardware Administration Manual](#) 또는 서버와 함께 제공된 하드웨어 설명서를 참조하십시오.

이 절차를 수행하면 시스템이 클러스터에 대한 권한이 있는 노드 목록에 노드 이름을 추가하여 클러스터에 자동으로 시스템을 설치할 수 있습니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문 형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 현재 전역 클러스터 구성원에서 root 역할로 전환합니다. 전역 클러스터의 한 노드에서 이러한 단계를 수행합니다.
2. 표 15. “작업 맵: 기존 전역 또는 영역 클러스터에 노드 추가”의 작업 맵에 나열된 필수 하드웨어 설치 및 구성 작업을 모두 올바르게 완료했는지 확인합니다.

3. 새 클러스터 노드에 소프트웨어를 설치하고 구성합니다.

[Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)에서 설명한 대로 새 노드의 설치 및 구성을 완료하려면 `scinstall` 유틸리티를 사용합니다.

4. 새 노드에 `scinstall` 유틸리티를 사용하여 클러스터에서 해당 노드를 구성합니다.

5. 수동으로 영역 클러스터에 노드를 추가하려면 Oracle Solaris 호스트와 가상 노드 이름을 지정해야 합니다.

또한 각 노드에서 공용 네트워크 통신에 사용할 네트워크 리소스를 지정해야 합니다. 다음 예에서 영역 이름은 `sczone`이고 `sc_ipmp0`은 IPMP 그룹 이름입니다.

```
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
```

노드 구성에 대한 자세한 지침은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “영역 클러스터 만들기 및 구성”을 참조하십시오.

6. 새로운 영역 클러스터 노드가 `solaris10` 브랜드 유형이 되고 Oracle Solaris Cluster 소프트웨어가 영역 클러스터에 설치되지 않은 경우 DVD 이미지 경로를 제공하고 소프트웨어를 설치하십시오.

```
# clzonecluster install-cluster -d dvd-image zone-cluster-name
```

7. 노드를 구성한 후에 클러스터 모드로 재부트하고 노드에 영역 클러스터를 설치합니다.

```
# clzonecluster install zone-cluster-name
```

8. 새 시스템이 클러스터에 추가되지 않도록 하려면 `clsetup` 유틸리티에서 클러스터가 새 시스템 추가에 대한 요청을 무시하도록 지시하는 옵션에 대한 번호를 입력합니다.

Enter 키를 누릅니다.

clsetup 프롬프트를 따릅니다. 이 옵션을 사용하면 새 시스템이 클러스터에 추가되려고 공용 네트워크를 통해 보내는 모든 요청을 클러스터가 무시합니다.

9. clsetup 유틸리티를 종료합니다.

참조 [clsetup\(1CL\)](#) 매뉴얼 페이지

클러스터 노드 추가에 대한 전체 작업 목록은 [표 15. “작업 맵: 기존 전역 또는 영역 클러스터에 노드 추가”](#), “작업 맵: 클러스터 노드 추가”를 참조하십시오.

기존 리소스 그룹에 노드를 추가하려면 [Oracle Solaris Cluster 4.3 데이터 서비스 계획 및 관리 설명서](#)를 참조하십시오.

클러스터 노드 복원

Unified Archives를 사용하여 아카이브와 정확히 동일하도록 클러스터 노드를 복원할 수 있습니다. 노드를 복원하기 전에 먼저 클러스터 노드에서 복구 아카이브를 만들어야 합니다. 복구 아카이브만 사용할 수 있습니다. 복제 아카이브를 통해서는 클러스터 노드를 복원할 수 없습니다. 복구 아카이브 만들기 지침은 아래 [1 단계를](#) 참조하십시오.

이 절차에서는 클러스터 이름, 노드 이름과 MAC 주소, Unified Archives에 대한 경로를 묻습니다. 각 지정된 아카이브에 대해 scinstall 유틸리티는 아카이브의 소스 노드 이름이 복원 중인 노드와 같은지 확인합니다. 통합 아카이브에서 클러스터 내 노드를 복원하는 것과 관련된 지침은 [통합 아카이브에서 노드를 복원하는 방법 \[202\]](#)을 참조하십시오.

▼ 통합 아카이브에서 노드를 복원하는 방법

이 절차에서는 자동 설치 프로그램 서버에 대해 scinstall 유틸리티의 대화식 형식을 사용합니다. 이미 AI 서버를 설정하고 Oracle Solaris Cluster 저장소에서 ha-cluster/system/install 패키지를 설치했어야 합니다. 아카이브의 노드 이름이 복원 중인 노드와 같아야 합니다.

이 절차에서 대화식 scinstall 유틸리티를 사용하려면 다음 지침을 준수하십시오.

- 대화식 scinstall 유틸리티에서는 사용자가 먼저 입력할 수 있습니다. 따라서 다음 메뉴 화면이 즉시 나타나지 않을 경우에 Enter 키를 두 번 이상 누르지 마십시오.
- 다른 지시가 없는 한 Ctrl-D를 눌러 관련 질문의 시작 부분이나 주 메뉴로 돌아갈 수 있습니다.
- 질문의 끝에 기본 응답이나 이전 세션에 대한 응답이 괄호([]) 안에 표시됩니다. Enter 키를 누르면 별도의 입력 없이 괄호 안의 응답을 선택할 수 있습니다.

1. 전역 클러스터의 노드에서 root 역할을 수행하고 복구 아카이브를 만듭니다.

```
phys-schost# archiveadm create -r archive-location
```

아카이브를 만들 때는 공유 저장소에 있는 ZFS 데이터 세트를 제외합니다. 공유 저장소의 데이터를 복원하려면 기존 방법을 사용합니다.

archiveadm 명령 사용에 대한 자세한 내용은 [archiveadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

2. 자동 설치 프로그램 서버에 로그인하고 root 역할을 수행합니다.

3. **scinstall** 유틸리티를 시작합니다.

```
phys-schost# scinstall
```

4. 클러스터를 복원할 옵션 번호를 입력합니다.

```
*** 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
```

비보안 AI 서버 설치를 사용하여 클러스터 노드를 복원하려면 옵션 1을 선택합니다. 보안 AI 서버 설치를 사용하여 클러스터 노드를 복원하려면 옵션 2를 선택합니다.

Custom Automated Installer(사용자정의 자동 설치 프로그램) 메뉴나 Custom Secure Automated Installer(사용자정의 보안 자동 설치 프로그램) 메뉴가 표시됩니다.

5. **Restore Cluster Nodes from Unified Archives(Unified Archives에서 클러스터 노드 복원)**에 해당하는 옵션 번호를 입력합니다.

Cluster Name(클러스터 이름) 화면이 표시됩니다.

6. **복원할 노드가 포함된 클러스터 이름을** 입력합니다.

Cluster Nodes(클러스터 노드) 화면이 표시됩니다.

7. **Unified Archives에서 복원할 클러스터 노드의 이름을** 입력합니다.

라인별로 노드 이름 하나를 입력합니다. 완료되면 Control-D 키를 누르고 yes를 입력하고 Return 키를 눌러 목록을 확인합니다. 클러스터의 모든 노드를 복원하려면 모든 노드를 지정합니다.

scinstall 유틸리티가 노드의 MAC 주소를 찾을 수 없을 때 프롬프트가 표시되면 각 주소를 입력합니다.

8. 복구 아카이브의 전체 경로를 입력합니다.

노드 복원에 사용되는 아카이브는 반드시 복구 아카이브여야 합니다. 동일한 노드에서 특정 노드 복원에 사용하는 아카이브 파일을 만들어야 합니다. 복원할 클러스터 노드마다 이 단계를 반복합니다.

9. **scinstall** 유틸리티가 이 AI 서버에서 클러스터 노드를 설치하는 데 필요한 구성을 수행하도록 각 노드에 대해 선택한 옵션을 확인합니다.

또한 유틸리티는 DHCP 서버에서 DHCP 매크로를 추가하기 위한 지침을 인쇄하고 SPARC 노드에 대한 보안 키를 추가하거나 지웁니다(보안 설치를 선택한 경우). 해당 지침을 따릅니다.

10. (선택사항) 대상 장치를 사용자정의하려면 각 노드에 대해 AI 매니페스트를 업데이트합니다. AI 매니페스트는 다음 디렉토리에 있습니다.

```
/var/cluster/logs/install/autoscinstall.d/ \
cluster-name/node-name/node-name_aimanifest.xml
```

a. 대상 장치를 사용자정의하려면 매니페스트 파일에서 **target** 요소를 업데이트합니다.

설치 대상 장치를 찾기 위해 지원되는 조건의 사용 방법에 따라 매니페스트 파일에서 target 요소를 업데이트합니다. 예를 들어, disk_name 하위 요소를 지정할 수 있습니다.

주 - scinstall에서는 매니페스트 파일의 기존 부트 디스크가 대상 장치가 된다고 가정합니다. 대상 장치를 사용자정의하려면 매니페스트 파일에서 target 요소를 업데이트합니다. 자세한 내용은 [Oracle Solaris 11.3 시스템 설치](#)의 제 III 부, “설치 서버를 사용하여 설치,” 및 [ai_manifest\(4\)](#) 매뉴얼 페이지를 참조하십시오.

b. 각 노드에 대해 **installadm** 명령을 실행합니다.

```
# 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
```

SPARC 및 i386은 클러스터 노드의 구조입니다.

11. 클러스터 관리 콘솔을 사용하고 있는 경우 클러스터의 각 노드에 대해 콘솔 화면을 표시합니다.

■ **pconsole** 소프트웨어를 관리 콘솔에 설치 및 구성한 경우 **pconsole** 유틸리티를 사용하여 개별 콘솔 화면을 표시합니다.

root 역할로 다음 명령을 사용하여 pconsole 유틸리티를 시작합니다.

```
adminconsole# pconsole host[:port] [...] &
```

pconsole 유틸리티를 사용하면 사용자 입력을 모든 개별 콘솔 창으로 동시에 전송할 수 있는 마스터 창도 열립니다.

■ **pconsole** 유틸리티를 사용하지 않는 경우 각 노드의 콘솔에 개별적으로 연결합니다.

12. 각 노드를 종료 후 부트하여 AI 설치를 시작합니다.

Oracle Solaris 소프트웨어가 기본 구성으로 설치됩니다.

주 - Oracle Solaris 설치를 사용자정의하려면 이 방법을 사용할 수 없습니다. Oracle Solaris 대화식 설치를 선택하면 자동 설치 프로그램이 무시되고 Oracle Solaris Cluster 소프트웨어가 설치 및 구성되지 않습니다.

설치 중 Oracle Solaris를 사용자정의하려면 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris 소프트웨어를 설치하는 방법”에 설명된 지침을 따르고 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 소프트웨어 패키지를 설치하는 방법”에 설명된 지침에 따라 클러스터를 설치 및 구성하십시오.

■ SPARC:

a. 각 노드를 종료합니다.

```
phys-schost# cluster shutdown -g 0 -y
```

b. 다음 명령으로 노드를 부트합니다.

```
ok boot net:dhcp - install
```

주 - 명령의 양쪽에 대시(-)를 표시하고 대시의 앞뒤를 한 칸씩 띄우십시오.

■ x86

a. 노드를 재부트합니다.

```
# reboot -p
```

b. PXE 부트 중 Ctrl-N을 누릅니다.

GRUB 메뉴가 표시됩니다.

c. Automated Install(자동 설치) 항목을 즉시 선택합니다.

주 - Automated Install(자동 설치) 항목을 20초 내에 선택하지 않으면 기본 대화식 텍스트 설치 프로그램 방법을 사용해서 설치가 진행되어 Oracle Solaris Cluster 소프트웨어가 설치 및 구성되지 않습니다.

설치가 완료되면 각 노드가 자동으로 재부트되어 클러스터가 결합됩니다. 노드는 아카이브가 만들어졌을 때와 동일한 상태로 복원됩니다. Oracle Solaris Cluster

설치 출력은 각 노드의 `/var/cluster/logs/install/sc_ai_config.log` 파일에 기록됩니다.

13. 한 노드에서 모든 노드가 클러스터를 결합했는지 확인합니다.

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

다음과 비슷한 결과가 출력됩니다:

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

```
--- Node Status ---
```

Node Name	Status
phys-schost-1	Online
phys-schost-2	Online
phys-schost-3	Online

자세한 내용은 [clnode\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

클러스터에서 노드 제거

이 절에서는 전역 클러스터 또는 영역 클러스터에서 노드를 제거하는 방법에 대한 지침을 제공합니다. 전역 클러스터에서 특정 영역 클러스터를 제거할 수도 있습니다. 다음 표는 기존 클러스터에서 노드를 제거하기 위해 수행하는 작업의 목록입니다. 표시된 순서대로 작업을 수행합니다.



주의 - RAC 구성에서 이 절차만 사용하여 노드를 제거하면 재부트 중에 노드 패닉이 발생할 수 있습니다. RAC 구성에서 노드를 제거하는 방법에 대한 지침은 [Oracle Real Application Clusters용 Oracle Solaris Cluster 데이터 서비스 설명서](#)의 “선택한 노드에서 Oracle RAC용 지원을 제거하는 방법”을 참조하십시오. 프로세스를 완료한 후 RAC 구성에 대한 노드를 제거하고 아래의 해당하는 단계를 따르십시오.

표 16 작업 맵: 노드 제거

작업	지침
제거할 노드에서 모든 리소스 그룹과 장치 그룹을 이동합니다. 영역 클러스터가 있는 경우 영역 클러스터에 로그인하고 제거할 물리적 노드에서 영역 클러스터 노드를 제외합니다. 그런 다음, 물리적 노드를 중단하기 전에 영역 클러스터에서 노드를 제거합니다.	<code>clnode evacuate node</code> 영역 클러스터에서 노드를 제거하는 방법 [207]
영향을 받는 물리적 노드가 이미 실패한 경우 클러스터에서 해당 노드를 제거하면 됩니다.	
허용된 호스트를 검사하여 노드를 제거할 수 있는지 확인합니다.	<code>claccess show</code> <code>claccess allow -h node-to-remove</code>

작업	지침
claccess show 명령으로 노드가 나열되지 않을 경우 노드를 제거할 수 없습니다. 노드에 클러스터 구성 액세스 권한을 부여합니다.	
모든 장치 그룹에서 노드를 제거합니다.	장치 그룹에서 노드를 제거하는 방법 (Solaris Volume Manager) [127]
제거할 노드에 연결된 모든 쿼럼 장치를 제거합니다.	쿼럼 장치를 제거하는 방법 [165]
2 노드 클러스터에서 노드를 제거하는 경우 이 단계는 선택사항입니다.	클러스터에서 마지막 쿼럼 장치를 제거하는 방법 [166]
다음 단계에서 저장 장치를 제거하기 전에 쿼럼 장치를 제거해야 하지만 이후에 바로 다시 쿼럼 장치를 추가할 수 있습니다.	
제거할 노드를 비클러스터 모드로 전환합니다.	노드를 유지 보수 상태로 전환하는 방법 [227]
클러스터 소프트웨어 구성에서 노드를 제거합니다.	클러스터 소프트웨어 구성에서 노드를 제거하는 방법 [208]
(선택사항) 클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거합니다.	클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법 [231]

▼ 영역 클러스터에서 노드를 제거하는 방법

노드를 중지하고 제거하고 구성에서 노드를 제거하면 영역 클러스터에서 노드를 제거할 수 있습니다. 나중에 노드를 다시 영역 클러스터에 추가하려면 [표 15. “작업 맵: 기존 전역 또는 영역 클러스터에 노드 추가”](#)의 지침을 따르십시오. 이러한 단계는 대부분 전역 클러스터 노드에서 수행됩니다.

주 - 또한 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서도 영역 클러스터 노드를 종료할 수 있지만 영역 클러스터 노드를 제거할 수 없습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 전역 클러스터 노드에서 root 역할을 수행합니다.
2. 노드 및 해당 영역 클러스터를 지정하여 제거할 영역 클러스터 노드를 종료합니다.

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

영역 클러스터 내에서 clnode evacuate 및 shutdown 명령을 사용할 수도 있습니다.

3. 영역 클러스터의 모든 리소스 그룹에서 노드를 제거합니다.

```
phys-schost# clresourcegroup remove-node -n zone-hostname -z zone-cluster-name rg-name
```

단계 2의 주에 설명된 절차를 수행하는 경우 리소스 그룹이 자동으로 제거되므로 이 단계를 건너뛸 수도 있습니다.

4. 영역 클러스터 노드를 제거합니다.

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

5. 구성에서 영역 클러스터 노드를 제거합니다.

다음 명령을 사용합니다.

```
phys-schost# clzonecluster configure zone-cluster-name
clzc:sczone> remove node physical-host=node
clzc:sczone> exit
```

주 - 제거할 영역 클러스터 노드가 액세스할 수 없거나 클러스터를 결합할 수 없는 시스템에 상주하는 경우 clzonecluster 대화식 셸을 사용하여 노드를 제거하십시오.

```
clzc:sczone> remove -F node physical-host=node
```

이 방법으로 마지막 영역 클러스터 노드를 제거하는 경우 영역 클러스터를 완전히 삭제할지 묻는 메시지가 표시됩니다. 완전히 삭제하지 않도록 선택하는 경우 마지막 노드가 제거되지 않습니다. 이 삭제 작업의 결과는 clzonecluster delete -F zone-cluster-name을 사용하는 것과 동일합니다.

6. 영역 클러스터에서 해당 노드가 제거되었는지 확인합니다.

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

▼ 클러스터 소프트웨어 구성에서 노드를 제거하는 방법

다음 절차를 수행하여 전역 클러스터에서 노드를 제거합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 이 절차를 계속 진행하기 전에 모든 리소스 그룹, 디스크 장치 그룹 및 쿼럼 장치 구성에서 노드를 제거하고 유지 보수 상태로 만들었는지 확인하십시오.
2. 제거할 노드에서 solaris.cluster.modify RBAC 권한 부여를 제공하는 역할로 전환합니다. 전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.
3. 제거할 전역 클러스터 노드를 비클러스터 모드로 부트합니다.

영역 클러스터 노드의 경우 이 단계를 수행하기 전에 [영역 클러스터에서 노드를 제거하는 방법 \[207\]](#)의 지침에 따릅니다.

■ SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot -x
```

- x86 기반 시스템에서는 다음 명령을 실행합니다.

```
shutdown -g -y -i0
```

Press any key to continue

- a. GRUB 메뉴에서 화살표 키를 사용하여 적절한 Oracle Solaris 항목을 선택하고 **e**를 입력하여 해당 명령을 편집합니다.

GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”](#)를 참조하십시오.

- b. 부트 매개 변수 화면에서 화살표 키를 사용하여 커널 항목을 선택하고 **e**를 입력하여 항목을 편집합니다.

- c. 명령에 **-x**를 추가하여 시스템 부트를 비클러스터 모드로 지정합니다.

```
[ 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. Enter 키를 눌러 변경사항을 수락하고 부트 매개 변수 화면으로 복귀합니다.

화면에 편집된 명령이 표시됩니다.

- e. **b**를 입력하여 비클러스터 모드로 노드를 부트합니다.

커널 부트 매개 변수 명령에 대한 변경사항은 시스템을 재부트하면 사라집니다. 다음에 노드를 재부트하면 클러스터 모드로 부트됩니다. 비클러스터 모드로 부트하려면, 이러한 단계를 다시 실행하여 **-x** 옵션을 커널 부트 매개 변수 명령에 추가합니다.

4. 클러스터에서 노드를 삭제합니다.

- a. 활성 노드에서 다음 명령을 실행합니다.

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

rg_system=true인 리소스 그룹이 있는 경우 rg_system=false로 변경해야 clnode clear -F 명령이 성공합니다. clnode clear -F를 실행한 후에 리소스 그룹을 rg_system=true로 다시 재설정합니다.

- b. 제거할 노드에서 다음 명령을 실행합니다.

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

주 - 제거할 노드를 사용할 수 없거나 더 이상 부트할 수 없는 경우, 활성 클러스터 노드에서 다음 명령을 실행합니다.

```
# clnode clear -F node-to-be-removed
```

`clnode status nodename`을 실행하여 노드 제거를 확인합니다.

클러스터의 마지막 노드를 제거하는 경우 클러스터에 활성 노드가 남아 있지 않은 상태로 해당 노드가 비클러스터 모드에 있어야 합니다.

5. 다른 클러스터 노드에서 노드 제거를 확인합니다.

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

6. 노드 제거를 완료합니다.

- 제거된 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하려는 경우 [클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법 \[231\]](#)으로 이동합니다.

클러스터에서 노드를 제거하는 동시에 Oracle Solaris Cluster 소프트웨어를 제거하도록 선택할 수도 있습니다. Oracle Solaris Cluster 파일이 없는 디렉토리로 변경하고 `scinstall -r`을 입력합니다.

- 제거된 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하지 않으려는 경우에는 하드웨어 연결을 제거해서 클러스터에서 노드를 물리적으로 제거할 수 있습니다.

지침은 [Oracle Solaris Cluster Hardware Administration Manual](#)을 참조하십시오.

예 73 클러스터 소프트웨어 구성에서 노드 제거

이 예에서는 클러스터에서 노드(phys-schost-2)를 제거하는 방법을 보여 줍니다. `clnode remove` 명령은 클러스터(phys-schost-2)에서 제거할 노드에서 비클러스터 모드로 실행됩니다.

클러스터에서 노드 제거:

```
phys-schost-2# clnode remove
phys-schost-1# clnode clear -F phys-schost-2
```

노드 제거 확인:

```
phys-schost-1# clnode status
-- Cluster Nodes --
Node name      Status
-----
Cluster node:  phys-schost-1  Online
```

참조 제거된 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하려면 [클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법 \[231\]](#)을 참조하십시오.

하드웨어 절차는 [Oracle Solaris Cluster Hardware Administration Manual](#)을 참조하십시오.

클러스터 노드 제거에 대한 전체 작업 목록은 표 16. “작업 맵: 노드 제거”를 참조하십시오.

기존 클러스터에 노드를 추가하려면 [기존 클러스터 또는 영역 클러스터에 노드를 추가하는 방법 \[200\]](#)을 참조하십시오.

▼ 2 노드보다 많은 연결이 있는 클러스터에서 어레이와 단일 노드 사이의 연결을 제거하는 방법

3 노드 또는 4 노드 연결이 있는 클러스터에서 저장소 어레이를 단일 클러스터 노드로부터 분리하려면 이 절차를 사용합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 제거할 저장소 어레이에 연결된 데이터베이스 테이블, 데이터 서비스 및 볼륨을 모두 백업합니다.
2. 연결을 끊을 노드에서 실행되는 리소스 그룹과 장치 그룹을 확인합니다.

```
phys-schost# clresourcegroup status
phys-schost# cldevicegroup status
```

3. 필요한 경우 연결을 끊을 노드에서 리소스 그룹과 장치 그룹을 모두 다른 노드로 이동합니다.



주의 (SPARC만 해당) - 클러스터에서 Oracle RAC 소프트웨어가 실행되고 있는 경우, 그룹을 노드에서 이동하기 전에 노드에서 실행되고 있는 Oracle RAC 데이터베이스 인스턴스를 종료합니다. 지침은 *Oracle Database Administration Guide*를 참조하십시오.

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

clnode evacuate 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 장치 그룹을 전환합니다. 또한 이 명령은 지정된 노드에서 다음 기본 노드로 모든 리소스 그룹을 전환합니다.

4. 장치 그룹을 유지 보수 상태로 만듭니다.

장치 그룹을 유지 보수 상태로 전환하는 절차는 [노드를 유지 보수 상태로 전환하는 방법 \[227\]](#)을 참조하십시오.

5. 장치 그룹에서 노드를 제거합니다.

원시 디스크를 사용하는 경우 cldevicegroup(1CL) 명령을 사용하여 장치 그룹을 제거합니다.

6. **HASStoragePlus** 리소스가 포함된 각 리소스 그룹에 대해 리소스 그룹의 노드 목록에서 노드를 제거합니다.

```
phys-schost# clresourcegroup remove-node -n node + | resourcegroup
```

리소스 그룹의 노드 목록 변경에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 데이터 서비스 계획 및 관리 설명서](#)를 참조하십시오.

주 - clresourcegroup 명령을 실행할 때 리소스 유형, 리소스 그룹 및 리소스 등록 정보의 이름은 대소문자를 구분합니다.

7. 제거하는 저장소 어레이가 노드에 연결된 마지막 저장소 어레이면 이 저장소 어레이에 연결된 허브 또는 스위치와 노드 사이의 광섬유 케이블 연결을 끊습니다.
그렇지 않으면 이 단계를 생략하십시오.
8. 연결을 끊을 노드에서 호스트 어댑터를 제거하려는 경우 해당 노드의 연결을 끊고 전원을 끕니다.
연결을 끊을 노드에서 호스트 어댑터를 제거하려는 경우 [11단계](#)로 건너뛰십시오.
9. 노드에서 호스트 어댑터를 제거합니다.
호스트 어댑터를 제거하는 절차는 해당 노드에 대한 설명서를 참조하십시오.
10. 노드를 부트하지 않고 노드의 전원을 켵니다.
11. Oracle RAC 소프트웨어가 설치된 경우, 연결을 끊을 노드에서 Oracle RAC 소프트웨어 패키지를 제거합니다.

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



주의 (SPARC만 해당) - 연결을 끊은 노드에서 Oracle RAC 소프트웨어를 제거하지 않을 경우, 노드가 클러스터에 다시 포함될 때 해당 노드는 패닉 상태가 되어 데이터 가용성이 손실될 수 있습니다.

12. 클러스터 모드로 노드를 부트합니다.
- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot
```
 - x86 기반 시스템에서는 다음 명령을 실행합니다.
GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다.
13. 노드에서 **/devices** 및 **/dev** 항목을 업데이트하여 장치 이름 공간을 업데이트합니다.

```
phys-schost# devfsadm -C  
phys-schost# cldevice refresh
```

14. 장치 그룹을 다시 온라인으로 전환합니다.

장치 그룹을 온라인 상태로 전환하는 방법에 대한 자세한 내용은 [노드를 유지 보수 상태에서 해제하는 방법 \[229\]](#)을 참조하십시오.

▼ 오류 메시지를 수정하는 방법

클러스터 노드 제거 절차를 수행하는 동안 발생한 오류 메시지를 수정하려면 다음 절차를 수행합니다.

1. 노드를 전역 클러스터에 다시 결합합니다.

전역 클러스터에서만 이 절차를 수행합니다.

```
phys-schost# boot
```

2. 노드가 클러스터에 연결되었습니까?

- 연결되지 않은 경우 [2b단계](#)를 진행합니다.
- 연결되었으면 다음 단계를 수행하여 장치 그룹에서 노드를 제거합니다.

a. 노드가 클러스터에 다시 연결되면 나머지 장치 그룹 또는 그룹에서 노드를 제거합니다.
[모든 장치 그룹에서 노드를 제거하는 방법 \[126\]](#)의 절차를 수행합니다.

b. 모든 장치 그룹에서 노드를 제거한 후 클러스터 노드에서 [Oracle Solaris Cluster 소프트웨어를 제거하는 방법 \[231\]](#)으로 돌아가 해당 절차를 반복합니다.

3. 노드가 클러스터를 다시 결합할 수 없는 경우 노드의 `/etc/cluster/ccr` 파일을 원하는 다른 이름(예: `ccr.old`)으로 변경합니다.

```
# mv /etc/cluster/ccr /etc/cluster/ccr.old
```

4. 클러스터 노드에서 [Oracle Solaris Cluster 소프트웨어를 제거하는 방법 \[231\]](#)으로 돌아가 해당 절차를 반복합니다.

클러스터 관리

이 장에서는 전체 전역 클러스터 또는 영역 클러스터에 영향을 주는 관리 절차에 대해 설명합니다.

- “클러스터 관리 개요” [215]
- “영역 클러스터 관리 작업 수행” [243]
- “테스트 목적으로 사용할 수 있는 문제 해결 절차” [253]

클러스터에서의 노드 추가 또는 제거에 대한 자세한 내용은 [8장. 클러스터 노드 관리](#)를 참조하십시오.

클러스터 관리 개요

이 절에서는 전체 전역 클러스터 또는 영역 클러스터에 대해 관리 작업을 수행하는 방법에 대해 설명합니다. 다음 표에는 해당 관리 작업과 관련 절차가 나열되어 있습니다. 일반적으로 전역 영역에서 클러스터 관리 작업을 수행합니다. 영역 클러스터를 관리하려면 영역 클러스터를 호스트할 시스템이 한 개 이상 클러스터 모드로 작동하고 있어야 합니다. 모든 영역 클러스터 노드가 작동되어 실행 중일 필요는 없습니다. Oracle Solaris Cluster는 현재 클러스터에 없는 노드가 클러스터를 다시 결합할 때 구성 변경사항을 재생합니다.

주 - 기본적으로 전원 관리는 사용 안함으로 설정되어 있으므로 클러스터에 영향을 주지 않습니다. 단일 노드 클러스터에 대한 전원 관리를 활성화하면 클러스터가 계속 실행되지만 몇 초 동안 사용할 수 없게 될 수 있습니다. 전원 관리 기능에서는 노드를 종료하려고 하지만 종료되지 않습니다.

이 장에서 `phys-schost#`는 전역 클러스터 프롬프트를 반영합니다. `clzonecluster` 대화식 셸 프롬프트는 `clzc:schost>`입니다.

표 17 작업 목록: 클러스터 관리

작업	지침
클러스터에서 노드 추가 또는 제거	8장. 클러스터 노드 관리
클러스터 이름 변경	클러스터 이름을 변경하는 방법 [216]

작업	지침
노드 ID 및 해당 노드 이름 나열	노드 ID를 노드 이름에 매핑하는 방법 [217]
클러스터에 새 노드 추가 허용 또는 거부	새 클러스터 노드에 대한 인증 작업 방법 [218]
NTP를 사용하여 클러스터의 시간 변경	클러스터에서 시간을 설정하는 방법 [220]
노드를 종료하여 OpenBoot PROM ok 프롬프트(SPARC 기반 시스템) 또는 Press any key to continue 메시지(x86 기반 시스템의 GRUB 메뉴) 표시	노드에서 OBP(OpenBoot PROM)를 표시하는 방법 [221]
개인 호스트 이름 추가 또는 변경	노드 개인 호스트 이름을 변경하는 방법 [222] Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 전역 클러스터 또는 영역 클러스터에 논리 호스트 이름을 추가할 수 있습니다. Tasks(작업)를 누른 후 Logical Hostname(논리 호스트 이름)을 눌러서 마법사를 시작합니다. Oracle Solaris Cluster Manager 로그인 지침은 Oracle Solaris Cluster Manager에 액세스하는 방법 [278] 을 참조하십시오.
클러스터 노드를 유지 보수 상태로 전환	노드를 유지 보수 상태로 전환하는 방법 [227]
노드 이름 바꾸기	노드의 이름을 바꾸는 방법 [225]
클러스터 노드의 유지 보수 상태 해제	노드를 유지 보수 상태에서 해제하는 방법 [229]
클러스터 노드에서 클러스터 소프트웨어 제거	클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법 [231]
SNMP 이벤트 MIB 추가 및 관리	SNMP 이벤트 MIB를 사용하여 설정하는 방법 [235] 노드에서 SNMP 사용자를 추가하는 방법 [239]
각 노드에 대한 로드 한계 구성	노드에 대해 로드 한계를 구성하는 방법 [241]
영역 클러스터 이동, 응용 프로그램에 대해 영역 클러스터 준비, 영역 클러스터 제거	“영역 클러스터 관리 작업 수행” [243]

▼ 클러스터 이름을 변경하는 방법

필요한 경우 설치한 후에 클러스터 이름을 변경할 수 있습니다.



주의 - 클러스터가 Oracle Solaris Cluster Geographic Edition 파트너십에 속한 경우 이 절차를 수행하지 마십시오. 대신 [Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide](#)의 “Renaming a Cluster That Is in a Partnership”의 절차를 따르십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 전역 클러스터의 임의 노드에서 root 역할로 전환합니다.

2. **clsetup** 유틸리티를 시작합니다.

```
phys-schost# clsetup
```

주 메뉴가 표시됩니다.

3. 클러스터 이름을 변경하려면 기타 클러스터 등록 정보 옵션에 대한 번호를 입력합니다.

기타 클러스터 등록 정보 메뉴가 나타납니다.

4. 메뉴에서 원하는 항목을 선택하고 화면의 지시를 따릅니다.

5. Oracle Solaris Cluster의 서비스 태그에 새 클러스터 이름을 반영하게 하려면 기존 Oracle Solaris Cluster 태그를 삭제하고 클러스터를 다시 시작합니다.

Oracle Solaris Cluster 서비스 태그 인스턴스를 삭제하려면 클러스터의 모든 노드에서 다음 하위 단계를 수행합니다.

a. 모든 서비스 태그를 나열합니다.

```
phys-schost# stclient -x
```

b. Oracle Solaris Cluster 서비스 태그 인스턴스 번호를 찾은 후 다음 명령을 수행합니다.

```
phys-schost# stclient -d -i service_tag_instance_number
```

c. 클러스터의 모든 노트를 재부트합니다.

```
phys-schost# reboot
```

예 74 클러스터 이름 변경

다음 예에서는 새 클러스터 이름 dromedary로 변경하기 위해 **clsetup** 유틸리티에서 생성되는 **cluster** 명령을 보여 줍니다.

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

자세한 내용은 [cluster\(1CL\)](#) 및 [clsetup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 노드 ID를 노드 이름에 매핑하는 방법

Oracle Solaris Cluster 설치 중 각 노드에 고유한 노드 ID 번호가 자동으로 지정됩니다. 노드 ID 번호는 처음으로 클러스터에 연결되는 순서대로 노드에 할당됩니다. 노드 ID 번호가 할당되고 나면 해당 번호를 변경할 수 없습니다. 노드 ID 번호는 오류 메시지에서 관련된 클러스터 노드를 나타내는 데 사용됩니다. 노드 ID와 노드 이름 사이의 매핑을 결정하려면 이 절차를 사용합니다.

전역 클러스터 또는 영역 클러스터의 구성 정보를 나열하기 위해 `root` 역할로 전환할 필요는 없습니다. 이 절차의 한 단계는 전역 클러스터의 노드에서 수행되고 다른 단계는 영역 클러스터 노드에서 수행됩니다.

1. **`clnode` 명령을 사용하여 전역 클러스터에 대한 클러스터 구성 정보를 나열합니다.**

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

자세한 내용은 [clnode\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

2. **(옵션) 영역 클러스터에 대한 노드 ID를 나열합니다.**

영역 클러스터 노드의 노드 ID는 실행되고 있는 전역 클러스터 노드와 동일합니다.

```
phys-schost# zlogin sczone clnode -v | grep Node
```

예 75 노드 ID를 노드 이름에 매핑

다음 예에서는 전역 클러스터에 대한 노드 ID 할당을 보여 줍니다.

```
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
```

▼ 새 클러스터 노드에 대한 인증 작업 방법

Oracle Solaris Cluster를 사용하면 새 노드가 전역 클러스터에 추가될 수 있는지 여부와 사용할 인증 유형을 결정할 수 있습니다. 새 노드가 공용 네트워크를 통해 클러스터에 연결되도록 허용하거나 클러스터에 연결되지 않도록 금지할 수도 있고 클러스터에 연결할 수 있는 특정 노드를 지정할 수도 있습니다.

새 노드는 표준 UNIX 또는 DES (Diffie-Hellman) 인증을 사용하여 인증될 수 있습니다. DES 인증을 선택하면 필요한 암호화 키를 모두 구성해야 노드가 연결할 수 있습니다. 자세한 내용은 [keyserv\(1M\)](#) 및 [publickey\(4\)](#) 매뉴얼 페이지를 참조하십시오.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **전역 클러스터의 임의 노드에서 `root` 역할로 전환합니다.**

2. **clsetup** 유틸리티를 시작합니다.

```
phys-schost# clsetup
```

주 메뉴가 표시됩니다.

3. 클러스터 인증 작업을 하려면 새 노드의 옵션에 대한 번호를 입력합니다.

새 노드 메뉴가 표시됩니다.

4. 메뉴에서 원하는 항목을 선택하고 화면의 지시를 따릅니다.

예 76 클러스터에 새 시스템 추가 방지

clsetup 유틸리티는 **claccess** 명령을 생성합니다. 다음 예에서는 새 시스템이 클러스터에 추가되지 않도록 하는 **claccess** 명령을 보여 줍니다.

```
phys-schost# claccess deny -h hostname
```

예 77 전역 클러스터에 모든 새 시스템 추가 허용

clsetup 유틸리티는 **claccess** 명령을 생성합니다. 다음 예에서는 모든 새 시스템을 클러스터에 추가할 수 있도록 하는 **claccess** 명령을 보여 줍니다.

```
phys-schost# claccess allow-all
```

예 78 전역 클러스터에 추가할 새 시스템 지정

clsetup 유틸리티는 **claccess** 명령을 생성합니다. 다음 예에서는 단일 새 시스템을 클러스터에 추가할 수 있도록 하는 **claccess** 명령을 보여 줍니다.

```
phys-schost# claccess allow -h hostname
```

예 79 인증을 표준 UNIX로 설정

clsetup 유틸리티는 **claccess** 명령을 생성합니다. 다음 예에서는 클러스터를 결합하는 새 노드에 대해 표준 UNIX 인증으로 재설정되는 **claccess** 명령을 보여 줍니다.

```
phys-schost# claccess set -p protocol=sys
```

예 80 인증을 DES로 설정

clsetup 유틸리티는 **claccess** 명령을 생성합니다. 다음 예에서는 클러스터를 결합하는 새 노드에 대해 DES 인증을 사용하는 **claccess** 명령을 보여 줍니다.

```
phys-schost# claccess set -p protocol=des
```

DES 인증을 사용할 경우에는 필요한 암호화 키도 모두 구성해야 노드가 클러스터에 연결할 수 있습니다. 자세한 내용은 [keyserv\(1M\)](#) 및 [publickey\(4\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 클러스터에서 시간을 설정하는 방법

Oracle Solaris Cluster 소프트웨어에서는 NTP를 사용하여 클러스터 노드 간에 시간 동기화를 유지합니다. 노드가 시간을 동기화하면 필요에 따라 전역 클러스터에서 조정 작업이 자동으로 수행됩니다. 자세한 내용은 *Oracle Solaris Cluster 4.3 Concepts Guide* 및 <http://download.oracle.com/docs/cd/E19065-01/servers.10k/>의 *Network Time Protocol's User's Guide*를 참조하십시오.



주의 - NTP를 사용할 경우에 클러스터가 실행되고 있을 때는 클러스터를 조정하지 마십시오. `date`, `rdate`, `svcadm` 명령을 대화식으로 사용하여 또는 `cron` 스크립트 내에서 시간을 조정하지 마십시오. 자세한 내용은 `date(1)`, `rdate(1M)`, `svcadm(1M)`, `cron(1M)` 매뉴얼 페이지를 참조하십시오. `ntpd(1M)` 매뉴얼 페이지는 `service/network/ntp` Oracle Solaris 11 패키지로 제공됩니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 전역 클러스터의 임의 노드에서 `root` 역할로 전환합니다.
2. 전역 클러스터를 종료합니다.

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

3. 노드에 `ok` 프롬프트(SPARC 기반 시스템) 또는 `Press any key to continue` 메시지(x86 기반 시스템의 GRUB 메뉴)가 표시되는지 확인합니다.
4. 비클러스터 모드로 노드를 부트합니다.

- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot -x
```

- x86 기반 시스템에서는 다음 명령을 실행합니다.

```
# shutdown -g -y -i0
```

```
Press any key to continue
```

- a. GRUB 메뉴에서 화살표 키를 사용하여 적절한 Oracle Solaris 항목을 선택하고 `e`를 입력하여 해당 명령을 편집합니다.

GRUB 메뉴가 나타납니다.

GRUB 기반 부트에 대한 자세한 내용은 *Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”*를 참조하십시오.

- b. 부트 매개 변수 화면에서 화살표 키를 사용하여 커널 항목을 선택하고 e를 입력하여 항목을 편집합니다.

GRUB 부트 매개 변수 화면이 나타납니다.

- c. 명령에 -x를 추가하여 시스템 부트를 비클러스터 모드로 지정합니다.

[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. Enter 키를 눌러 변경사항을 수락하고 부트 매개 변수 화면으로 복귀합니다.

화면에 편집된 명령이 표시됩니다.

- e. b를 입력하여 비클러스터 모드로 노드를 부트합니다.

주 - 커널 부트 매개 변수 명령에 대한 변경사항은 시스템을 재부트하면 사라집니다. 다음에 노드를 재부트하면 클러스터 모드로 부트됩니다. 대신 비클러스터 모드로 부트하려면, 이러한 단계를 다시 실행하여 -x 옵션을 커널 부트 매개 변수 명령에 추가합니다.

5. 단일 노드에서 **date** 명령을 실행하여 시간을 설정합니다.

```
phys-schost# date HHMM.SS
```

6. 다른 시스템에서 **rdate(1M)** 명령을 실행하여 시간을 위의 노드와 동기화합니다.

```
phys-schost# rdate hostname
```

7. 각 노드를 부트하여 클러스터를 다시 시작합니다.

```
phys-schost# reboot
```

8. 모든 클러스터 노드에서 변경되었는지 확인합니다.

각 노드에서 date 명령을 실행합니다.

```
phys-schost# date
```

▼ SPARC: 노드에서 OBP(OpenBoot PROM)를 표시하는 방법

OpenBoot™ PROM 설정을 구성하거나 변경해야 하는 경우 이 절차를 사용합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 종료할 노드의 콘솔에 연결합니다.

```
# telnet tc_name tc_port_number
```

tc_name 터미널 콘센트레이터의 이름을 지정합니다.

tc_port_number 터미널 콘센트레이터에 포트 번호를 지정합니다. 포트 번호는 구성에 따라 다릅니다. 일반적으로 포트 2와 3(5002 및 5003)은 사이트에 설치된 첫번째 클러스터에 사용됩니다.

2. **clnode evacuate** 명령을 사용한 후 **shutdown** 명령을 사용하여 클러스터 노드를 정상적으로 종료합니다.

clnode evacuate 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 장치 그룹을 전환합니다. 또한 이 명령은 전역 클러스터의 지정된 노드에서 다음 기본 노드로 모든 리소스 그룹을 전환합니다.

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



주의 - 클러스터 노드를 종료하기 위해 클러스터 콘솔에서 **send brk**를 사용하지 마십시오.

3. OBP 명령을 실행합니다.

▼ 노드 개인 호스트 이름을 변경하는 방법

설치가 완료된 후 이 절차를 사용하여 클러스터 노드의 개인 호스트 이름을 변경합니다.

처음 클러스터를 설치할 때 개인 호스트 이름으로 기본값이 할당됩니다. 기본 개인 호스트 이름은 **clusternode***nodeid*-priv 형식(예: **clusternode3-priv**)을 사용합니다. 해당 이름을 도메인에서 이미 사용 중인 경우에만 개인 호스트 이름을 변경합니다.



주의 - 새 개인 호스트 이름에 IP 주소를 할당하지 마십시오. IP 주소는 클러스터링 소프트웨어에서 지정합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문 형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터의 모든 노드에서 개인 호스트 이름을 캐시할 수 있는 데이터 서비스 리소스나 기타 응용 프로그램을 사용 안함으로 설정합니다.

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

사용 안함으로 설정할 응용 프로그램에 다음을 포함하십시오.

- HA-DNS 및 HA-NFS 서비스(구성된 경우)
- 개인 호스트 이름을 사용하도록 사용자정의 구성된 모든 응용 프로그램
- 개인용 상호 연결을 통해 클라이언트가 사용하는 모든 응용 프로그램

clresource 명령 사용에 대한 자세한 내용은 [clresource\(1CL\)](#) 매뉴얼 페이지 및 [Oracle Solaris Cluster 4.3 데이터 서비스 계획 및 관리 설명서](#)를 참조하십시오.

2. 사용 중인 NTP 구성 파일이 변경할 개인 호스트 이름을 참조하는 경우 클러스터의 각 노드에서 NTP 데몬을 중지합니다.

svcadm 명령을 사용하여 NTP 데몬을 종료합니다. NTP 데몬에 대한 자세한 내용은 [svcadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

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

3. clsetup 유틸리티를 실행하여 해당 노드의 개인 호스트 이름을 변경합니다.

클러스터의 노드 중 하나에서만 유틸리티를 실행합니다. 자세한 내용은 [clsetup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - 새 개인 호스트 이름을 선택할 경우에는 이름이 클러스터 노드에서 고유해야 합니다.

clsetup 유틸리티 대신 clnode 명령을 실행하여 개인 호스트 이름을 변경할 수도 있습니다. 아래 예에서 클러스터 노드 이름은 phys-schost-1입니다. 다음 clnode 명령을 실행한 후 [6단계](#)로 이동합니다.

```
phys-schost# clnode set -p privatehostname=New-private-nodename phys-schost-1
```

4. clsetup 유틸리티에서 개인 호스트 이름 옵션에 대한 번호를 입력합니다.
5. clsetup 유틸리티에서 개인 호스트 이름을 변경하는 옵션에 대한 번호를 입력합니다.

화면에 표시되는 질문에 답하십시오. 개인 호스트 이름을 변경할 노드의 이름 (clusternodenodeid-priv)과 새 개인 호스트 이름을 묻는 메시지가 표시됩니다.

6. 이름 서비스 캐시를 비웁니다.

클러스터의 각 노드에서 이 단계를 수행합니다. 해당 캐시를 비우면 클러스터 응용 프로그램 및 데이터 서비스에서 이전의 개인 호스트 이름에 액세스할 수 없습니다.

```
phys-schost# nscd -i hosts
```

7. NTP 구성이나 include 파일에서 개인 호스트 이름을 변경한 경우 각 노드에서 NTP 파일을 업데이트합니다.

NTP 구성 파일(/etc/inet/ntp.conf)에서 개인 호스트 이름을 변경했고 피어 호스트 항목 또는 피어 호스트의 include 파일을 가리키는 포인터가 NTP 구성 파일(/etc/inet/ntp.conf.include)에 있는 경우 각 노드에서 파일을 업데이트합니다. NTP include 파일에서 개인 호스트 이름을 변경한 경우 각 노드에서 /etc/inet/ntp.conf.sc 파일을 업데이트합니다.

- a. 원하는 편집 도구를 사용합니다.

설치 시 이 단계를 수행할 경우에는 구성된 노드의 이름도 제거해야 합니다. 일반적으로 각 클러스터 노드에 있는 ntp.conf.sc 파일은 동일합니다.

- b. 모든 클러스터 노드에서 새 개인 호스트 이름에 대해 ping 명령을 수행하여 성공하는지 확인합니다.

- c. NTP 데몬을 다시 시작합니다.

클러스터의 각 노드에서 이 단계를 수행하십시오.

svcadm 명령을 사용하여 NTP 데몬을 다시 시작합니다.

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

8. 1단계에서 사용 안함으로 설정된 모든 데이터 서비스 리소스와 다른 응용 프로그램을 사용으로 설정합니다.

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

clresource 명령 사용에 대한 자세한 내용은 [clresource\(1CL\)](#) 매뉴얼 페이지 및 [Oracle Solaris Cluster 4.3 데이터 서비스 계획 및 관리 설명서](#)를 참조하십시오.

예 81 개인 호스트 이름 변경

다음 예에서는 phys-schost-2 노드의 개인 호스트 이름을 clusternode2-priv에서 clusternode4-priv로 변경합니다. 각 노드에서 이 작업을 수행합니다.

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
```

▼ 노드의 이름을 바꾸는 방법

Oracle Solaris Cluster 구성의 일부인 노드의 이름을 변경할 수 있습니다. 노드의 이름을 바꾸기 전에 Oracle Solaris 호스트 이름을 바꾸어야 합니다. `clnode rename` 명령을 사용하여 노드의 이름을 바꿉니다.

다음 지침은 전역 클러스터에서 실행되는 모든 응용 프로그램에 적용됩니다.

1. 전역 클러스터에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. (옵션) 파트너십에 속하는 Oracle Solaris Cluster Geographic Edition 클러스터에 있는 노드의 이름을 바꾸려면 보호 그룹을 스위치오버할지 여부를 결정합니다.

이름 바꾸기 절차를 수행 중인 클러스터가 보호 그룹의 기본 클러스터이고 보호 그룹의 응용 프로그램을 온라인으로 유지하려는 경우 이름 바꾸기 절차 중 보호 그룹을 보조 클러스터로 전환할 수 있습니다.

Geographic Edition 클러스터 및 노드에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 Geographic Edition System Administration Guide](#) 의 5 장, “Administering Cluster Partnerships”를 참조하십시오.

3. Oracle Solaris 호스트 이름을 바꿉니다.
[Oracle Solaris 11.3의 시스템 정보, 프로세스, 성능 관리](#) 의 “시스템의 ID를 변경하는 방법”에 나와 있는 단계를 완료합니다. 단, 해당 절차의 끝에 있는 재부트는 수행하지 마십시오.

대신 이 단계를 완료한 후 클러스터 종료를 수행합니다.

4. 모든 클러스터 노드를 비클러스터 모드로 부트합니다.

```
ok> boot -x
```

5. Oracle Solaris 호스트 이름을 바꾼 노드에서 비클러스터 모드로 노드의 이름을 바꾸고 이름을 바꾼 각 호스트에서 `cmd` 명령을 실행합니다.

한번에 하나씩 노드 이름을 바꿉니다.

```
# clnode rename -n new-node old-node
```

6. 클러스터에서 실행되는 응용 프로그램에 이전 호스트 이름을 참조하는 기존 항목이 있으면 업데이트합니다.
7. 명령 메시지 및 로그 파일을 확인하여 노드 이름이 바뀌었는지 확인합니다.
8. 모든 노드를 클러스터 모드로 재부트합니다.

```
# sync;sync;sync;reboot
```
9. 노드에 새 이름이 표시되는지 확인합니다.

```
# clnode status -v
```
10. 새 클러스터 노드 이름을 사용하도록 Geographic Edition 구성을 업데이트합니다.
보호 그룹 및 데이터 복제 제품에서 사용되는 구성 정보에 노드 이름이 지정될 수도 있습니다.
11. 논리 호스트 이름 리소스의 `hostnameList` 등록 정보를 변경하도록 선택할 수 있습니다.
이 선택적 단계에 대한 지침은 [기존 Oracle Solaris Cluster 논리 호스트 이름 리소스에서 사용하는 논리 호스트 이름을 변경하는 방법 \[226\]](#)을 참조하십시오.

▼ 기존 Oracle Solaris Cluster 논리 호스트 이름 리소스에서 사용하는 논리 호스트 이름을 변경하는 방법

[노드의 이름을 바꾸는 방법 \[225\]](#)에 나와 있는 단계에 따라 노드의 이름을 바꾸기 전후에 논리 호스트 이름 리소스의 `HostnameList` 등록 정보를 변경하도록 선택할 수 있습니다. 이 단계는 선택사항입니다.

1. 전역 클러스터에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 선택적으로 기존 Oracle Solaris Cluster 논리 호스트 이름 리소스에서 사용하는 논리 호스트 이름을 변경할 수 있습니다.
다음 단계에서는 새 논리 호스트 이름을 사용하도록 `apache-lh-res` 리소스를 구성하는 방법을 보여 주며 이 방법은 클러스터 모드로 실행되어야 합니다.
 - a. 클러스터 모드에서 논리 호스트 이름을 포함하는 Apache 리소스 그룹을 오프라인 상태로 전환합니다.

```
# clresourcegroup offline apache-rg
```
 - b. Apache 논리 호스트 이름 리소스를 사용 안함으로 설정합니다.

```
# clresource disable apache-lh-res
```

- c. 새 호스트 이름 목록을 제공합니다.

```
# clresource set -p HostnameList=test-2 apache-lh-res
```

- d. `hostnamelist` 등록 정보에서 이전 항목에 대한 응용 프로그램 참조를 새 항목으로 변경합니다.

- e. 새 Apache 논리 호스트 이름 리소스를 사용으로 설정합니다.

```
# clresource enable apache-lh-res
```

- f. Apache 리소스 그룹을 온라인 상태로 전환합니다.

```
# clresourcegroup online -eM apache-rg
```

- g. 클라이언트를 확인하는 다음 명령을 실행하여 응용 프로그램이 제대로 시작되었는지 확인합니다.

```
# clresource status apache-rs
```

▼ 노드를 유지 보수 상태로 전환하는 방법

오랫동안 노드의 서비스를 중단하는 경우 전역 클러스터 노드를 유지 보수 상태로 전환합니다. 이 방법을 사용하면 노드가 서비스를 받고 있지만 쿼럼 수에는 포함되지 않습니다. 노드를 유지 보수 상태로 전환하려면 `clnode evacuate` 및 `shutdown` 명령을 사용하여 해당 노드를 종료해야 합니다. 자세한 내용은 [clnode\(1CL\)](#) 및 [cluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 노드를 비우고 모든 리소스 그룹 및 장치 그룹을 다음 선호 노드로 전환할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

클러스터 노드가 중지되어 유지 보수 상태가 되면 노드에 대한 포트를 사용하여 구성된 모든 쿼럼 장치의 투표 수가 하나씩 감소됩니다. 노드를 유지 보수 모드에서 제거하여 다시 온라인 상태로 전환하면 노드와 쿼럼 장치 투표 수가 하나씩 증가됩니다.

주 - Oracle Solaris `shutdown` 명령은 단일 노드를 종료하지만 `cluster shutdown` 명령은 전체 클러스터를 종료합니다.

아직 클러스터 멤버인 다른 노드에서 `clquorum disable` 명령을 사용하여 클러스터 노드를 유지 보수 상태로 전환할 수 있습니다. 자세한 내용은 [clquorum\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 유지 보수 상태로 전환할 전역 클러스터 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 노드에서 모든 리소스 그룹 또는 장치 그룹을 제거합니다.

`clnode evacuate` 명령은 지정된 노드에서 다음 우선 순위 노드로 모든 리소스 그룹 및 장치 그룹을 전환합니다.

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

3. 제거한 노드를 종료합니다.

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

4. 클러스터의 다른 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환하고 3단계에서 종료한 노드를 유지 보수 상태로 전환합니다.

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

5. 전역 클러스터 노드가 현재 유지 보수 상태에 있는지 확인합니다.

```
phys-schost# clquorum status node
```

유지 보수 상태로 전환한 노드의 Status는 Present 및 Possible 쿼럼 투표에 대해 offline 및 0이어야 합니다.

예 82 전역 클러스터 노드를 유지 보수 상태로 전환

다음 예에서는 클러스터 노드를 유지 보수 상태로 전환한 후에 결과를 확인합니다. `clnode status` 명령을 실행하면 phys-schost-1에 대한 Node votes는 0으로 출력되고 상태는 Offline으로 출력됩니다. Quorum Summary에 줄어든 투표 수도 표시되어야 합니다. 구성에 따라 Quorum Votes by Device 출력에 일부 쿼럼 디스크 장치가 오프라인 상태인 것도 표시될 수 있습니다.

[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

참조 노드를 다시 온라인 상태로 전환하려면 [노드를 유지 보수 상태에서 해제하는 방법 \[229\]](#)을 참조하십시오.

▼ 노드를 유지 보수 상태에서 해제하는 방법

전역 클러스터 노드를 다시 온라인 상태로 전환하고 쿼럼 투표 수를 기본값으로 재설정하려면 다음 절차를 사용합니다. 클러스터 노드의 경우에 기본 쿼럼 수는 하나입니다. 쿼럼 장치의 경우 기본 쿼럼 수는 $N-1$ 이며, 여기서 N 은 쿼럼 장치에 대한 포트가 있으면서 투표 수가 0이 아닌 노드의 수입니다.

노드가 유지 보수 상태로 전환되었으면 노드의 쿼럼 투표 수가 하나씩 감소됩니다. 또한 노드에 대한 포트를 사용하여 쿼럼 장치가 구성되면 쿼럼 투표 수가 하나씩 감소합니다. 쿼럼 투표 수가 재설정되고 노드가 유지 보수 상태에서 해제되면, 노드의 쿼럼 투표 수 및 쿼럼 장치 투표 수가 하나씩 증가합니다.

유지 보수 상태에 있던 전역 클러스터 노드를 유지 보수 상태에서 해제하려면 이 절차를 수행합니다.



주의 - globaldev 또는 node 옵션을 지정하지 않으면 쿼럼 계수가 전체 클러스터에 대해 재설정됩니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 유지 보수 상태에 있는 노드 이외의 전역 클러스터 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 전역 클러스터 구성에 있는 노드의 수에 따라 다음 단계 중 하나를 수행합니다.
 - 클러스터 구성의 노드가 2개인 경우 [4단계](#)로 이동합니다.
 - 클러스터 구성의 노드가 3개 이상인 경우 [3단계](#)로 이동합니다.
3. 유지 보수 상태를 제거하려는 노드에 쿼럼 장치가 있을 경우 유지 보수 상태에 있지 않은 노드에서 클러스터 쿼럼 수를 재설정합니다.

노드를 재부트하기 전에 유지 보수 상태의 노드가 아닌 다른 노드에서 쿼럼 수를 재설정해야 합니다. 재설정하지 않으면 해당 노드가 쿼럼 대기 중에 멈출 수도 있습니다.

```
phys-schost# clquorum reset
```

4. 유지 보수 상태에서 해제할 노드를 부트합니다.

5. 쿼럼 투표 수를 확인하십시오.

```
phys-schost# clquorum status
```

유지 보수 상태에서 해제된 노드는 online 상태이고 Present 및 Possible 쿼럼 투표에 대해 필요한 투표 수가 표시되어야 합니다.

예 83 클러스터 노드의 유지 보수 상태 해제 및 쿼럼 투표 수 재설정

다음 예에서는 클러스터 노드 및 해당 쿼럼 장치에 대한 쿼럼 수를 다시 기본값으로 재설정하고 결과를 확인합니다. cluster status 명령을 실행하면 phys-schost-1에 대한 Node votes는 1로 출력되고 상태는 online으로 출력됩니다. Quorum Summary에 투표 수 증가도 표시되어야 합니다.

```
phys-schost-2# clquorum reset
```

- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot
```

- x86 기반 시스템에서는 다음 명령을 실행합니다.

GRUB 메뉴가 표시되면 적절한 Oracle Solaris 항목을 선택하고 Enter 키를 누릅니다.

```
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

▼ 클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법

완전히 설정된 클러스터 구성에서 연결을 끊기 전에 전역 클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 구성 해제하려면 이 절차를 수행합니다. 이 절차를 사용하여 클러스터에 남은 마지막 노드에서 소프트웨어를 제거할 수 있습니다.

주 - 아직 클러스터에 연결되지 않았거나 설치 모드 상태인 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 경우에는 이 절차를 수행하지 마십시오. 대신 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 소프트웨어를 구성 해제하여 설치 문제를 해결하는 방법”으로 이동합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 클러스터 노드를 제거하려면 작업 맵에 있는 모든 사전 작업을 정확하게 완료해야 합니다.

[표 16. “작업 맵: 노드 제거”](#)를 참조하십시오.

이 절차를 계속 진행하기 전에 `clnode remove`를 사용하여 클러스터 구성에서 노드를 제거했는지 확인합니다. 제거할 노드를 클러스터의 노드 인증 목록에 추가하고 영역 클러스터를 제거하는 등의 다른 단계가 포함될 수 있습니다.

주 - 노드 구성을 해제하되 Oracle Solaris Cluster 소프트웨어는 노드에 설치된 채 두려면 `clnode remove` 명령을 실행한 후에 더 이상 진행하지 마십시오.

2. 설치 제거할 노드에서 `root` 역할로 전환합니다.
3. 노드에 전역 장치 이름 공간 전용 파티션이 있는 경우 전역 클러스터 노드를 비클러스터 모드로 재부팅합니다.

- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

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

- x86 기반 시스템에서는 다음 명령을 실행합니다.

```
# 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. `/etc/vfstab` 파일에서 `/global/devices` 전역 마운트를 제외하고 전역으로 마운트된 파일 시스템 항목을 모두 제거합니다.

5. 비클러스터 모드로 노드를 재부트합니다.

- SPARC 기반 시스템에서는 다음 명령을 실행합니다.

```
ok boot -x
```

- x86 기반 시스템에서는 다음 명령을 실행합니다.

- a. GRUB 메뉴에서 화살표 키를 사용하여 적절한 Oracle Solaris 항목을 선택하고 `e`를 입력하여 해당 명령을 편집합니다.

GRUB 기반 부트에 대한 자세한 내용은 [Oracle Solaris 11.3 시스템 부트 및 종료의 “시스템 부트”](#)를 참조하십시오.

- b. 부트 매개변수 화면에서 화살표 키를 사용하여 `kernel` 항목을 선택하고 `e`를 입력하여 항목을 편집합니다.

- c. 명령에 `-x`를 추가하여 시스템 부트를 비클러스터 모드로 지정합니다.

- d. Enter 키를 눌러 변경사항을 적용하고 부트 매개 변수 화면으로 돌아갑니다.
화면에 편집된 명령이 표시됩니다.

- e. `b`를 입력하여 비클러스터 모드로 노드를 부트합니다.

주 - 커널 부트 매개 변수 명령에 대한 변경사항은 시스템을 재부트하면 사라집니다. 다음에 노드를 재부트하면 클러스터 모드로 부트됩니다. 비클러스터 모드로 부트하려면 이 단계를 다시 실행하여 `-x` 옵션을 커널 부트 매개 변수 명령에 추가합니다.

6. Oracle Solaris Cluster 패키지에서 제공하는 파일이 없는 디렉토리(예: 루트(/) 디렉토리)로 변경합니다.

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

7. 노드 구성을 해제하고 Oracle Solaris Cluster 소프트웨어를 제거하려면 다음 명령을 실행합니다.

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

-r

클러스터 구성 정보를 제거하고 클러스터 노드에서 Oracle Solaris Cluster 프레임워크 및 데이터 서비스 소프트웨어를 제거합니다. 그런 다음 노드를 다시 설치하거나 클러스터에서 노드를 제거할 수 있습니다.

-b *bootenvironmentname*

새 부트 환경(제거 프로세스를 완료한 후에 부트할 위치)의 이름을 지정합니다. 이름 지정은 선택사항입니다. 부트 환경의 이름을 지정하지 않으면 자동으로 생성됩니다.

자세한 내용은 [scinstall\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

8. 제거를 완료한 후에 이 노드에 Oracle Solaris Cluster 소프트웨어를 다시 설치하려면 새 부트 환경으로 부트하도록 노드를 재부팅합니다.
9. 이 클러스터에 Oracle Solaris Cluster 소프트웨어를 다시 설치하지 않으려는 경우 다른 클러스터 장치와 연결된 전송 케이블 및 전송 스위치가 있으면 연결 해제합니다.
- a. 제거된 노드가 병렬 SCSI 인터페이스를 사용하는 저장 장치에 연결되어 있으면 전송 케이블 연결을 제거한 후에 저장 장치의 열린 SCSI 커넥터에 SCSI 터미네이터를 설치하십시오.
제거되는 노드가 광섬유 채널 인터페이스를 사용하는 저장 장치에 연결되어 있으면 터미네이터 장치가 없어도 됩니다.
 - b. 연결 해제 절차는 호스트 어댑터 및 서버와 함께 제공된 설명서를 따릅니다.

팁 - 전역 장치 이름 공간을 *lofi* 장치로 마이그레이션에 대한 자세한 내용은 “[전역 장치 이름 공간 마이그레이션](#)” [117]을 참조하십시오.

노드 제거 문제 해결

이 절에서는 `clnode remove` 명령을 실행할 때 표시될 수 있는 오류 메시지와 해결 방법에 대해 설명합니다.

제거되지 않은 클러스터 파일 시스템 항목

다음 오류 메시지가 표시되면 제거한 전역 클러스터 노드의 `vfstab` 파일에 클러스터 파일 시스템 참조 항목이 아직 남아 있는 것입니다.

```
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.
```

이 오류를 해결하려면 [클러스터 노드에서 Oracle Solaris Cluster 소프트웨어를 제거하는 방법 \[231\]](#)으로 돌아가 해당 절차를 반복합니다. clnode remove 명령을 다시 실행하기 전에 이 절차의 4단계를 성공적으로 완료해야 합니다.

장치 그룹의 목록에서 제거되지 않은 항목

다음 오류 메시지가 표시되면 제거한 노드가 장치 그룹 목록에 아직 남아 있는 것입니다.

```
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.
```

Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관리

이 절에서는 SNMP(Simple Network Management Protocol) 이벤트 MIB (Management Information Base)를 만들기, 설정 및 관리하는 방법을 설명합니다. 또한 Oracle Solaris Cluster SNMP 이벤트 MIB를 사용 및 사용 안함으로 설정하고 변경하는 방법에 대해서도 설명합니다.

Oracle Solaris Cluster 소프트웨어는 현재 MIB 중 이벤트 MIB만 지원합니다. SNMP 관리자 소프트웨어는 실시간으로 클러스터 이벤트를 트랩합니다. 사용 가능한 경우 SNMP 관리자는 clsnmphost 명령을 통해 정의된 모든 호스트에 트랩 알림을 자동으로 전송합니다. 클러스터가 많은 수의 알림을 생성하므로 min_severity 이상의 심각도를 포함하는 이벤트만 트랩 알림으로 전송됩니다. 기본적으로 min_severity 값은 NOTICE로 설정되어 있습니다. log_number 값은 오래된 항목을 폐기하기 전에 MIB 테이블에 로깅할 이벤트 수를 지정합니다. MIB는 트랩을 보낸 최근 이벤트의 읽기 전용 테이블을 유지합니다. 이벤트 수는 log_number 값으로 제한됩니다. 재부트 시 이 정보는 지속되지 않습니다.

SNMP 이벤트 MIB는 sun-cluster-event-mib.mib 파일에 정의되어 있으며 /usr/cluster/lib/mib 디렉토리에 위치합니다. 이 정의는 SNMP 트랩 정보를 해석하는 데 사용할 수 있습니다.

클러스터 이벤트 SNMP 인터페이스는 공통 에이전트 컨테이너(cacao) SNMP 어댑터를 해당 SNMP 에이전트 기반구조로 사용합니다. 기본적으로 SNMP의 포트 번호는 11161이고 SNMP 트랩의 기본 포트 번호는 11162입니다. 이 포트 번호는 `cacaoadm` 명령을 사용하여 변경할 수 있습니다. 자세한 내용은 `cacaoadm(1M)` 매뉴얼 페이지를 참조하십시오.

Oracle Solaris Cluster SNMP 이벤트 MIB의 만들기, 설정 및 관리에는 다음 작업이 수반될 수 있습니다.

표 18 작업 맵: Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관리

작업	지침
SNMP 이벤트 MIB 사용	SNMP 이벤트 MIB를 사용으로 설정하는 방법 [235]
SNMP 이벤트 MIB 사용 안함	SNMP 이벤트 MIB를 사용 안함으로 설정하는 방법 [236]
SNMP 이벤트 MIB 변경	SNMP 이벤트 MIB를 변경하는 방법 [236]
SNMP 호스트를 MIB에 대한 트랩 통지를 수신할 호스트 목록에 추가	SNMP 호스트가 노드에서 SNMP 트랩을 수신하도록 설정하는 방법 [237]
SNMP 호스트 제거	SNMP 호스트가 노드에서 SNMP 트랩을 수신하지 않도록 설정하는 방법 [238]
SNMP 사용자 추가	노드에서 SNMP 사용자를 추가하는 방법 [239]
SNMP 사용자 제거	노드에서 SNMP 사용자를 제거하는 방법 [240]

▼ SNMP 이벤트 MIB를 사용으로 설정하는 방법

이 절차에서는 SNMP 이벤트 MIB를 사용으로 설정하는 방법을 보여 줍니다.

`phys-schost#` 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. SNMP 이벤트 MIB를 사용으로 설정합니다.

```
phys-schost-1# clsnmpmib enable [-n node] MIB
```

`[-n node]` 사용으로 설정하려는 이벤트 MIB가 있는 *node*를 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우, 현재 노드가 기본값으로 사용됩니다.

MIB 사용으로 설정하려는 MIB의 이름을 지정합니다. 이 경우, MIB 이름은 `event`여야 합니다.

▼ SNMP 이벤트 MIB를 사용 안함으로 설정하는 방법

이 절차에서는 SNMP 이벤트 MIB를 사용 안함으로 설정하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. **SNMP 이벤트 MIB를 사용 안함으로 설정합니다.**

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

-n node 사용 안함으로 설정하려는 이벤트 MIB가 있는 *node*를 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우, 현재 노드가 기본값으로 사용됩니다.

MIB 사용 안함으로 설정하려는 MIB의 유형을 지정합니다. 이 경우, *event*를 지정해야 합니다.

▼ SNMP 이벤트 MIB를 변경하는 방법

이 절차에서는 SNMP 이벤트 MIB의 프로토콜, 최소 심각도 값 및 이벤트 로깅을 변경하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. **SNMP 이벤트 MIB의 프로토콜, 최소 심각도 값 및 이벤트 로깅을 변경합니다.**

```
phys-schost-1# clsnmpmib set -n node
```

```
-p version=SNMPv3 \
-p min_severity=WARNING \
-p log_number=100 MIB
```

-n node

변경하려는 이벤트 MIB가 있는 *node*를 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우, 현재 노드가 기본값으로 사용됩니다.

-p version=*value*

MIB와 함께 사용할 SNMP 프로토콜의 버전을 지정합니다. 다음과 같이 *value*를 지정합니다.

- version=SNMPv2
- version=snmpv2
- version=2
- version=SNMPv3
- version=snmpv3
- version=3

-p min_severity=*value*

MIB와 함께 사용할 최소 심각도 값을 지정합니다. 다음과 같이 *value*를 지정합니다.

- min_severity=NOTICE
- min_severity=WARNING
- min_severity=ERROR
- min_severity=CRITICAL
- min_severity=FATAL

-p log_number=*number*

오래된 항목을 폐기하기 전에 MIB 테이블에 로깅할 이벤트 수를 지정합니다. 기본값은 100입니다. 값은 100-500 범위여야 합니다. 다음과 같이 *value*를 지정합니다.
log_number=100

MIB

MIB 또는 하위 명령을 적용할 MIB의 이름을 지정합니다. 이 경우, event를 지정해야 합니다. 이 피연산자를 지정하지 않을 경우 하위 명령에서 모든 MIB를 의미하는 기본값 더하기 기호(+)를 사용합니다. *MIB* 피연산자를 사용하는 경우 다른 모든 명령줄 옵션 뒤의 공백으로 구분된 목록에 MIB를 지정하십시오.

자세한 내용은 [clsnmpmib\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

▼ SNMP 호스트가 노드에서 SNMP 트랩을 수신하도록 설정하는 방법

이 절차에서는 MIB에 대한 트랩 알림을 수신할 호스트 목록에 노드의 SNMP 호스트를 추가하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 다른 노드에 있는 커뮤니티의 SNMP 호스트 목록에 호스트를 추가합니다.

```
phys-schost-1# clsnmphost add -c SNMPcommunity [-n node] host
```

-c SNMPcommunity

호스트 이름과 함께 사용되는 SNMP 커뮤니티 이름을 지정합니다. 호스트는 트랩을 수신하도록 구성할 수 있는 네트워크의 시스템입니다.

public 이외의 커뮤니티에 호스트를 추가할 때에는 SNMP 커뮤니티 이름 (SNMPcommunity)을 지정해야 합니다. -c 옵션 없이 add 하위 명령을 사용할 경우, 하위 명령은 기본 커뮤니티 이름으로 public을 사용합니다.

지정한 커뮤니티 이름이 존재하지 않을 경우, 이 명령은 커뮤니티를 생성합니다.

-n node

클러스터의 SNMP MIB에 대한 액세스 권한이 있는 SNMP 호스트의 클러스터 *node* 이름을 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우 기본값은 명령이 실행된 노드입니다.

host

클러스터의 SNMP MIB에 대한 액세스 권한이 있는 호스트의 이름, IP 주소 또는 IPv6 주소를 지정합니다. SNMP 트랩을 가져오려고 시도 중인 클러스터 또는 클러스터 노드 자체의 외부 호스트일 수 있습니다.

▼ SNMP 호스트가 노드에서 SNMP 트랩을 수신하지 않도록 설정하는 방법

이 절차에서는 MIB에 대한 트랩 알림을 수신할 호스트 목록에서 노드의 SNMP 호스트를 제거하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 지정한 노드에 있는 커뮤니티의 SNMP 호스트 목록에서 호스트를 제거합니다.

```
phys-schost-1# clsnmphost remove -c SNMPcommunity -n node host
```

remove

지정한 노드에서 지정한 SNMP 호스트를 제거합니다.

-c *SNMPcommunity*

SNMP 호스트가 제거된 SNMP 커뮤니티의 이름을 지정합니다.

-n *node*

구성에서 SNMP 호스트가 제거되는 클러스터 *node* 이름을 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우 기본값은 명령이 실행된 노드입니다.

host

구성에서 제거된 호스트의 이름, IP 주소 또는 IPv6 주소를 지정합니다. SNMP 트랩을 가져오려고 시도 중인 클러스터 또는 클러스터 노드 자체의 외부 호스트일 수 있습니다.

지정한 SNMP 커뮤니티에서 모든 호스트를 제거하려면 *host*에 -c 옵션과 함께 덧셈 부호(+)를 사용합니다. 모든 호스트를 제거하려면 *host*에 덧셈 부호(+)를 사용합니다.

▼ 노드에서 SNMP 사용자를 추가하는 방법

이 절차에서는 노드의 SNMP 사용자 구성에 SNMP 사용자를 추가하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. ***solaris.cluster.modify*** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. **SNMP 사용자를 추가합니다.**

```
phys-schost-1# clsnmpuser create -n node -a authentication -f password user
```

-n *node*

SNMP 사용자를 추가할 노드를 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우, 현재 노드가 기본값으로 사용됩니다.

-a *authentication*

사용자에게 권한을 부여하는데 사용되는 인증 프로토콜을 지정합니다. 인증 프로토콜 값은 SHA 또는 MD5가 될 수 있습니다.

-f *password*

SNMP 사용자 암호가 포함된 파일을 지정합니다. 새 사용자를 만들 때 이 옵션을 지정하지 않을 경우, 해당 명령이 암호를 묻는 메시지를 표시합니다. 이 옵션은 add 하위 명령에서만 유효합니다.

다음 형식과 같이 사용자 암호를 별도의 행에 지정해야 합니다.

user:password

암호에는 다음 문자 또는 공백이 포함될 수 없습니다.

- ; (세미콜론)
- : (콜론)
- \ (백슬래시)
- \n (새 줄)

user

추가하려는 SNMP 사용자의 이름을 지정합니다.

▼ 노드에서 SNMP 사용자를 제거하는 방법

이 절차에서는 노드의 SNMP 사용자 구성에서 SNMP 사용자를 제거하는 방법을 보여 줍니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. **SNMP 사용자를 제거합니다.**

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

-n node SNMP 사용자를 제거할 노드를 지정합니다. 노드 이름 또는 노드 ID를 지정할 수 있습니다. 이 옵션을 지정하지 않을 경우, 현재 노드가 기본값으로 사용됩니다.

user 제거하려는 SNMP 사용자의 이름을 지정합니다.

로드 한계 구성

로드 한계를 설정하여 노드 간에 리소스 그룹 로드 자동 배포를 사용으로 설정할 수 있습니다. 각 클러스터 노드에 대해 일련의 로드 한계를 구성할 수 있습니다. 리소스 그룹에 로드 요소를 할당하며 로드 요소는 노드의 정의된 로드 한계에 해당합니다. 기본 동작은 리소스 그룹의 노드 목록에 있는 사용 가능한 모든 노드에서 리소스 그룹 로드를 균등하게 배포하는 것입니다.

리소스 그룹은 RGM에 의해 리소스 그룹의 노드 목록에 있는 노드에서 시작되어 노드의 로드 한계가 초과되지 않도록 합니다. RGM에서 노드에 리소스 그룹을 할당하면 각 노드에서

리소스 그룹의 로드 요소가 합계되어 총 로드를 제공합니다. 그런 다음 총 로드는 노드의 로드 한계와 비교됩니다.

로드 한계는 다음 항목으로 구성됩니다.

- 사용자가 지정한 이름
- 소프트 한계 값 - 소프트 로드 한계를 일시적으로 초과할 수 있습니다.
- 하드 한계 값 - 하드 로드 한계는 초과할 수 없으며 엄격하게 적용됩니다.

명령 하나만으로 하드 한계와 소프트 한계를 모두 설정할 수 있습니다. 한계 중 하나가 명시적으로 설정되지 않으면 기본값이 사용됩니다. `clnode create-loadlimit`, `clnode set-loadlimit` 및 `clnode delete-loadlimit` 명령을 사용하여 각 노드에 대한 하드 및 소프트 로드 한계를 만들고 수정합니다. 자세한 내용은 [clnode\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

특정 노드에서 재배치되지 않도록 리소스 그룹의 우선 순위를 더 높게 구성할 수 있습니다. 또한 `preemption_mode` 등록 정보를 설정하여 노드 과부화로 인해 리소스 그룹이 우선 순위가 더 높은 리소스 그룹에 의해 노드에서 선점되는지 여부를 결정할 수 있습니다. `concentrate_load` 등록 정보를 사용하면 리소스 그룹 로드를 최대한 적은 수의 노드로 집중할 수도 있습니다. `concentrate_load` 등록 정보의 기본값은 FALSE입니다.

주 - 전역 클러스터 또는 영역 클러스터의 노드에 대해 로드 한계를 구성할 수 있습니다. 명령줄, `clsetup` 유틸리티 또는 Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 로드 한계를 구성할 수 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오. 다음 절차에서는 명령줄을 사용하여 로드 한계를 구성하는 방법을 보여 줍니다.

▼ 노드에 대해 로드 한계를 구성하는 방법

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 전역 클러스터 노드 또는 영역 클러스터 노드에서 로드 한계를 만들고 구성하거나 기존 노드 로드 한계를 편집 또는 삭제할 수도 있습니다. Nodes(노드) 또는 Zone Clusters(영역 클러스터)를 누른 후 노드 이름을 눌러서 해당 페이지에 액세스합니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 전역 클러스터의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 로드 균형을 사용할 노드에 대해 로드 한계를 만들고 설정합니다.

다음 예제 명령에서 영역 클러스터 이름은 `zc1`입니다. 샘플 등록 정보는 `mem_load`이고 소프트 한계는 11이며 하드 로드 한계는 20입니다. 하드 및 소프트 한계는 선택 인수이며 특별히 정의하지 않는 한 기본값은 무제한으로 설정됩니다. 자세한 내용은 [clnode\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

```
# clnode create-loadlimit -p limitname=mem_load -Z zc1 \
-p softlimit=11 -p hardlimit=20 node1 node2 node3
```

3. 각 리소스 그룹에 로드 요소 값을 지정합니다.

다음 예제 명령에서 로드 비율은 리소스 그룹 두 개(rg1 및 rg2)에 설정됩니다. 로드 요소 설정은 노드의 정의된 로드 한계에 해당합니다.

```
# clresourcegroup set -p load_factors=mem_load@50,factor2@1 rg1 rg2
```

clresourcegroup create 명령을 사용하여 리소스 그룹을 만드는 동안 이 단계를 수행할 수도 있습니다. 자세한 내용은 [clresourcegroup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

4. 필요에 따라 하나 이상의 선택적인 추가 구성 작업을 수행합니다.

■ 기존 로드를 재배포합니다.

```
# clresourcegroup remaster rg1 rg2
```

이 명령은 리소스 그룹을 현재 마스터에서 다른 노드로 이동하여 로드를 균일하게 배포합니다.

■ 일부 리소스 그룹에 다른 리소스 그룹보다 높은 우선순위를 지정합니다.

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

기본 우선 순위는 500입니다. 우선 순위 값이 더 높은 리소스 그룹은 노드 할당에서 우선 순위가 낮은 리소스 그룹보다 우선하게 됩니다.

■ Preemption_mode 등록 정보를 설정합니다.

```
# clresourcegroup set -p Preemption_mode=No_cost rg1
```

HAS_COST, NO_COST, NEVER 옵션에 대한 자세한 내용은 [clresourcegroup\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

■ Concentrate_load 플래그를 설정합니다.

```
# cluster set -p Concentrate_load=TRUE
```

■ 리소스 그룹 간 유사성을 지정합니다.

강력한 양수 또는 음수 유사성은 로드 배포보다 우선합니다. 강력한 유사성을 위반하거나 하드 로드 한계로 제한할 수 없습니다. 강력한 유사성과 하드 로드 한계를 모두 설정한 경우 두 제약 조건을 충족할 수 없으면 일부 리소스 그룹이 강제로 오프라인 상태로 유지될 수 있습니다.

다음 예에서는 영역 클러스터 zc1의 리소스 그룹 rg1과 영역 클러스터 zc2의 리소스 그룹 rg2 사이에 강력한 양의 유사성을 지정합니다.

```
# clresourcegroup set -p RG_affinities=++zc2:rg2 zc1:rg1
```

5. 클러스터에서 모든 전역 클러스터 노드와 영역 클러스터 노드의 상태를 확인합니다.

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

출력에는 노드에 정의된 로드 한계 설정이 포함됩니다.

영역 클러스터 관리 작업 수행

영역 경로 이동, 응용 프로그램을 실행하도록 영역 클러스터 준비, 영역 클러스터 복제 등 영역 클러스터에서 기타 관리 작업을 수행할 수 있습니다. 이러한 모든 명령은 전역 클러스터의 노드에서 수행해야 합니다.

clsetup 유틸리티를 통해 영역 클러스터 구성 마법사를 실행하여 새 영역 클러스터를 만들거나 기존 영역 클러스터에 파일 시스템 또는 저장 장치를 추가할 수 있습니다. 프로파일을 구성하기 위해 clzonecluster install -c를 실행할 때 영역 클러스터의 영역이 구성됩니다. clsetup 유틸리티 또는 -c config_profile 옵션 사용 지침은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “영역 클러스터 만들기 및 구성”을 참조하십시오.

Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터를 만들거나, 파일 시스템 또는 저장 장치를 추가할 수도 있습니다. Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터의 Resource Security 등록 정보를 편집할 수 있습니다. Zone Clusters(영역 클러스터)를 누르고 영역 클러스터 이름을 눌러서 해당 페이지로 이동한 후 Solaris Resources(Solaris 리소스) 탭을 눌러서 영역 클러스터 구성요소를 관리하거나 Properties(등록 정보)를 눌러서 영역 클러스터 등록 정보를 관리합니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

주 - 전역 클러스터의 노드에서만 실행되는 Oracle Solaris Cluster 명령은 영역 클러스터에서 사용할 수 없습니다. 영역 클러스터에서의 올바른 명령 사용에 대한 자세한 내용은 해당 Oracle Solaris Cluster 매뉴얼 페이지를 참조하십시오.

표 19 기타 영역 클러스터 작업

작업	지침
새 영역 경로로 영역 경로 이동	<code>clzonecluster move -f zonepath zone-cluster-name</code>
응용 프로그램을 실행할 영역 클러스터 준비	<code>clzonecluster ready -n nodename zone-cluster-name</code>
통합 아카이브에서 노드 복원	통합 아카이브에서 노드를 복원하는 방법 [202]
통합 아카이브에서 영역 클러스터 구성 또는 설치	통합 아카이브에서 영역 클러스터를 구성하는 방법 [244] 통합 아카이브에서 영역 클러스터를 설치하는 방법 [245]
	명령 사용:
	<code>clzonecluster clone -Z target- zone-cluster-name [-m copymethod] source-zone-cluster-name</code>

작업	지침
영역 클러스터에 네트워크 주소 추가	영역 클러스터에 네트워크 주소를 추가하는 방법 [246]
영역 클러스터에 노드 추가	기존 클러스터 또는 영역 클러스터에 노드를 추가하는 방법 [200]
영역 클러스터에서 노드를 제거합니다	영역 클러스터에서 노드를 제거하는 방법 [207]
영역 클러스터 제거	영역 클러스터를 제거하는 방법 [247]
영역 클러스터에서 파일 시스템 제거	영역 클러스터에서 파일 시스템을 제거하는 방법 [248]
영역 클러스터에서 저장 장치 제거	영역 클러스터에서 저장 장치를 제거하는 방법 [251]
통합 아카이브에서 영역 클러스터 노드 복원	통합 아카이브에서 노드를 복원하는 방법 [202]
노드 제거 문제 해결	“노드 제거 문제 해결” [233]
Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관리	“Oracle Solaris Cluster SNMP 이벤트 MIB 만들기, 설정 및 관리” [234]

▼ 통합 아카이브에서 영역 클러스터를 구성하는 방법

clzonecluster 명령을 통해 대화식 유틸리티를 실행하여 Unified Archive에서 solaris10 또는 labeled 브랜드 영역 클러스터를 구성할 수 있습니다. clzonecluster configure 유틸리티에서는 복구 아카이브 또는 복제 아카이브를 지정할 수 있습니다.

대화식 유틸리티 대신 명령줄을 사용하여 아카이브에서 영역 클러스터를 구성하려는 경우 clzonecluster configure -f *command-file* 명령을 사용합니다. 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - 설치할 영역 클러스터가 지원되는 다른 방법으로 이미 구성된 경우 통합 아카이브에서 영역 클러스터를 구성할 필요가 없습니다.

1. 복구 또는 복제 아카이브를 만듭니다.

```
phys-schost# archiveadm create -r archive-location
```

create 명령을 사용하여 복제 아카이브를 만들거나 -r 옵션을 사용하여 복구 아카이브를 만듭니다. archiveadm 명령 사용에 대한 자세한 내용은 [archiveadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

2. 영역 클러스터를 호스트할 전역 클러스터의 노드에서 root 역할을 수행합니다.

3. 통합 아카이브의 복구 또는 복제된 아카이브에서 영역 클러스터를 구성합니다.

```
phys-schost-1# clzonecluster configure zone-cluster-name
```

`clzonecluster configure zone-cluster-name` 명령은 `create -a archive [other-options-such-as-"-x"]`를 지정할 수 있는 대화식 유틸리티를 실행합니다. 아카이브는 복제 아카이브 또는 복구 아카이브일 수 있습니다.

주 - 영역 클러스터 멤버를 구성에 추가해야만 영역 클러스터를 만들 수 있습니다.

`configure` 하위 명령은 `zonecfg` 명령을 사용하여 지정된 각 시스템에서 영역을 구성합니다. `configure` 하위 명령을 통해 영역 클러스터의 각 노드에 적용되는 등록 정보를 지정할 수 있습니다. 해당 등록 정보는 개별 영역에 대해 `zonecfg` 명령으로 설정된 것과 동일한 의미를 갖습니다. `configure` 하위 명령은 `zonecfg` 명령에서 확인할 수 없는 등록 정보 구성을 지원합니다. `-f` 옵션을 지정하지 않을 경우 `configure` 하위 명령은 대화식 셸을 실행합니다. `-f` 옵션은 명령 파일을 인수로 사용합니다. `configure` 하위 명령은 이 파일을 사용하여 영역 클러스터를 비대화식으로 만들거나 수정할 수 있습니다.

▼ 통합 아카이브에서 영역 클러스터를 설치하는 방법

통합 아카이브에서 영역 클러스터를 설치할 수 있습니다. `clzonecluster install` 유틸리티를 통해 설치에 사용할 아카이브 또는 Oracle Solaris 10 이미지 아카이브의 절대 경로를 지정할 수 있습니다. 지원되는 아카이브 유형과 관련된 자세한 내용은 [solaris10\(5\)](#) 매뉴얼 페이지를 참조하십시오. 영역 클러스터가 설치될 클러스터의 모든 물리적 노드에서 아카이브의 절대 경로에 액세스할 수 있어야 합니다. 통합 아카이브 설치의 복구 아카이브 또는 복제 아카이브를 사용할 수 있습니다.

대화식 유틸리티 대신 명령줄을 사용하여 아카이브에서 영역 클러스터를 설치하려는 경우 `clzonecluster create -a archive -z archived-zone` 명령을 사용합니다. 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

1. 복구 또는 복제 아카이브를 만듭니다.

```
phys-schost# archiveadm create -r archive-location
```

`create` 명령을 사용하여 복제 아카이브를 만들거나 `-r` 옵션을 사용하여 복구 아카이브를 만듭니다. `archiveadm` 명령 사용에 대한 자세한 내용은 [archiveadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

2. 영역 클러스터를 호스트할 전역 클러스터의 노드에서 root 역할을 수행합니다.

3. 통합 아카이브의 복구 또는 복제된 아카이브에서 영역 클러스터를 설치합니다.

```
phys-schost-1# clzonecluster install -a absolute_path_to_archive zone-cluster-name
```

영역 클러스터가 설치될 클러스터의 모든 물리적 노드에서 아카이브의 절대 경로에 액세스할 수 있어야 합니다. HTTPS 통합 아카이브 위치가 사용되는 경우 `-x cert|ca-cert|key=file`을 사용하여 SSL 인증서, CA(인증 기관) 인증서 및 키 파일을 지정합니다.

통합 아카이브에는 영역 클러스터 노드 리소스가 포함되지 않습니다. 클러스터가 구성되면 노드 리소스가 지정됩니다. Unified Archives를 사용하여 전역 영역에서 영역 클러스터를 구성하는 경우 영역 경로를 설정해야 합니다.

Unified Archive에 다중 영역이 포함되는 경우 *zone-cluster-name*을 사용하여 설치 소스의 영역 이름을 지정합니다. 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - Unified Archives를 만드는 데 사용한 소스에 Oracle Solaris Cluster 패키지가 포함되지 않을 경우 `pkg install ha-cluster-packages(ha-cluster-minimal 또는 ha-cluster-framework-full)`과 같은 특정 패키지 이름 대체를 실행해야 합니다. 영역을 부트하고 `zlogin` 명령을 실행한 다음 `pkg install` 명령을 실행해야 합니다. 이 작업을 수행하면 대상 영역 클러스터에 전역 클러스터와 동일한 패키지가 설치됩니다.

4. 새 영역 클러스터를 부트합니다.

```
phys-schost-1# clzonecluster boot zone-cluster-name
```

▼ 영역 클러스터에 네트워크 주소를 추가하는 방법

이 절차에서는 기존 영역 클러스터에 사용할 네트워크 주소를 추가합니다. 네트워크 주소는 영역 클러스터에서 논리 호스트 또는 공유 IP 주소 리소스를 구성하는 데 사용됩니다. `clsetup` 유틸리티를 여러 번 실행하여 네트워크 주소를 필요한 만큼 추가할 수 있습니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터에 네트워크 주소를 추가할 수도 있습니다. Zone Clusters(영역 클러스터)를 누르고, 영역 클러스터 이름을 눌러서 해당 페이지로 이동한 후, Solaris Resources(Solaris 리소스) 탭을 눌러서 영역 클러스터 구성요소를 관리합니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 영역 클러스터를 호스트하는 전역 클러스터의 노드에서 **root** 역할로 전환합니다.
2. 전역 클러스터에서, 영역 클러스터에서 사용할 클러스터 파일 시스템을 구성합니다.
`clsetup` 유틸리티를 시작합니다.

phys-schost# **clsetup**

주 메뉴가 표시됩니다.
3. Zone Cluster(영역 클러스터) 메뉴 항목을 선택합니다.
4. Add Network Address to a Zone Cluster(영역 클러스터에 네트워크 주소 추가) 메뉴 항목을 선택합니다.

5. 네트워크 주소를 추가할 영역 클러스터를 선택합니다.
6. 등록 정보를 선택하여 추가할 네트워크 주소를 지정합니다.

`address=value`

영역 클러스터에서 논리 호스트 또는 공유 IP 주소 리소스를 구성하는 데 사용되는 네트워크 주소를 지정합니다. 예: 192.168.100.101.

다음과 같은 유형의 네트워크 주소가 지원됩니다.

- 선택적으로 /와 접두어 길이가 뒤에 나오는 유효한 IPv4 주소.
- 뒤에 /와 접두어 길이가 필요한 유효한 IPv6 주소
- IPv4 주소로 확인되는 호스트 이름. IPv6 주소로 확인되는 호스트 이름은 지원되지 않습니다.

네트워크 주소에 대한 자세한 내용은 [zonecfg\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

7. 네트워크 주소를 더 추가하려면 `a`를 입력합니다.
 8. 구성 변경을 저장하려면 `c`를 입력합니다.
- 구성 변경 결과가 표시됩니다. 예를 들면 다음과 같습니다.

```
>>> 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. 모두 완료되면 `clsetup` 유틸리티를 종료합니다.

▼ 영역 클러스터를 제거하는 방법

특정 영역 클러스터를 삭제하거나 와일드카드를 사용하여 전역 클러스터에 구성된 모든 영역 클러스터를 제거할 수 있습니다. 제거하기 전에 영역 클러스터가 구성되어 있어야 합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하여 영역 클러스터를 삭제할 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 전역 클러스터의 노드에서 `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환합니다.
전역 클러스터의 한 노드에서 이 절차의 모든 단계를 수행합니다.
2. 영역 클러스터에서 모든 리소스 그룹과 해당 리소스를 삭제합니다.

```
phys-schost# clresourcegroup delete -F -Z zone-cluster-name +
```

주 - 이 단계는 전역 클러스터 노드에서 수행됩니다. 대신 영역 클러스터의 노드에서 이 단계를 수행하려면 영역 클러스터 노드에 로그인하고 명령에서 `-z zonecluster`를 생략합니다.

3. 영역 클러스터를 중지합니다.

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

4. 영역 클러스터를 제거합니다.

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

5. 영역 클러스터를 구성 해제합니다.

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

예 84 전역 클러스터에서 영역 클러스터 제거

```
phys-schost# clresourcegroup delete -F -Z SCZone +
phys-schost# clzonecluster halt SCZone
phys-schost# clzonecluster uninstall SCZone
phys-schost# clzonecluster delete SCZone
```

▼ 영역 클러스터에서 파일 시스템을 제거하는 방법

직접 마운트 또는 루프백 마운트를 사용하여 영역 클러스터로 파일 시스템을 내보낼 수 있습니다.

영역 클러스터는 다음에 대한 직접 마운트를 지원합니다.

- UFS 로컬 파일 시스템
- StorageTek QFS 독립형 파일 시스템
- StorageTek QFS 공유 파일 시스템(Oracle RAC 지원에 사용되는 경우)
- Oracle Solaris ZFS(데이터 세트로 내보냄)
- 지원되는 NAS 장치의 NFS

영역 클러스터는 다음에 대한 루프백 마운트를 관리할 수 있습니다.

- UFS 로컬 파일 시스템

- StorageTek QFS 독립형 파일 시스템
- StorageTek QFS 공유 파일 시스템(Oracle RAC 지원에 사용되는 경우에만)
- UFS 클러스터 파일 시스템

HASStoragePlus 또는 ScalMountPoint 리소스를 구성하여 파일 시스템의 마운트를 관리합니다. 영역 클러스터에 파일 시스템을 추가하는 지침은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “영역 클러스터에 파일 시스템 추가”를 참조하십시오.

HASStoragePlus 리소스는 ZFS 파일 시스템의 mountpoint 등록 정보가 none 또는 legacy로 설정되거나 canmount 등록 정보가 off로 설정된 경우 파일 시스템을 모니터하지 않습니다. 그 밖의 다른 ZFS 파일 시스템의 경우 HASStoragePlus 리소스 결함 모니터에서 파일 시스템이 마운트되는지 확인합니다. 파일 시스템이 마운트되면 HASStoragePlus 리소스는 IOOption 등록 정보 값 ReadOnly/ReadWrite에 따라 읽기/쓰기를 수행하여 파일 시스템의 접근성을 프로브합니다.

ZFS 파일 시스템이 마운트되지 않거나 파일 시스템의 프로브를 실패할 경우 리소스 결함 모니터를 실패하고 리소스가 Faulted로 설정됩니다. RGM은 리소스의 retry_count 및 retry_interval 등록 정보에 설정된 대로 다시 시작하려고 시도합니다. 위에 설명된 특정 mountpoint 및 canmount 등록 정보 설정이 작동하지 않으면 이 작업으로 파일 시스템이 다시 마운트될 수 있습니다. 결함 모니터가 계속 실패하고 retry_interval 내에서 retry_count를 초과하면 RGM은 리소스를 다른 노드로 페일오버합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터에서 파일 시스템을 제거할 수도 있습니다. Zone Clusters(영역 클러스터)를 누르고, 영역 클러스터 이름을 눌러서 해당 페이지로 이동한 후, Solaris Resources(Solaris 리소스) 탭을 눌러서 영역 클러스터 구성요소를 관리합니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. **영역 클러스터를 호스트하는 전역 클러스터의 노드에서 root 역할로 전환합니다.**
이 절차의 일부 단계는 전역 클러스터의 노드에서 수행됩니다. 다른 단계는 영역 클러스터의 노드에서 수행됩니다.
2. **제거할 파일 시스템과 관련된 리소스를 삭제합니다.**
 - a. **제거할 영역 클러스터의 파일 시스템에 대해 구성된 Oracle Solaris Cluster 리소스 유형(예: HASStoragePlus 및 SUNW.ScalMountPoint)을 식별하고 제거합니다.**


```
phys-schost# clresource delete -F -Z zone-cluster-name fs_zone_resources
```
 - b. **해당되는 경우 제거할 파일 시스템에 대해 전역 클러스터에 구성된 SUNW.qfs 유형의 Oracle Solaris Cluster 리소스를 식별하고 제거합니다.**

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

-F 옵션은 먼저 사용 안함으로 설정하지 않은 경우에도 지정한 모든 리소스의 삭제를 강제 실행하므로 주의해서 사용해야 합니다. 지정한 모든 리소스가 다른 리소스의 리소스 종속성 설정에서 제거되므로 클러스터에서 서비스 손실이 발생할 수 있습니다. 삭제되지 않은 종속 리소스는 잘못된 상태나 오류 상태로 남아 있을 수 있습니다. 자세한 내용은 [clresource\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

팁 - 제거된 리소스에 대한 리소스 그룹이 나중에 비게 되면 리소스 그룹을 삭제해도 됩니다.

3. 파일 시스템 마운트 지점 디렉토리에 대한 경로를 결정합니다.

예를 들면 다음과 같습니다.

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

4. 영역 클러스터 구성에서 파일 시스템을 제거합니다.

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

```
clzc:zone-cluster-name> remove fs dir=filesystemdirectory
```

```
clzc:zone-cluster-name> commit
```

파일 시스템 마운트 지점은 **dir=**로 지정됩니다.

5. 파일 시스템 제거를 확인합니다.

```
phys-schost# clzonecluster show -v zone-cluster-name
```

예 85 영역 클러스터에서 고가용성 로컬 파일 시스템 제거

이 예에서는 sczone이라는 영역 클러스터에 구성된 마운트 지점 디렉토리(/local/ufs-1)가 있는 파일 시스템을 제거하는 방법을 보여 줍니다. 리소스는 hasp-rs이며 HASToragePlus 유형입니다.

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

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

```
phys-schost# clresource delete -F -Z sczone hasp-rs
```

```
phys-schost# clzonecluster configure sczone
```

```
clzc:sczone> remove fs dir=/local/ufs-1
```

```
clzc:sczone> commit
```

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

예 86 영역 클러스터에서 고가용성 ZFS 파일 시스템 제거

이 예에서는 SUNW.HAStoragePlus 유형의 hasp-rs 리소스에 있는 sczone 영역 클러스터에 구성된 HAZpool이라는 ZFS 풀에서 ZFS 파일 시스템을 제거하는 방법을 보여 줍니다.

```
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
```

▼ 영역 클러스터에서 저장 장치를 제거하는 방법

영역 클러스터에서 Solaris Volume Manager 디스크 세트 및 DID 장치와 같은 저장 장치를 제거할 수 있습니다. 영역 클러스터에서 저장 장치를 제거하려면 이 절차를 수행합니다.

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터에서 저장 장치를 제거할 수도 있습니다. Zone Clusters(영역 클러스터)를 누르고, 영역 클러스터 이름을 눌러서 해당 페이지로 이동한 후, Solaris Resources(Solaris 리소스) 탭을 눌러서 영역 클러스터 구성요소를 관리합니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

1. 영역 클러스터를 호스트하는 전역 클러스터의 노드에서 **root** 역할로 전환합니다.

이 절차의 일부 단계는 전역 클러스터의 노드에서 수행됩니다. 다른 단계는 영역 클러스터의 노드에서 수행할 수 있습니다.

2. 제거할 장치와 관련된 리소스를 삭제합니다.

제거할 영역 클러스터의 장치에 대해 구성된 Oracle Solaris Cluster 리소스 유형(예: SUNW.HAStoragePlus 및 SUNW.ScalDeviceGroup)을 식별하고 제거합니다.

```
phys-schost# clresource delete -F -Z zone-cluster dev_zone_resources
```

3. 제거할 장치에 일치하는 항목을 확인합니다.

```
phys-schost# clzonecluster show -v zone-cluster
...
Resource Name:      device
match:              <device_match>
...
```

4. 영역 클러스터 구성에서 장치를 제거합니다.

```
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. 영역 클러스터를 재부트합니다.

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

6. 장치 제거를 확인합니다.

```
phys-schost# clzonecluster show -v zone-cluster
```

예 87 영역 클러스터에서 Solaris Volume Manager 디스크 세트 제거

이 예에서는 sczone이라는 영역 클러스터에 구성된 apachedg라는 Solaris Volume Manager 디스크 세트를 제거하는 방법을 보여 줍니다. apachedg 디스크 세트의 세트 번호는 3입니다. 이 장치는 클러스터에 구성된 zc_rs 리소스에서 사용됩니다.

```
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
```

예 88 영역 클러스터에서 DID 장치 제거

이 예에서는 sczone이라는 영역 클러스터에 구성된 DID 장치 d10 및 d11을 제거하는 방법을 보여 줍니다. 이 장치는 클러스터에 구성된 zc_rs 리소스에서 사용됩니다.

```
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

```

테스트 목적으로 사용할 수 있는 문제 해결 절차

이 절에서는 테스트를 목적으로 사용할 수 있는 문제 해결 절차에 대해 설명합니다.

전역 클러스터 외부에서 응용 프로그램 실행

▼ 비클러스터 모드로 부트된 노드에서 Solaris Volume Manager 메타 세트를 가져오는 방법

테스트 목적으로 전역 클러스터 외부의 응용 프로그램을 실행하려면 이 절차를 수행합니다.

1. 쿼럼 장치가 Solaris Volume Manager 메타 세트에서 사용되고 있는지 확인하고 쿼럼 장치에서 SCSI2 또는 SCSI3 예약을 사용하는지 확인합니다.

```
phys-schost# clquorum show
```

- a. 쿼럼 장치가 Solaris Volume Manager 메타 세트에 있는 경우 나중에 비클러스터 모드로 가져올 메타 세트의 일부가 아닌 새 쿼럼 장치를 추가합니다.

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

- b. 이전의 쿼럼 장치를 제거합니다.

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

- c. 쿼럼 장치에서 SCSI2 예약을 사용하는 경우 이전 쿼럼에서 SCSI2 예약을 스크럽하고 남아 있는 SCSI2 예약이 없는지 확인합니다.

다음 명령은 PGRE(Persistent Group Reservation Emulation) 키를 찾습니다. 디스크에 키가 없으면 `errno=22` 메시지가 표시됩니다.

```
# /usr/cluster/lib/sc/pgre -c pgre_inkeys -d /dev/did/rdisk/dids2
```

키를 찾은 후 PGRE 키를 스크럽합니다.

```
# /usr/cluster/lib/sc/pgre -c pgre_scrub -d /dev/did/rdisk/dids2
```



주의 - 디스크에서 활성 쿼럼 장치 키를 스크럽할 경우 다음 번 재구성에서 클러스터 패닉이 발생하고 작동 중인 쿼럼 상실 메시지가 표시됩니다.

2. 비클러스터 모드에서 부트할 전역 클러스터 노드를 비웁니다.

```
phys-schost# clresourcegroup evacuate -n target-node
```

3. HAStorage 또는 HAStoragePlus 리소스가 포함되어 있고 나중에 비클러스터 모드에서 가져올 메타 세트에 의해 영향을 받는 장치 또는 파일 시스템이 포함된 리소스 그룹을 오프라인으로 전환합니다.

```
phys-schost# clresourcegroup offline resource-group
```

4. 오프라인으로 전환한 리소스 그룹의 모든 리소스를 사용 안함으로 설정합니다.

```
phys-schost# clresource disable resource
```

5. 리소스 그룹을 관리 해제합니다.

```
phys-schost# clresourcegroup unmanage resource-group
```

6. 해당하는 장치 그룹을 오프라인으로 전환합니다.

```
phys-schost# cldevicegroup offline device-group
```

7. 장치 그룹을 사용 안함으로 설정합니다.

```
phys-schost# cldevicegroup disable device-group
```

8. 수동 노드를 비클러스터 모드로 부트합니다.

```
phys-schost# shutdown -g0 -i0 -y
ok> boot -x
```

9. 진행하기 전에 부트 프로세스가 수동 노드에서 완료되었는지 확인합니다.

```
phys-schost# svcs -x
```

10. 메타 세트의 디스크에 SCSI3 예약이 있는지 확인합니다.

메타 세트의 모든 디스크에서 다음 명령을 실행합니다.

```
phys-schost# /usr/cluster/lib/sc/scsi -c inkeys -d /dev/did/rdisk/dids2
```

11. SCSI3 예약이 디스크에 있는 경우 스크럽합니다.

```
phys-schost# /usr/cluster/lib/sc/scsi -c scrub -d /dev/did/rdisk/dids2
```

12. 비워진 노드에 메타 세트를 가져옵니다.

```
phys-schost# metaset -s name -C take -f
```

13. 메타 세트에 정의된 장치가 포함된 파일 시스템을 마운트합니다.

```
phys-schost# mount device mountpoint
```

14. 응용 프로그램을 시작하고 원하는 테스트를 수행합니다. 테스트를 완료하면 응용 프로그램을 중지합니다.

15. 노드를 재부트하고 부트 프로세스가 종료될 때까지 기다립니다.

```
phys-schost# reboot
```

16. 장치 그룹을 온라인으로 전환합니다.

```
phys-schost# cldevicegroup online -e device-group
```

17. 리소스 그룹을 시작합니다.

```
phys-schost# clresourcegroup online -eM resource-group
```

손상된 디스크 세트 복원

디스크 세트가 손상되었거나 클러스터의 노드가 디스크 세트를 소유할 수 없는 상태인 경우 이 절차를 수행합니다. 상태를 지우려는 시도가 실패한 경우 디스크 세트를 수정하기 위한 마지막 시도로 이 절차를 사용합니다.

이 절차는 Solaris Volume Manager 메타 세트 및 다중 소유자 Solaris Volume Manager 메타 세트에 적용됩니다.

▼ Solaris Volume Manager 소프트웨어 구성을 저장하는 방법

처음부터 디스크 세트를 복원하면 시간이 오래 걸리고 오류가 발생하기 쉽습니다. `metastat` 명령을 사용하여 복제본을 정기적으로 백업하거나 Oracle Explorer(SUNWexplo)를 사용하여 백업을 만드는 방법이 더 좋습니다. 그러면 저장된 구성을 사용하여 디스크 세트를 다시 만들 수 있습니다. `prvtoc` 및 `metastat` 명령을 사용하여 현재 구성을 파일에 저장한 다음 디스크 세트와 해당 구성요소를 다시 만들어야 합니다. [Solaris Volume Manager 소프트웨어 구성을 다시 만드는 방법 \[257\]](#)을 참조하십시오.

1. 디스크 세트의 각 디스크에 대한 파티션 테이블을 저장합니다.

```
# /usr/sbin/prvtoc /dev/global/rdisk/disk-name > /etc/lvm/disk-name.vtoc
```

2. Solaris Volume Manager 소프트웨어 구성을 저장합니다.

```
# /bin/cp /etc/lvm/md.tab /etc/lvm/md.tab_ORIGINAL
# /usr/sbin/metastat -p -s set-name >> /etc/lvm/md.tab
```

주 - /etc/vfstab 파일과 같은 다른 구성 파일에서 Solaris Volume Manager 소프트웨어를 참조할 수 있습니다. 이 절차에서는 동일한 Solaris Volume Manager 소프트웨어 구성이 재구축된다고 가정하므로 마운트 정보가 동일합니다. 세트를 소유하는 노드에서 Oracle Explorer(SUNWexplo)를 실행하면 prtvtoc 및 metaset -p 정보를 검색합니다.

▼ 손상된 디스크 세트를 비우는 방법

하나의 노드 또는 모든 노드에서 세트를 지우면 구성이 제거됩니다. 노드에서 디스크 세트를 비우려면 해당 노드가 디스크 세트를 소유하지 않아야 합니다.

1. 모든 노드에서 purge 명령을 실행합니다.

```
# /usr/sbin/metaset -s set-name -P
```

이 명령을 실행하면 Oracle Solaris Cluster 저장소뿐만 아니라 데이터베이스 복제본에서도 디스크 세트 정보가 제거됩니다. -P 및 -C 옵션을 사용하면 Solaris Volume Manager 환경을 완전히 재구축하지 않고도 디스크 세트를 비울 수 있습니다.

주 - 노드가 클러스터 모드에서 부트되는 동안 다중 소유자 디스크 세트를 비우면 DCS 구성 파일에서 정보를 제거해야 할 수 있습니다.

```
# /usr/cluster/lib/sc/dcs_config -c remove -s set-name
```

자세한 내용은 [dcs_config\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

2. 데이터베이스 복제본에서 디스크 세트 정보만 제거하려면 다음 명령을 사용합니다.

```
# /usr/sbin/metaset -s set-name -C purge
```

일반적으로 -c 옵션이 아닌 -p 옵션을 사용해야 합니다. -c 옵션을 사용하면 Oracle Solaris Cluster 소프트웨어가 여전히 디스크 세트를 인식하기 때문에 디스크 세트를 다시 만들 때 문제가 발생할 수 있습니다.

a. **metaset** 명령과 함께 -c 옵션을 사용한 경우 먼저 디스크 세트를 만들어 문제가 발생하는지 확인합니다.

b. 문제가 있으면 dcs 구성 파일에서 정보를 제거합니다.

```
# /usr/cluster/lib/sc/dcs_config -c remove -s setname
```

purge 옵션이 실패하면 최신 커널 및 메타 장치 업데이트를 설치했는지 확인하고 [My Oracle Support](#)에 문의하십시오.

▼ Solaris Volume Manager 소프트웨어 구성을 다시 만드는 방법

Solaris Volume Manager 소프트웨어 구성이 완전히 손실된 경우에만 이 절차를 사용합니다. 이 단계에서는 현재 Solaris Volume Manager 구성 및 해당 구성요소를 저장하고 손상된 디스크 세트를 비웠다고 가정합니다.

주 - 2 노드 클러스터에서만 중개자를 사용합니다.

1. 새 디스크 세트를 만듭니다.

```
# /usr/sbin/metaset -s set-name -a -h node1 node2
```

다중 소유자 디스크 세트인 경우 다음 명령을 사용해서 새 디스크 세트를 만듭니다.

```
/usr/sbin/metaset -s set-name -aM -h node1 node2
```

2. 세트를 만든 것과 동일한 호스트에서 필요한 경우 중재자 호스트를 추가합니다(2 노드만 해당).

```
/usr/sbin/metaset -s set-name -a -m node1 node2
```

3. 동일한 디스크를 동일한 이 호스트의 디스크 세트에 추가합니다.

```
/usr/sbin/metaset -s set-name -a /dev/did/rdisk/disk-name /dev/did/rdisk/disk-name
```

4. 디스크 세트를 비운 경우 다시 만들려면 VTOC(Volume Table of Contents)가 디스크에 유지되어야 합니다. 그러면 이 단계를 건너뛸 수 있습니다.

그러나 복구할 세트를 다시 만드는 경우 `/etc/lvm/disk-name.vtoc` 파일에 저장된 구성에 따라 디스크의 형식을 지정해야 합니다. 예를 들면 다음과 같습니다.

```
# /usr/sbin/fmthard -s /etc/lvm/d4.vtoc /dev/global/rdisk/d4s2
```

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

이 명령은 모든 노드에서 실행할 수 있습니다.

5. 각 메타 장치에 대한 기존 `/etc/lvm/md.tab` 파일에서 구문을 확인합니다.

```
# /usr/sbin/metainit -s set-name -n -a metadvice
```

6. 저장된 구성으로 각 메타 장치를 만듭니다.

```
# /usr/sbin/metainit -s set-name -a metadvice
```

7. 파일 시스템이 메타 장치에 있는 경우 `fsck` 명령을 실행합니다.

```
# /usr/sbin/fsck -n /dev/md/set-name/rdisk/metadvice
```

`fsck` 명령에 수퍼 블록 개수와 같은 몇 가지 오류만 표시되면 장치가 올바르게 재구성된 것입니다. 그러면 `-n` 옵션 없이 `fsck` 명령을 실행할 수 있습니다. 여러 가지 오류가 나타나면 메타

장치를 올바르게 재구성했는지 확인합니다. 올바르게 재구성한 경우 `fsck` 오류를 검토하여 파일 시스템을 복구할 수 있는지 확인합니다. 복구할 수 없는 경우 백업에서 데이터를 복구해야 합니다.

8. 모든 클러스터 노드의 다른 모든 메타 세트를 `/etc/lvm/md.tab` 파일에 연결(concatenate)한 다음 로컬 디스크 세트를 연결(concatenate)합니다.

```
# /usr/sbin/metastat -p >> /etc/lvm/md.tab
```

CPU 사용 제어 구성

CPU의 사용을 제어하려면 CPU 제어 기능을 구성해야 합니다. CPU 제어 기능 구성에 대한 자세한 내용은 [rg_properties\(5\)](#) 매뉴얼 페이지를 참조하십시오. 이 장에서는 다음 주제에 대한 정보를 제공합니다.

- “CPU 제어 소개” [259]
- “CPU 제어 구성” [260]

CPU 제어 소개

Oracle Solaris Cluster 소프트웨어는 CPU 사용을 제어할 수 있습니다.

CPU 제어 기능은 Oracle Solaris OS에서 사용 가능한 기능을 기반으로 구축됩니다. 영역, 프로젝트, 리소스 풀, 프로세서 세트 및 예약 클래스에 대한 자세한 내용은 [Oracle Solaris 영역 소개](#)를 참조하십시오.

Oracle Solaris OS에서는 다음을 수행할 수 있습니다.

- CPU 공유를 리소스 그룹에 할당
- 프로세서를 리소스 그룹에 할당

주 - Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용해서 영역 클러스터의 구성을 볼 수도 있습니다. Oracle Solaris Cluster Manager 로그인 지침은 [Oracle Solaris Cluster Manager에 액세스하는 방법 \[278\]](#)을 참조하십시오.

시나리오 선택

선택하는 구성 및 운영체제의 버전에 따라 CPU 제어 레벨이 달라집니다. 이 장에 설명된 CPU 제어의 모든 측면은 `automated`로 설정된 리소스 그룹 등록 정보 `RG_SLM_TYPE`에 따라 달라집니다.

페어 쉘어 스케줄러

리소스 그룹에 CPU 공유를 지정하는 절차의 첫번째 단계는 시스템의 스케줄러가 페어 쉘어 스케줄러(FSS)가 되도록 설정하는 것입니다. 기본적으로 Oracle Solaris OS의 일정 계획 클래스는 TS(Timesharing Schedule)입니다. 스케줄러를 FSS로 설정하여 공유 구성을 적용합니다.

선택하는 스케줄러 클래스와 관계없이 전용 프로세서 세트를 만들 수 있습니다.

CPU 제어 구성

이 절에서는 전역 클러스터 노드에서 CPU 사용을 제어하는 방법에 대해 설명합니다.

▼ 전역 클러스터 노드에서 CPU 사용을 제어하는 방법

이 절차를 수행하여 전역 클러스터 노드에서 실행될 리소스 그룹에 CPU 공유를 지정합니다.

리소스 그룹에 CPU 공유를 지정하면 Oracle Solaris Cluster 소프트웨어가 전역 클러스터 노드에서 리소스 그룹의 리소스를 시작할 때 다음 작업을 수행합니다.

- 아직 수행하지 않은 경우 노드에 지정된 CPU 공유 수(*zone.cpu-shares*)를 지정한 CPU 공유 수로 늘립니다.
- 아직 수행하지 않은 경우 노드에 *SCSLM_resource-group*이라는 프로젝트를 만듭니다. 이 프로젝트는 해당 리소스 그룹에만 해당하는 것이며 지정한 수의 CPU 공유(*project.cpu-shares*)가 지정됩니다.
- *SCSLM_resource-group* 프로젝트에서 리소스를 시작합니다.

CPU 제어 기능 구성에 대한 자세한 내용은 [rg_properties\(5\)](#) 매뉴얼 페이지를 참조하십시오.

1. 시스템의 기본 스케줄러를 페어 쉘어 스케줄러(FSS)로 설정합니다.

```
# dispadmin -d FSS
```

다음 부트 시 FSS가 기본 스케줄러가 됩니다. 이 구성을 즉시 적용하려면 `priocntl` 명령을 사용합니다.

```
# priocntl -s -C FSS
```

`priocntl` 및 `dispadmin` 명령 조합을 사용하면 FSS가 즉시 기본 스케줄러가 되고 재부트 후에도 계속 유지됩니다. 예약 클래스 설정에 대한 자세한 내용은 [dispadmin\(1M\)](#) 및 [priocntl\(1\)](#) 매뉴얼 페이지를 참조하십시오.

주 - FSS가 기본 스케줄러가 아닐 경우, CPU 공유 할당은 적용되지 않습니다.

2. 각 노드에서 CPU 제어를 사용하려면 전역 클러스터 노드에 대한 공유 수 및 기본 프로세서 세트에서 사용 가능한 최소 CPU 수를 구성합니다.

globalzoneshares 및 defaultpsetmin 등록 정보에 값을 지정하지 않으면 이러한 등록 정보의 기본값이 사용됩니다.

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

```
-p defaultpsetmin=integer
```

기본 프로세서 세트에서 사용 가능한 최소 CPU 수를 설정합니다. 기본값은 1입니다.

```
-p globalzoneshares=integer
```

노드에 지정되는 공유 수를 설정합니다. 기본값은 1입니다.

```
node
```

설정할 등록 정보의 노드를 지정합니다.

이러한 등록 정보를 설정하면 노드에 대한 등록 정보를 설정하는 것입니다.

3. 해당 등록 정보를 제대로 설정했는지 확인합니다.

```
# clnode show node
```

지정하는 노드에 대해 clnode 명령은 등록 정보 세트와 이러한 등록 정보에 대해 설정되는 값을 출력합니다. clnode를 사용하여 CPU 제어 등록 정보를 설정하지 않으면 기본값이 사용됩니다.

4. CPU 제어 기능을 구성합니다.

```
# clresourcegroup create -p RG_SLM_TYPE=automated [-p RG_SLM_CPU_SHARES=value] resource-group
```

```
-p RG_SLM_TYPE=automated
```

CPU 사용을 제어하고 시스템 리소스 관리를 위한 Oracle Solaris OS의 일부 구성 단계를 자동으로 수행할 수 있도록 해줍니다.

```
-p RG_SLM_CPU_SHARES=value
```

리소스 그룹 고유 프로젝트인 project.cpu-shares에 지정되는 CPU 공유 수를 지정하고 zone.cpu-shares 노드에 지정되는 CPU 공유 수를 결정합니다.

```
resource-group
```

리소스 그룹의 이름을 지정합니다.

이 절차에서는 RG_SLM_PSET_TYPE 등록 정보를 설정하지 않습니다. 노드에서 이 등록 정보는 default 값을 사용합니다.

이 단계에서는 리소스 그룹을 생성합니다. 또는 `clresourcegroup set` 명령을 사용하여 기존 리소스 그룹을 수정할 수 있습니다.

5. 구성 변경사항을 활성화합니다.

```
# clresourcegroup online -eM resource-group
```

resource-group 리소스 그룹의 이름을 지정합니다.

주 - `SCSLM_resource-group` 프로젝트를 제거하거나 수정하지 마십시오. 예를 들어 `project.max-lwps` 등록 정보를 구성하여 프로젝트에 더 많은 리소스 제어를 수동으로 추가할 수 있습니다. 자세한 내용은 [projmod\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

소프트웨어 업데이트

이 장에서는 다음 절에서 Oracle Solaris Cluster 소프트웨어를 업데이트하기 위한 정보와 지침을 제공합니다.

- “Oracle Solaris Cluster 소프트웨어 업데이트 개요” [263]
- “Oracle Solaris Cluster 소프트웨어 업데이트” [264]
- “패키지 제거” [269]

Oracle Solaris Cluster 소프트웨어 업데이트 개요

모든 클러스터 구성원 노드가 적절한 클러스터 작업을 수행하려면 동일한 업데이트가 적용되어야 합니다. 노드를 업데이트하는 경우 업데이트를 수행하기 전에 클러스터 멤버십에서 노드를 일시적으로 제거하거나 전체 클러스터를 중지해야 할 수 있습니다.

Oracle Solaris Cluster 소프트웨어를 업데이트하는 두 가지 방법이 있습니다.

- **업그레이드** - 클러스터를 최신 주 또는 부 Oracle Solaris Cluster 릴리스로 업그레이드하고 모든 패키지를 업데이트하여 Oracle Solaris OS를 업데이트합니다. 주 릴리스의 예는 Oracle Solaris Cluster 4.0에서 5.0으로 업그레이드하는 것입니다. 부 릴리스의 예는 Oracle Solaris Cluster 4.1에서 4.2로 업그레이드하는 것입니다.

scinstall 유틸리티 또는 `scinstall -u update` 명령을 실행하여 새 부트 환경(부트 가능한 이미지 인스턴스)을 만들고, 사용 중이 아닌 마운트 지점에 부트 환경을 마운트하고, 비트를 업데이트하고, 새 부트 환경을 활성화합니다. 처음에 복제 환경을 만들면 추가 공간 없이 순간적으로 발생합니다. 이 업데이트를 수행한 후에 클러스터를 재부트해야 합니다. 또한 Oracle Solaris OS를 최신 호환 버전으로 업그레이드합니다. 자세한 지침은 [Oracle Solaris Cluster 4.3 Upgrade Guide](#)를 참조하십시오.

solaris 브랜드 유형의 파일오버 영역이 있는 경우 [Oracle Solaris Cluster 4.3 Upgrade Guide](#)의 “How to Upgrade a solaris Branded Failover Zone”에 설명된 지침을 따르십시오.

영역 클러스터에 solaris10 브랜드 영역이 있는 경우 [Oracle Solaris Cluster 4.3 Upgrade Guide](#)의 “Upgrading a solaris10 Branded Zone in a Zone Cluster”에 설명된 업그레이드 지침을 따르십시오.

브랜드되지 않은 영역 클러스터를 업데이트하려면 “[특정 패키지 업데이트](#)” [265]의 절차를 따라 기본 전역 클러스터를 업데이트합니다. 전역 클러스터를 업데이트할 때 해당 영역 클러스터도 자동으로 업데이트됩니다.

주 - Oracle Solaris Cluster 코어 SRU를 적용하면 다른 Oracle Solaris Cluster 릴리스로 소프트웨어를 업그레이드하는 경우와 동일한 결과가 나오지 않습니다.

- **업데이트** - 특정 Oracle Solaris Cluster 패키지를 다양한 SRU 레벨로 업데이트합니다. pkg 명령 중 하나를 사용하여 SRU(Support Repository Update)의 IPS(Image Packaging System) 패키지를 업데이트할 수 있습니다.

SRU는 대개 정기적으로 릴리스되고 업데이트된 패키지 및 결함 수정을 포함합니다. 저장소에 모든 IPS 패키지 및 업데이트된 패키지가 들어 있습니다. pkg update 명령을 실행하면 Oracle Solaris 운영체제와 Oracle Solaris Cluster 소프트웨어가 모두 호환 버전으로 업데이트됩니다. 이 업데이트를 수행한 후에 클러스터를 재부트해야 할 수 있습니다. 지침은 [특정 패키지를 업데이트하는 방법](#) [266]을 참조하십시오.

Oracle Solaris Cluster 제품에 필요한 소프트웨어 업데이트를 보고 다운로드하려면 My Oracle Support의 등록된 사용자에게 문의하십시오. My Oracle Support 계정이 없으면 Oracle 서비스 담당자 또는 판매 담당 기술자에게 문의하거나 <http://support.oracle.com>에서 온라인으로 등록하십시오. 펌웨어 업데이트에 대한 내용은 하드웨어 설명서를 참조하십시오.

주 - 어떤 업데이트를 적용하거나 제거하기 전에 소프트웨어 업데이트 README를 읽어보십시오.

Oracle Solaris 패키지 관리 유틸리티인 pkg 사용에 대한 정보는 [Oracle Solaris 11.3의 소프트웨어 추가 및 업데이트](#)의 3 장, “[소프트웨어 패키지 설치 및 업데이트](#)”에서 제공됩니다.

Oracle Solaris Cluster 소프트웨어 업데이트

다음 표에서 Oracle Solaris Cluster 릴리스 또는 Oracle Solaris Cluster 소프트웨어의 패키지를 업그레이드/업데이트하는 방법을 확인하십시오.

표 20 Oracle Solaris Cluster 소프트웨어 업데이트

작업	지침
전체 클러스터를 새로운 주 또는 부 릴리스로 업그레이드	Oracle Solaris Cluster 4.3 Upgrade Guide 의 “How to Upgrade the Software (Standard Upgrade)”
성공적인 업데이트를 위한 팁 검토	“ 업데이트 팁 ” [265]
특정 패키지 업데이트	특정 패키지를 업데이트하는 방법 [266]

작업	지침
쿼럼 서버 또는 AI 설치 서버 업데이트	쿼럼 서버 또는 AI 설치 서버를 업데이트하는 방법 [269]
solaris 브랜드 영역 클러스터 업데이트	solaris 브랜드 영역 클러스터를 업데이트하는 방법 [267]
solaris10 브랜드 영역 클러스터 패치	solaris10 브랜드 영역 클러스터를 패치하는 방법 (clzonecluster) [267]
	영역 클러스터에서 solaris10 브랜드 영역을 패치하는 방법 (patchadd) [268]
	solaris10 브랜드 영역 클러스터에서 영역에 대해 대체 부트 환경을 패치하는 방법 [268]
Oracle Solaris Cluster 패키지 제거	패키지를 제거하는 방법 [269]
	쿼럼 서버 또는 AI 설치 서버 패키지를 제거하는 방법 [270]

업데이트 팁

Oracle Solaris Cluster 업데이트를 더 효율적으로 관리하려면 다음 팁을 사용합니다.

- 업데이트를 수행하기 전에 SRU의 README 파일을 읽어보십시오.
- 저장 장치의 업데이트 요구사항을 확인하십시오.
- 운영 환경에서 클러스터를 실행하기 전에 모든 업데이트를 적용하십시오.
- 하드웨어 펌웨어 레벨을 검사하고 필요한 펌웨어 업데이트를 설치하십시오. 펌웨어 업데이트에 대한 정보는 하드웨어 설명서를 참조하십시오.
- 클러스터 구성원 기능을 하는 모든 노드에 동일한 업데이트가 있어야 합니다.
- 클러스터 부속 시스템에 항상 최신 업데이트를 설치하십시오. 업데이트에는 볼륨 관리, 저장 장치 펌웨어 및 클러스터 전송 등이 포함되어 있습니다.
- 주요 업데이트 후에 파일오버를 테스트하십시오. 클러스터 작동이 저하되거나 기능이 떨어지면 업데이트를 취소하십시오.

특정 패키지 업데이트

IPS 패키지는 Oracle Solaris 11 운영체제와 함께 도입되었습니다. 각 IPS 패키지는 FMRI (Fault Managed Resource Indicator)로 설명되고 pkg 명령을 사용하여 SRU 업데이트를 수행할 수 있습니다. 다른 방법으로, scinstall -u 명령을 사용하여 SRU 업데이트를 수행할 수도 있습니다.

업데이트된 Oracle Solaris Cluster 데이터 서비스 에이전트를 사용하기 위해 특정 패키지의 업데이트가 필요할 수 있습니다.

▼ 특정 패키지를 업데이트하는 방법

업데이트하려는 각 전역 클러스터 노드에서 이 절차를 수행합니다. 클러스터 노드에서 브랜드 해제된 영역 클러스터는 자동으로 이 업데이트를 수신합니다.

1. **solaris.cluster.admin** RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 패키지를 업데이트합니다.

예를 들어, 특정 게시자의 패키지를 업데이트하려면 *pkg-fmri*에 게시자 이름을 지정합니다.

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



주의 - *pkg-fmri*를 지정하지 않은 채 `pkg update` 명령을 사용하면 업데이트 가능한 모든 설치된 패키지가 업데이트됩니다.

설치된 패키지의 최신 버전을 사용할 수 있고 나머지 이미지와 호환되는 경우 패키지가 해당 버전으로 업데이트됩니다. 패키지에 `reboot-needed` 플래그가 `true`로 설정된 이진 파일이 있는 경우 `pkg update pkg-fmri`를 실행하면 자동으로 새 부트 환경이 생성되므로 업데이트 후에 새 부트 환경으로 부트합니다. 업데이트 중인 패키지에 재부트를 강제하는 이진 파일이 없는 경우 `pkg update` 명령이 라이브 이미지를 업데이트하고 재부트가 필요하지 않습니다.

3. 데이터 서비스 에이전트(`ha-cluster/data-service/*` 또는 일반 데이터 서비스 에이전트인 `ha-cluster/ha-service/gds`)를 업데이트하는 경우 다음 명령을 실행합니다.

```
# pkg change-facet facet.version-lock.pkg-name=false
# pkg update pkg-name
```

예를 들면 다음과 같습니다.

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

주 - 에이전트가 업데이트되지 않도록 하려면 다음 명령을 실행합니다.

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

특정 에이전트 고정에 대한 자세한 내용은 [Oracle Solaris 11.3의 소프트웨어 추가 및 업데이트의 “선택적 구성 요소의 설치 제어”](#)를 참조하십시오.

4. 패키지가 업데이트되었는지 확인합니다.

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

영역 클러스터 업데이트 또는 패치 적용

solaris 브랜드 영역 클러스터를 업데이트하려면 `scinstall-u update` 명령을 사용하여 SRU를 적용합니다.

solaris10 브랜드 영역 클러스터를 패치하려면 `clzonecluster install-cluster -p` 명령 또는 `patchadd` 명령을 사용하거나 대체 부트 환경을 패치해서 패치를 적용합니다.

이 절에서는 다음 절차를 제공합니다.

- [solaris 브랜드 영역 클러스터를 업데이트하는 방법 \[267\]](#)
- [solaris10 브랜드 영역 클러스터를 패치하는 방법\(clzonecluster\) \[267\]](#)
- [영역 클러스터에서 solaris10 브랜드 영역을 패치하는 방법\(patchadd\) \[268\]](#)
- [solaris10 브랜드 영역 클러스터에서 영역에 대해 대체 부트 환경을 패치하는 방법 \[268\]](#)

▼ solaris 브랜드 영역 클러스터를 업데이트하는 방법

이 절차를 수행하면 `scinstall-u update` 명령을 사용해서 SRU를 적용하여 solaris 브랜드 영역 클러스터를 업데이트할 수 있습니다.

1. 전역 클러스터의 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할을 수행합니다.
2. 전역 클러스터 노드에서 전체 노드를 업데이트합니다.

```
phys-schost# scinstall -u update [-b be-name]
```

각 클러스터 노드에서 이 단계를 반복합니다.

3. 클러스터를 재부트합니다.

```
phys-schost# clzonecluster reboot
```

▼ solaris10 브랜드 영역 클러스터를 패치하는 방법(clzonecluster)

이 절차는 전역 클러스터의 한 노드에서 수행하십시오.

1. 전역 클러스터의 노드에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할을 수행합니다.
2. 영역 클러스터를 오프라인 실행 모드로 부트합니다.

- 영역 클러스터가 부트되지 않았으면 다음 명령을 실행합니다.

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

- 영역 클러스터가 부트되었으면 다음 명령을 실행합니다.

```
phys-schost# clzonecluster reboot -o zone-cluster
```

3. 전역 클러스터의 한 노드에서 전체 solaris10 브랜드 영역 클러스터를 패치합니다.

```
phys-schost# clzonecluster install-cluster -p patch-spec [options] zone-cluster
```

install-cluster 하위 명령에 대한 자세한 내용은 [clzonecluster\(1CL\)](#) 매뉴얼 페이지를 참조하십시오.

주 - pkg 명령이 오래되었다는 메시지와 함께 `scinstall -u update` 명령이 실패하면 메시지의 지침을 따릅니다.

4. 영역 클러스터를 재부트합니다.

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

▼ 영역 클러스터에서 solaris10 브랜드 영역을 패치하는 방법(patchadd)

구성된 각 영역 클러스터 노드에서 다음 단계를 수행합니다.

1. 영역 클러스터를 오프라인 실행 상태로 설정합니다.

```
phys-schost# clzonecluster reboot -o zone-cluster
```

2. 영역 클러스터 노드에 로그인합니다.

```
phys-schost# zlogin zone-cluster
```

3. 명령줄에서 patchadd를 입력합니다.

```
zchost# patchadd
```

4. 새 릴리스에 해당하는 패치를 설치합니다.

▼ solaris10 브랜드 영역 클러스터에서 영역에 대해 대체 부트 환경을 패치하는 방법

영역 클러스터 노드의 영역에 대해 대체 BE(부트 환경)를 패치하려면 이 절차를 따릅니다. 이 방법은 필요한 경우 패치된 BE를 취소하는 옵션을 제공합니다.

1. 패치하려는 영역 클러스터 노드의 영역에 대해 관리자로 전환합니다.
2. 새 부트 환경을 만들고, 패치하고, 활성화한 후 `shutdown` 명령을 사용해서 영역을 재부트합니다.
이러한 작업에 대해 [Oracle Solaris 10 영역 만들기 및 사용](#)의 “solaris10 브랜드 영역 부트 정보” 절, “How to Create and Activate Multiple Boot Environments on a solaris10 Branded Zone”에 포함된 지침을 따르십시오.
3. 패치해야 하는 다른 영역 클러스터 노드의 영역에 대해 위 단계를 반복합니다.

쿼럼 서버 또는 AI 설치 서버 업데이트

아래 절차에 따라 쿼럼 서버 또는 Oracle Solaris 11 AI(자동 설치 프로그램) 설치 서버에 대한 패키지를 업데이트할 수 있습니다. 쿼럼 서버에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris Cluster 쿼럼 서버 소프트웨어를 설치하고 구성하는 방법”을 참조하십시오. AI 사용에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 “Oracle Solaris 및 Oracle Solaris Cluster 소프트웨어를 설치 및 구성하는 방법(IPS 저장소)”을 참조하십시오.

▼ 쿼럼 서버 또는 AI 설치 서버를 업데이트하는 방법

1. `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 쿼럼 서버 또는 AI 설치 서버 패키지를 업데이트합니다.

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

설치된 `ha-cluster` 패키지의 최신 버전을 사용할 수 있고 나머지 이미지와 호환되는 경우 패키지가 해당 버전으로 업데이트됩니다.



주의 - `pkg update` 명령을 실행하면 시스템에 설치된 모든 `ha-cluster` 패키지가 업데이트됩니다.

패키지 제거

단일 패키지 또는 여러 패키지를 제거할 수 있습니다.

▼ 패키지를 제거하는 방법

1. `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. 하나 이상의 기존 패키지를 제거합니다.

```
# pkg uninstall pkg-fmri [pkg-fmri ...]
```

제거하려는 *pkg-fmri*에 다른 설치된 패키지가 종속되는 경우 `pkg uninstall` 명령을 실패합니다. *pkg-fmri*를 제거하려면 `pkg uninstall` 명령과 함께 모든 *pkg-fmri* 종속 항목을 제공해야 합니다. 패키지 제거에 대한 추가 정보는 [Oracle Solaris 11.3의 소프트웨어 추가 및 업데이트](#) 및 [pkg\(1\)](#) 매뉴얼 페이지를 참조하십시오.

▼ 쿼럼 서버 또는 AI 설치 서버 패키지를 제거하는 방법

1. `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.
2. 쿼럼 서버 또는 AI 설치 서버 패키지를 제거합니다.

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



주의 - 이 명령은 시스템에 설치된 모든 `ha-cluster` 패키지를 제거합니다.

클러스터 백업 및 복원

이 장은 다음 절로 구성되어 있습니다.

- “클러스터 백업” [271]
- “클러스터 파일 복원” [273]
- “클러스터 노드 복원” [202]

클러스터 백업

클러스터를 백업하기 전에 백업할 파일 시스템의 이름을 찾아보고, 전체 백업을 포함할 데이터 수량을 계산하고, ZFS 루트 파일 시스템을 백업하십시오.

표 21 작업 맵: 클러스터 파일 백업

작업	지침
미러 또는 플렉스 파일 시스템에 대해 온라인 백업 수행	미러를 온라인으로 백업하는 방법(Solaris Volume Manager) [271]
클러스터 구성 백업	클러스터 구성을 백업하는 방법 [273]
스토리지 디스크에 대한 디스크 분할 영역 구성 백업	사용 중인 스토리지 디스크에 대한 문서 참조

▼ 미러를 온라인으로 백업하는 방법(Solaris Volume Manager)

미러된 Solaris Volume Manager 볼륨은 마운트 해제하거나 전체 미러를 오프라인 상태로 전환하지 않고도 백업할 수 있습니다. 하위 미러 중 하나는 일시적으로 오프라인으로 전환하여 미러링을 제거해야 하지만 백업이 완료되면 바로 온라인으로 전환되어 동기화되므로 시스템이 중단되거나 데이터에 대한 사용자의 액세스를 거부하지 않습니다. 미러를 사용하여 온라인 백업을 수행하면 현재 작동하는 파일 시스템의 "스냅샷"이 백업됩니다.

lockfs 명령이 실행되기 직전에 프로그램에서 볼륨에 데이터를 쓰면 문제가 발생할 수 있습니다. 이 문제를 방지하려면, 이 노드에서 실행되는 모든 서비스를 일시적으로 중지하십시오. 또한 백업 절차를 수행하기 전에 클러스터가 오류 없이 실행되는지 확인합니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 백업할 클러스터 노드에서 동등한 역할로 전환합니다.
2. **metaset** 명령을 사용하여 백업된 볼륨에 소유권을 가진 노드를 확인합니다.

```
# metaset -s setname
```

-s *setname* 디스크 세트 이름을 지정합니다.

자세한 내용은 [metaset\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

3. **lockfs** 명령을 -w 옵션과 함께 사용하여 파일 시스템의 쓰기를 잠급니다.

```
# lockfs -w mountpoint
```

자세한 내용은 [lockfs\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

4. **metastat** 명령을 사용하여 하위 미러의 이름을 확인합니다.

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

-p md.tab 파일과 유사한 형식으로 상태를 표시합니다.

자세한 내용은 [metastat\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

5. **metadetach** 명령을 사용하여 미러에서 한 하위 미러를 오프라인으로 가져옵니다.

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

자세한 내용은 [metadetach\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

주 - 읽기 기능은 다른 하위 미러에서 계속 제공됩니다. 그러나 미러에 첫번째 쓰기 작업이 수행되면 그 때부터 오프라인 상태의 하위 미러에 대한 동기화가 수행되지 않습니다. 오프라인 상태의 하위 미러가 다시 온라인 상태로 전환되면 이러한 불일치 문제가 해결됩니다. fsck 명령을 수행할 필요가 없습니다.

6. **lockfs** 명령에 -u 옵션을 사용하여 파일 시스템 잠금을 해제하고 쓰기를 허용하여 계속합니다.

```
# lockfs -u mountpoint
```

7. 파일 시스템 검사를 수행합니다.

```
# fsck /dev/md/diskset/rdsk/submirror
```

- 오프라인 상태의 하위 미러를 테이프나 다른 백업 매체에 백업합니다.

주 - 하위 미러에 대해 블록 장치(/dsk) 이름이 아닌 원시 장치(/rdsk) 이름을 사용합니다.

- metattach** 명령을 사용하여 메타 장치 또는 볼륨을 다시 온라인으로 전환합니다.

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

메타 장치 또는 볼륨이 온라인으로 전환되면 자동으로 미러와 다시 동기화됩니다. 자세한 내용은 [metattach\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

- metastat** 명령을 사용하여 하위 미러가 재동기화되는지 확인합니다.

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

자세한 내용은 [Oracle Solaris 11.3의 ZFS 파일 시스템 관리](#)를 참조하십시오.

▼ 클러스터 구성을 백업하는 방법

클러스터 구성을 아카이브하고 사용 중인 클러스터 구성을 쉽게 복구하려면 클러스터 구성을 주기적으로 백업합니다. Oracle Solaris Cluster에서는 사용 중인 클러스터 구성을 XML (eXtensible Markup Language) 파일로 내보내는 기능을 제공합니다.

- 클러스터의 노드에 로그인하여 **solaris.cluster.read** RBAC 권한 부여를 제공하는 역할로 전환합니다.
- 클러스터 구성 정보를 파일로 내보냅니다.

```
# cluster export -o configfile
```

configfile 클러스터 명령이 클러스터 구성 정보를 내보내는 XML 구성 파일의 이름입니다. XML 구성 파일에 대한 자세한 내용은 [clconfiguration\(5CL\)](#) 매뉴얼 페이지를 참조하십시오.

- 클러스터 구성 정보를 XML 파일로 성공적으로 내보냈는지 확인합니다.

```
# pfedit configfile
```

클러스터 파일 복원

ZFS 루트 파일 시스템을 새 디스크로 복원할 수 있습니다.

Unified Archive에서 클러스터 또는 노드를 복원하거나, 특정 파일이나 파일 시스템을 복원할 수 있습니다.

파일이나 파일 시스템을 복원하려면 먼저 다음 정보를 알아야 합니다.

- 필요한 테이프
- 파일 시스템을 복원할 원시 장치 이름
- 사용 중인 테이프 드라이브의 유형
- 테이프 드라이브에 대한 장치 이름(로컬 또는 원격)
- 실패한 모든 디스크에 대한 분할 영역 구조(분할 영역 및 파일 시스템이 대체 디스크에 정확히 복사되어야 하기 때문임)

표 22 작업 맵: 클러스터 파일 복원

작업	지침
Solaris Volume Manager의 경우 ZFS 루트(/) 파일 시스템 복원	ZFS 루트(/) 파일 시스템을 복원하는 방법(Solaris Volume Manager) [274]

▼ ZFS 루트(/) 파일 시스템을 복원하는 방법(Solaris Volume Manager)

이 절차를 사용하여 ZFS 루트(/) 파일 시스템을 새 디스크로 복원합니다(예: 잘못된 루트 디스크 교체 후). 복원하는 노드를 부트하면 안됩니다. 복원 절차를 수행하기 전에 클러스터가 오류 없이 실행되는지 확인합니다. 루트 파일 시스템을 제외하면 UFS가 지원됩니다. UFS는 공유 디스크에서 Solaris Volume Manager 메타 세트의 메타 장치에 사용할 수 있습니다.

주 - 새 디스크는 오류가 발생한 디스크와 같은 형식으로 분할해야 하므로 이 절차를 시작하기 전에 분할 영역 구조를 확인하고 적절한 형식으로 파일 시스템을 다시 만듭니다.

phys-schost# 프롬프트는 전역 클러스터 프롬프트를 반영합니다. 전역 클러스터에서 이 절차를 수행합니다.

이 절차에서는 장문형 Oracle Solaris Cluster 명령을 제공합니다. 대부분의 명령에는 단문형도 있습니다. 명령은 명령 이름이 장문형과 단문형인 것을 제외하면 동일합니다.

1. 복원할 노드도 연결되어 있는 디스크 세트에 액세스할 권한이 있는 클러스터 노드에 **solaris.cluster.modify** RBAC 권한 부여를 제공하는 역할로 전환합니다.

복원할 노드가 아닌 다른 노드를 사용합니다.

2. 모든 메타 세트에서 복원 중인 노드의 호스트 이름을 제거합니다.

제거하는 노드가 아닌 메타 세트의 노드에서 이 명령을 실행합니다. 복원할 노드가 오프라인 상태이므로 시스템에는 RPC: Rpcbind failure - RPC: Timed out 오류가 표시됩니다. 이 오류를 무시하고 다음 단계를 수행합니다.

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

- s setname 디스크 세트 이름을 지정합니다.
- f 디스크 세트에서 마지막 호스트를 삭제합니다.
- d 디스크 세트에서 삭제합니다.
- h nodelist 디스크 세트에서 삭제할 노드의 이름을 지정합니다.

3. ZFS 루트 파일 시스템(/)을 복원합니다.

자세한 내용은 [Oracle Solaris 11.3의 ZFS 파일 시스템 관리](#)의 “ZFS 루트 풀에서 디스크 교체”를 참조하십시오.

ZFS 루트 풀 또는 루트 풀 스냅샷을 복구하려면 [Oracle Solaris 11.3의 ZFS 파일 시스템 관리](#)의 “ZFS 루트 풀에서 디스크 교체”에 설명된 절차를 따르십시오.

주 - /global/.devices/node@nodeid 파일 시스템을 만들어야 합니다.

/globaldevices 백업 파일이 백업 디렉토리에 존재하면 ZFS 루트 복원과 함께 복원됩니다. globaldevices SMF 서비스에 의해 파일이 자동으로 생성되지 않습니다.

4. 노드를 복수 사용자 모드로 재부트합니다.

```
# reboot
```

5. 장치 ID를 대체합니다.

```
# cldevice repair root-disk
```

6. metadb 명령을 사용하여 상태 데이터베이스 복제본을 다시 만듭니다.

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

-c copies 만들 복제본의 수를 지정합니다.

-f raw-disk-device 복제본을 만들 원시 디스크 장치입니다.

-a 복제본을 추가합니다.

자세한 내용은 [metadb\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

7. 복원된 노드가 아닌 다른 클러스터 노드에서 복원된 노드를 모든 디스크 세트에 추가합니다.

```
phys-schost-2# metaset -s setname -a -h nodelist
```

-a 호스트를 만들어 디스크 세트에 추가합니다.

노드가 클러스터 모드로 재부트됩니다. 이제 클러스터를 사용할 준비가 되었습니다.

예 89 ZFS 루트(/) 파일 시스템 복원(Solaris Volume Manager)

다음 예에서는 phys-schost-1 노드로 복원된 루트(/) 파일 시스템을 보여 줍니다. phys-schost-1 노드를 제거하고 나중에 schost-1 디스크 세트에 다시 추가하기 위해 클러스터의 다른 노드인 phys-schost-2에서 metaset 명령을 실행합니다. 다른 명령은 모두 phys-schost-1에서 실행됩니다. 새 부트 블록은 /dev/rdisk/c0t0d0s0에 만들어지고 상태 데이터베이스 복제본 세 개는 /dev/rdisk/c0t0d0s4에 다시 만들어집니다. 데이터 복원에 대한 자세한 내용은 [Oracle Solaris 11.3의 ZFS 파일 시스템 관리](#)의 “ZFS 저장소 풀의 데이터 문제 해결”을 참조하십시오.

메타 세트에서 노드 제거

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

오류 디스크 교체 후 노드 부트

Oracle Solaris 설명서의 절차에 따라 루트(/) 및 /usr 파일 시스템 복원

노드 재부트

```
# reboot
```

디스크 ID 교체

```
# cldevice repair /dev/dsk/c0t0d0
```

상태 데이터베이스 복제본 다시 만들기

```
# metadb -c 3 -af /dev/rdisk/c0t0d0s4
```

메타 세트에 다시 노드 추가

```
phys-schost-2# metaset -s schost-1 -a -h phys-schost-1
```

Oracle Solaris Cluster Manager 브라우저 인터페이스 사용

이 장에서는 클러스터를 여러 측면을 관리하는 데 사용할 수 있는 Oracle Solaris Cluster Manager 브라우저 인터페이스에 대해 설명합니다. 또한 Oracle Solaris Cluster Manager에 액세스하고 사용하기 위한 절차가 포함되어 있습니다.

주 - Oracle Solaris Cluster Manager는 Oracle Solaris Cluster 제품과 함께 제공되는 개인 버전의 Oracle GlassFish Server 소프트웨어를 사용합니다. 공용 버전의 Oracle GlassFish Server 소프트웨어를 설치하거나 패치 세트 업데이트를 시도하지 마십시오. 그렇게 하면 Oracle Solaris Cluster 소프트웨어를 업데이트하거나 Oracle Solaris Cluster SRU를 설치할 때 패키지 문제가 발생할 수 있습니다. Oracle Solaris Cluster Manager에 필요한 모든 Oracle GlassFish Server 개인 버전의 버그 수정은 Oracle Solaris Cluster SRU에 제공됩니다.

이 장의 내용:

- [“Oracle Solaris Cluster Manager 개요” \[277\]](#)
- [“Oracle Solaris Cluster Manager 소프트웨어 액세스” \[278\]](#)
- [“Oracle Solaris Cluster Manager에 대한 접근성 지원 구성” \[281\]](#)
- [“토폴로지를 사용하여 클러스터 모니터링” \[281\]](#)

Oracle Solaris Cluster Manager 개요

Oracle Solaris Cluster Manager 브라우저 인터페이스에서는 클러스터 정보를 그래픽으로 표시하고 클러스터 구성요소의 상태를 확인하며 구성 변경사항을 모니터링할 수 있습니다. 또한 Oracle Solaris Cluster Manager를 사용하면 다음 Oracle Solaris Cluster 구성요소 관리를 포함해서 여러 가지 관리 작업을 수행할 수 있습니다.

- 데이터 서비스
- 영역 클러스터
- 노드
- 개인 어댑터, 케이블 및 스위치

- 쿼럼 장치
- 장치 그룹
- 디스크
- NAS 장치
- 노드 로드 한계
- 리소스 그룹 및 리소스
- Geographic Edition 파트너십

Oracle Solaris Cluster Manager는 현재 모든 Oracle Solaris Cluster 관리 작업을 수행할 수 없습니다. 다른 작업에는 명령줄 인터페이스를 사용해야 합니다.

Oracle Solaris Cluster Manager 소프트웨어 액세스

Oracle Solaris Cluster Manager 브라우저 인터페이스에서는 Oracle Solaris Cluster 소프트웨어의 여러 작업을 쉽게 관리할 수 있습니다. 자세한 내용은 Oracle Solaris Cluster Manager 온라인 도움말을 참조하십시오.

클러스터를 부트하면 공통 에이전트 컨테이너가 자동으로 시작됩니다. 공통 에이전트 컨테이너가 실행 중인지 확인해야 하는 경우 [“Oracle Solaris Cluster Manager 문제 해결” \[279\]](#)을 참조하십시오.

팁 - Oracle Solaris Cluster Manager를 종료하기 위해 브라우저에서 Back(뒤로)을 누르지 마십시오.

▼ Oracle Solaris Cluster Manager에 액세스하는 방법

이 절차에서는 클러스터에서 Oracle Solaris Cluster Manager에 액세스하는 방법을 보여줍니다.

1. 클러스터 멤버에서 **root** 역할을 수행합니다.
2. 관리 콘솔 또는 클러스터 외부의 다른 시스템에서 브라우저를 시작합니다.
3. 브라우저의 디스크 및 메모리 캐시 크기는 0보다 큰 값으로 설정해야 합니다.
4. Java와 Javascript가 브라우저에서 사용으로 설정되었는지 확인합니다.
5. 브라우저에서 클러스터의 한 노드에 있는 Oracle Solaris Cluster Manager 포트에 연결합니다.
기본 포트 번호는 8998입니다.

`https://node:8998/scm`

6. 웹 브라우저에 나타나는 모든 인증서를 허용합니다.
Oracle Solaris Cluster Manager 로그인 페이지가 표시됩니다.
7. 관리할 클러스터 노드의 이름을 입력하거나 현재 클러스터를 관리하도록 localhost의 기본 값을 수락합니다.
8. 노드의 사용자 이름과 암호를 입력합니다.
9. Sign In(사인인)을 누릅니다.
Oracle Solaris Cluster Manager 응용 프로그램 실행 페이지가 표시됩니다.

주 - 클러스터가 2개 이상 구성된 경우 드롭다운 목록에서 Other(기타)를 선택하고 다른 클러스터에 로그인하여 해당 클러스터에 대한 정보를 확인할 수 있습니다. 클러스터가 파트너쉽 하나 이상의 멤버인 경우 Partnerships(파트너쉽) 폴더를 방문하면 모든 파트너 이름이 자동으로 드롭다운 목록에 추가됩니다. 인증된 후에는 Switch Cluster(클러스터 전환)를 선택할 수 있습니다.

Oracle Solaris Cluster Manager에 연결할 수 없으면 [“Oracle Solaris Cluster Manager 문제 해결” \[279\]](#)을 참조하십시오. Oracle Solaris 설치 시 제한된 네트워크 프로파일을 선택하는 경우 Oracle Solaris Cluster Manager에 대한 외부 액세스가 제한됩니다. Oracle Solaris Cluster Manager 브라우저 인터페이스를 사용하려면 이 네트워크가 필요합니다.

Oracle Solaris Cluster Manager 문제 해결

- 두 관리자 서비스가 실행 중인지 확인합니다.

```
# svcs system/cluster/manager\*
```

```
STATE      STIME      FMRI
online     Oct_30     svc:/system/cluster/manager-glassfish3:default
online     Oct_30     svc:/system/cluster/manager:default
```

svcadm 명령을 사용하여 system/cluster/manager-glassfish3을 사용 안함 또는 사용으로 설정합니다. 이 작업은 각각 애플리케이션 서버를 중지했다가 다시 시작합니다. system/cluster/manager를 온라인으로 유지해야 합니다. 사용 안함 또는 사용으로 설정할 필요는 없습니다.

- Oracle Solaris Cluster Manager에 연결할 수 없으면 `usr/sbin/cacoadm status`를 입력하여 공통 에이전트 컨테이너가 실행되고 있는지 확인합니다. 공통 에이전트 컨테이너가 실행 중이 아니면 로그인 페이지가 나타나지만 인증할 수 없습니다.
`/usr/sbin/cacoadm start`를 입력하여 공통 에이전트 컨테이너를 수동으로 시작할 수 있습니다.

▼ 공통 에이전트 컨테이너 보안 키를 구성하는 방법

Oracle Solaris Cluster Manager는 강력한 암호화 기술을 사용하여 Oracle Solaris Cluster Manager 웹 서버와 각 클러스터 노드 간 보안 통신을 보장합니다.

Oracle Solaris Cluster Manager에서 데이터 서비스 구성 마법사를 사용하거나 다른 Oracle Solaris Cluster Manager 작업을 수행하는 경우 Cacao 연결 오류가 발생할 수 있습니다. 이 절차에서는 공통 에이전트 컨테이너에 대한 보안 파일을 모든 클러스터 노드에 복사합니다. 따라서 공통 에이전트 컨테이너에 대한 보안 파일이 모든 클러스터 노드에서 동일하며 복사된 파일에서 올바른 파일 권한이 유지됩니다. 이 절차를 수행하면 보안 키가 동기화됩니다.

1. 각 노드에서 보안 파일 에이전트를 중지합니다.

```
phys-schost# /usr/sbin/cacaoadm stop
```

2. 한 노드에서 `/etc/cacao/instances/default/` 디렉토리로 변경합니다.

```
phys-schost-1# cd /etc/cacao/instances/default/
```

3. `/etc/cacao/instances/default/` 디렉토리를 tar 파일로 만듭니다.

```
phys-schost-1# tar cf /tmp/SECURITY.tar security
```

4. `/tmp/Security.tar` 파일을 각 클러스터 노드에 복사합니다.

5. `/tmp/SECURITY.tar` 파일을 복사한 각 노드에서 보안 파일을 추출합니다.

`/etc/cacao/instances/default/` 디렉토리에 이미 보안 파일이 있으면 덮어씁니다.

```
phys-schost-2# cd /etc/cacao/instances/default/
phys-schost-2# tar xf /tmp/SECURITY.tar
```

6. 보안 위험을 방지하려면 tar 파일의 각 복사본을 삭제합니다.

보안 위험을 방지하려면 tar 파일의 각 복사본을 삭제해야 합니다.

```
phys-schost-1# rm /tmp/SECURITY.tar
phys-schost-2# rm /tmp/SECURITY.tar
```

7. 각 노드에서 보안 파일 에이전트를 시작합니다.

```
phys-schost# /usr/sbin/cacaoadm start
```

▼ 네트워크 바인드 주소를 확인하는 방법

Oracle Solaris Cluster Manager가 실행되는 노드가 아닌 다른 노드에 대한 정보를 보려고 할 때 시스템 오류 메시지가 표시되면 공통 에이전트 컨테이너 `network-bind-address` 매개 변수가 올바른 값 `0.0.0.0`으로 설정되어 있는지 확인합니다.

클러스터의 각 노드에서 다음 단계를 수행합니다.

1. 네트워크 바인드 주소를 확인합니다.

```
phys-schost# cacaoadm list-params | grep network
network-bind-address=0.0.0.0
```

네트워크 바인드 주소가 0.0.0.0 이외의 값으로 설정되면 원하는 주소로 변경해야 합니다.

2. 변경 전후 Cacao를 중지하고 시작합니다.

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

Oracle Solaris Cluster Manager에 대한 접근성 지원 구성

Oracle Solaris Cluster Manager는 다음과 같은 접근성 설정을 제공합니다.

- 스크린 리더
- 고대비
- 큰 글꼴

이러한 컨트롤 중 하나 이상을 설정하려면 Oracle Solaris Cluster Manager 최상위 메뉴 표시줄에서 Accessibility(접근성)를 누릅니다. 그런 다음 사용 또는 사용 안함으로 설정하려는 접근성 컨트롤의 확인란을 누릅니다.

접근성 컨트롤 설정을 변경할 경우에는 로그인 세션이 만료된 후 해당 설정이 지속되지 않습니다. 동일한 로그인 세션 중에 다른 클러스터 노드에서 인증을 수행하면 설정이 지속됩니다.

토폴로지를 사용하여 클러스터 모니터링

토폴로지 보기에서 클러스터를 모니터링하고 문제를 식별할 수 있습니다. 객체 사이의 관계를 빠르게 살펴보고 각 노드에 속하는 리소스 및 리소스 그룹을 확인할 수 있습니다.

▼ 토폴로지를 사용하여 클러스터를 모니터하고 업데이트하는 방법

Topology(토폴로지) 페이지에 도달하려면 Oracle Solaris Cluster Manager에 로그인하여 Resource Groups(리소스 그룹)를 누르고 Topology(토폴로지) 탭을 누릅니다. 라인은 종속성 및 코로케이션 관계를 나타냅니다.

온라인 도움말에서 뷰의 요소에 대한 상세한 지침과 필터링 객체를 선택하는 방법을 제공합니다. 마우스 오른쪽 단추를 누르면 해당 객체의 작업에 관련된 컨텍스트 메뉴를 볼 수 있습니다. 옆에 있는 화살표를 누르면 온라인 도움말을 축소하거나 복원할 수 있습니다. 필터를 축소하거나 복원할 수도 있습니다.

다음 표는 Resource Topology(리소스 토폴로지) 페이지의 컨트롤 목록을 제공합니다.

컨트롤	기능	설명
	확대/축소	페이지의 일부분을 확대하거나 축소합니다.
	개요	다이어그램 위로 뷰포트를 끌어서 화면을 이동합니다.
	격리	리소스 그룹 또는 리소스를 한 번 눌러서 디스플레이에서 다른 모든 객체를 제거합니다.
	드릴	리소스 그룹을 한 번 눌러서 해당 리소스로 드릴합니다.
	재설정	격리/드릴 후에 전체 뷰로 돌아갑니다.
	필터	유형, 인스턴스, 상태별로 객체를 선택하여 범위를 좁힙니다.

다음 절차에서는 클러스터 노드에 치명적 오류가 있는지 모니터하는 방법을 보여줍니다.

1. Topology(토폴로지) 탭에서 Potential Masters(잠재적 마스터) 영역을 찾습니다.
2. 클러스터에서 각 노드의 상태를 보기 위해 확대합니다.
3. 빨간색 Critical(위기) 상태 아이콘이 표시된 노드를 찾아서 마우스 오른쪽 단추로 누르고 Show Details(세부 사항 표시)를 선택합니다.

4. 노드의 *Status*(상태) 페이지에서 *System Log*(시스템 로그)를 눌러 로그 메시지를 보고 필터링합니다.



배치 예제: Availability Suite 소프트웨어를 사용하여 클러스터 간 호스트 기반 데이터 복제 구성

이 부록에서는 Oracle Solaris Cluster Geographic Edition을 사용하지 않고 호스트 기반 복제를 사용하는 방법을 제공합니다. 호스트 기반 복제에 Oracle Solaris Cluster Geographic Edition을 사용하여 클러스터 간에 호스트 기반 복제의 구성 및 작업을 단순화 하십시오. [“데이터 복제 이해” \[93\]](#)를 참조하십시오.

이 부록의 예에서는 Oracle Solaris의 가용성 제품군 기능 소프트웨어를 사용하는 클러스터 간에 호스트 기반 데이터 복제를 구성하는 방법을 보여 줍니다. 이 예에서는 NFS 응용 프로그램의 전체 클러스터 구성을 보여주어 개별 작업을 수행할 수 있는 방법을 알려줍니다. 모든 작업은 전역 클러스터에서 수행되어야 합니다. 이 예에 다른 응용 프로그램이나 클러스터 구성에 필요한 단계가 모두 포함되어 있지는 않습니다.

역할 기반 액세스 제어(Role-based Access Control, RBAC)를 사용하여 클러스터 노드에 액세스할 경우, 모든 Oracle Solaris Cluster 명령에 대한 권한 부여를 제공하는 RBAC 역할로 전환할 수 있어야 합니다. 이러한 일련의 데이터 복제 절차에는 다음과 같은 Oracle Solaris Cluster RBAC 권한 부여가 필요합니다.

- `solaris.cluster.modify`
- `solaris.cluster.admin`
- `solaris.cluster.read`

RBAC 역할 사용에 대한 자세한 내용은 [Oracle Solaris 11.3의 사용자 및 프로세스 보안](#)을 참조하십시오. 각 Oracle Solaris Cluster 하위 명령에 필요한 RBAC 인증에 대해서는 Oracle Solaris Cluster 매뉴얼 페이지를 참조하십시오.

클러스터의 Availability Suite 소프트웨어 이해

이 절에서는 재해 허용 한계를 소개하고 Availability Suite 소프트웨어에서 사용하는 데이터 복제 방법에 대해 설명합니다.

재해 허용 한계는 기본 클러스터를 실패할 때 대체 클러스터에 응용 프로그램을 복원하는 기능입니다. 재해 허용 한계는 데이터 복제 및 테이크오버를 기반으로 합니다. 테이크오버는

하나 이상의 리소스 그룹 및 장치 그룹을 온라인으로 전환하여 응용 프로그램 서비스를 보조 클러스터로 재배치합니다.

기본 클러스터와 보조 클러스터 간에 동기적으로 데이터를 복제할 경우 기본 사이트를 실패할 때 커밋된 데이터가 손실되지 않습니다. 그러나 비동기적으로 데이터를 복제할 경우 기본 사이트를 실패할 때 일부 데이터가 보조 클러스터로 복제되지 않을 수 있으므로 데이터가 손실됩니다.

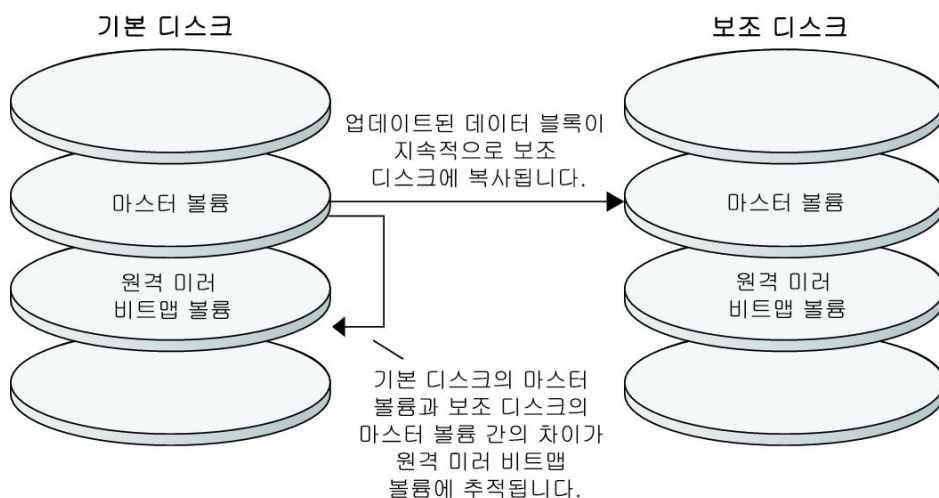
Availability Suite 소프트웨어에서 사용하는 데이터 복제 방법

이 절에서는 Availability Suite 소프트웨어가 사용하는 원격 미러 복제 방법 및 포인트 인 타임 스냅샷 방법을 설명합니다. 이 소프트웨어에서는 `sndradm` 및 `iiadm` 명령을 사용하여 데이터를 복제합니다. 자세한 내용은 [sndradm\(1M\)](#) 및 [iiadm\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

원격 미러 복제

[그림 2](#)에서는 원격 미러 복제를 보여 줍니다. 기본 디스크의 마스터 볼륨 데이터가 TCP/IP 연결을 통해 보조 디스크의 마스터 볼륨으로 복제됩니다. 원격 미러 비트맵이 1차 디스크의 마스터 볼륨과 2차 디스크의 마스터 볼륨 사이의 차이를 추적합니다.

그림 2 원격 미러 복제



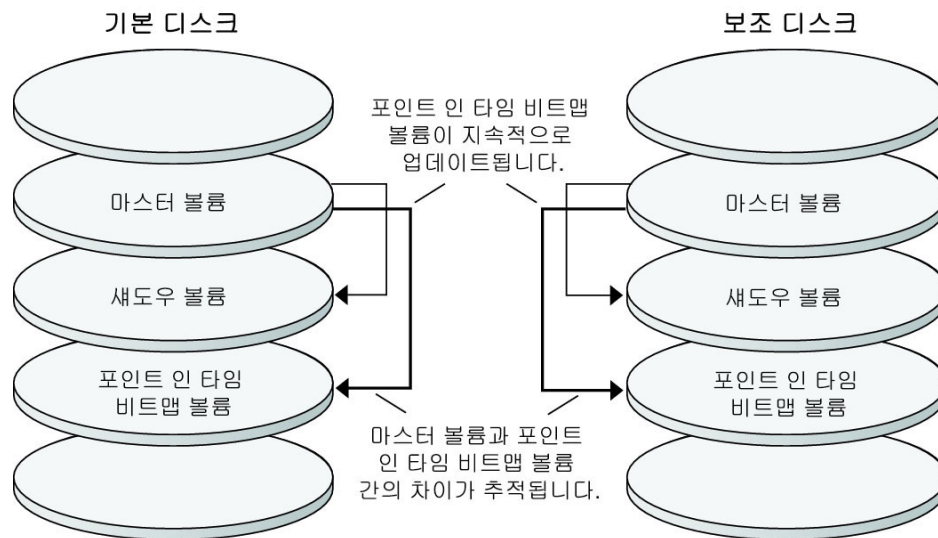
원격 미러 복제는 동기식 또는 비동기식으로 실시간 수행할 수 있습니다. 각 클러스터에 설정된 볼륨 각각은 동기식 복제나 비동기식 복제로 개별 구성될 수 있습니다.

- 동기식 데이터 복제에서는 원격 볼륨이 업데이트된 다음에야 쓰기 작업이 완료된 것으로 확인됩니다.
- 비동기식 데이터 복제에서는 원격 볼륨이 업데이트되기 전에 쓰기 작업이 완료된 것으로 확인됩니다. 비동기식 데이터 복제는 원거리 및 낮은 대역폭 환경에서 보다 융통성있게 활용할 수 있습니다.

포인트 인 타임 스냅샷

그림 3에서는 포인트 인 타임 스냅샷을 보여 줍니다. 각 디스크의 마스터 볼륨 데이터가 동일한 디스크의 새도우 볼륨으로 복사됩니다. 포인트 인 타임 비트맵은 마스터 볼륨과 새도우 볼륨 간의 차이를 추적합니다. 새도우 볼륨에 데이터가 복사되면 포인트 인 타임 비트맵이 다시 설정됩니다.

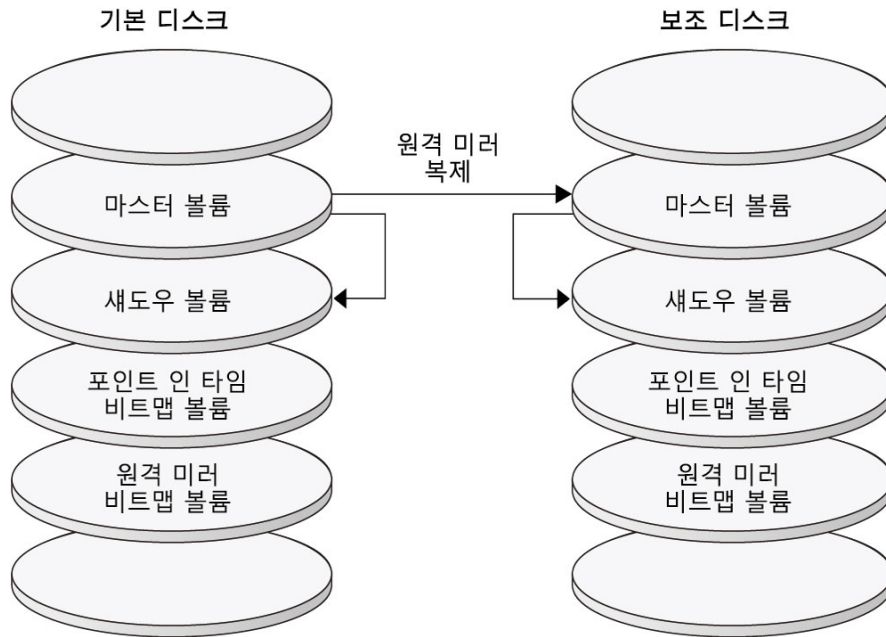
그림 3 포인트 인 타임 스냅샷



구성 예에서의 복제

그림 4에서는 이 구성 예에서 원격 미러 복제 및 포인트 인 타임 스냅샷을 사용하는 방법을 보여 줍니다.

그림 4 구성 예에서의 복제



클러스터 간 호스트 기반 데이터 복제 구성 지침

이 절에서는 클러스터 간에 데이터 복제를 구성하는 지침을 제공합니다. 또한 복제 리소스 그룹 및 응용 프로그램 리소스 그룹 구성에 대한 팁을 제공합니다. 클러스터에서 데이터 복제를 구성할 때 이 지침을 사용합니다.

이 절은 다음 내용으로 구성되어 있습니다.

- “복제 리소스 그룹 구성” [289]
- “응용 프로그램 리소스 그룹 구성” [289]
 - “폐일오버 응용 프로그램에서 리소스 그룹 구성” [290]
 - “확장 가능 응용 프로그램에서 리소스 그룹 구성” [291]
- “테이크오버 관리 지침” [292]

복제 리소스 그룹 구성

복제 리소스 그룹은 Availability Suite 소프트웨어 통제하의 장치 그룹을 논리 호스트 이름 리소스와 함께 배치합니다. 논리적 호스트 이름은 각 데이터 복제 스트림 끝에 존재해야 하고, 장치의 기본 I/O 경로와 동일한 클러스터 노드에 있어야 합니다. 복제 리소스 그룹은 다음과 같은 특징을 가져야 합니다.

- 파일오버 리소스 그룹 만들기

파일오버 리소스는 한 번에 한 노드에서만 실행할 수 있습니다. 파일오버가 발생하면 파일오버 리소스가 파일오버에 참여합니다.

- 논리 호스트 이름 리소스 보유

논리적 호스트 이름은 각 클러스터(기본 및 보조)의 한 노드에서 호스트되고, Availability Suite 소프트웨어 데이터 복제 스트림의 소스 및 대상 주소를 제공할 수 있습니다.

- HASToragePlus 리소스 소유

HASToragePlus 리소스는 복제 리소스 그룹이 스위치오버 또는 파일오버될 때 장치 그룹의 파일오버를 적용합니다. Oracle Solaris Cluster 소프트웨어 또한 장치 그룹이 스위치오버될 때 복제 리소스 그룹의 파일오버를 적용합니다. 이런 식으로 복제 리소스 그룹과 장치 그룹은 항상 동일한 노드에서 나란히 배열되거나 마스터됩니다.

다음 확장 등록 정보는 HASToragePlus 리소스에 정의되어 있어야 합니다.

- GlobalDevicePaths. 이 확장 등록 정보는 볼륨이 속한 장치 그룹을 정의합니다.

- AffinityOn 등록 정보 = True. 이 확장 등록 정보는 복제 리소스 그룹이 스위치오버 또는 파일오버될 때 해당 리소스 그룹을 스위치오버 또는 파일오버합니다. 이 기능을 유사성 스위치오버라고 합니다.

HASToragePlus에 대한 자세한 내용은 [SUNW.HASToragePlus\(5\)](#) 매뉴얼 페이지를 참조하십시오.

- 함께 배치되는 장치 그룹의 이름에 따라 이름 지정 후 -stor-rg 추가

예를 들어 devgrp-stor-rg입니다.

- 기본 클러스터 및 보조 클러스터 모두에서 온라인화

응용 프로그램 리소스 그룹 구성

고가용성을 제공하려면 응용 프로그램을 응용 프로그램 리소스 그룹의 리소스로 관리해야 합니다. 응용 프로그램 리소스 그룹은 파일오버 응용 프로그램이나 확장 가능 응용 프로그램으로 구성할 수 있습니다.

ZPoolsSearchDir 확장 등록 정보가 HASToragePlus 리소스에 정의되어야 합니다. 이 확장 등록 정보는 ZFS 파일 시스템을 사용하기 위해 필요합니다.

기본 클러스터에서 구성된 응용 프로그램 리소스 및 응용 프로그램 리소스 그룹은 보조 클러스터에서도 구성되어야 합니다. 또한 응용 프로그램 리소스에서 액세스하는 데이터는 보조 클러스터에 복제되어야 합니다.

이 절에서는 다음 응용 프로그램 리소스 그룹 구성에 대한 지침을 제공합니다.

- “[페일오버 응용 프로그램에서 리소스 그룹 구성](#)” [290]
- “[확장 가능 응용 프로그램에서 리소스 그룹 구성](#)” [291]

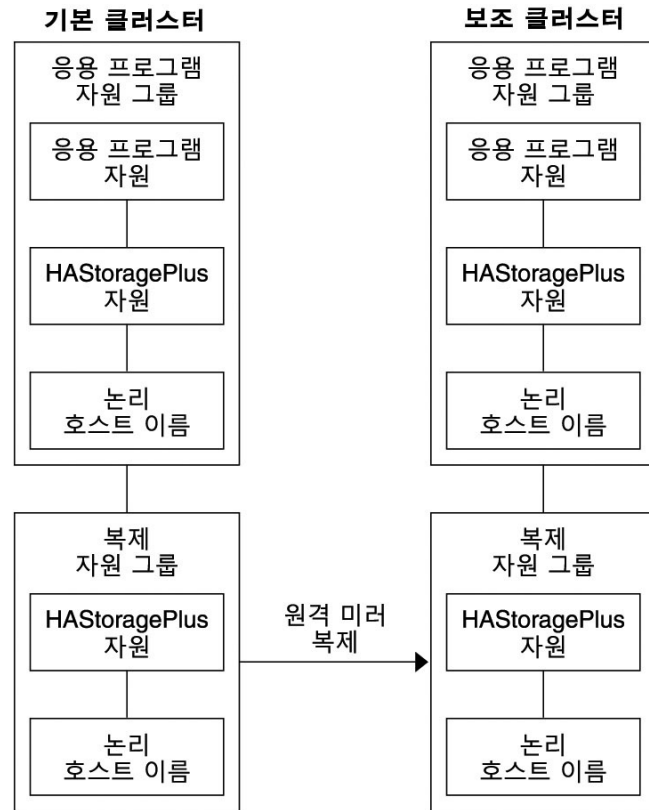
페일오버 응용 프로그램에서 리소스 그룹 구성

페일오버 응용 프로그램에서는 하나의 응용 프로그램이 한 번에 한 노드에서 실행됩니다. 해당 노드가 실패하면 응용 프로그램은 동일한 클러스터의 다른 노드로 페일오버됩니다. 페일오버 응용 프로그램의 리소스 그룹은 다음과 같은 특징을 가져야 합니다.

- 응용 프로그램 리소스 그룹이 스위치오버 또는 페일오버될 때 파일 시스템이나 zpools의 페일오버를 적용할 HAStoragePlus 리소스를 소유합니다.
장치 그룹은 복제 리소스 그룹 및 응용 프로그램 리소스 그룹과 함께 나란히 배열됩니다. 따라서 응용 프로그램 리소스 그룹이 페일오버되면 장치 그룹 및 복제 리소스 그룹도 페일오버됩니다. 응용 프로그램 리소스 그룹, 복제 리소스 그룹 및 장치 그룹은 동일한 노드에서 마스터됩니다.
그러나 장치 그룹이나 복제 리소스 그룹이 페일오버되면 응용 프로그램 리소스 그룹이 페일오버되지 않습니다.
 - 응용 프로그램 데이터가 전역으로 마운트될 경우, 응용 프로그램 리소스 그룹에 HAStoragePlus 리소스가 반드시 있을 필요는 없지만, 있는 것이 좋습니다.
 - 응용 프로그램 데이터가 로컬로 마운트될 경우, 응용 프로그램 리소스 그룹에 HAStoragePlus 리소스가 반드시 있어야 합니다.
- HAStoragePlus에 대한 자세한 내용은 [SUNW.HAStoragePlus\(5\)](#) 매뉴얼 페이지를 참조하십시오.
- 기본 클러스터에서는 온라인, 보조 클러스터에서는 오프라인이어야 합니다.
보조 클러스터가 기본 클러스터를 대신하는 경우 응용 프로그램 리소스 그룹은 보조 클러스터에서 온라인화되어야 합니다.

[그림 5](#)에서는 페일오버 응용 프로그램에 있는 응용 프로그램 리소스 그룹 및 복제 리소스 그룹의 구성을 보여 줍니다.

그림 5 페일오버 응용 프로그램에서 리소스 그룹 구성



확장 가능 응용 프로그램에서 리소스 그룹 구성

확장 가능 응용 프로그램에서는 하나의 응용 프로그램이 여러 노드에서 실행되어 단일한 논리 서비스를 만듭니다. 확장 가능 응용 프로그램을 실행하는 노드가 실패할 경우 페일오버가 발생하지 않습니다. 응용 프로그램은 다른 노드에서 계속 실행됩니다.

확장 가능 응용 프로그램이 응용 프로그램 리소스 그룹의 리소스로 관리되는 경우 응용 프로그램 리소스 그룹을 장치 그룹과 함께 배치할 필요는 없습니다. 따라서 응용 프로그램 리소스 그룹에 대해 HASToragePlus 리소스를 만들지 않아도 됩니다.

확장 가능 응용 프로그램의 리소스 그룹은 다음과 같은 특징을 가져야 합니다.

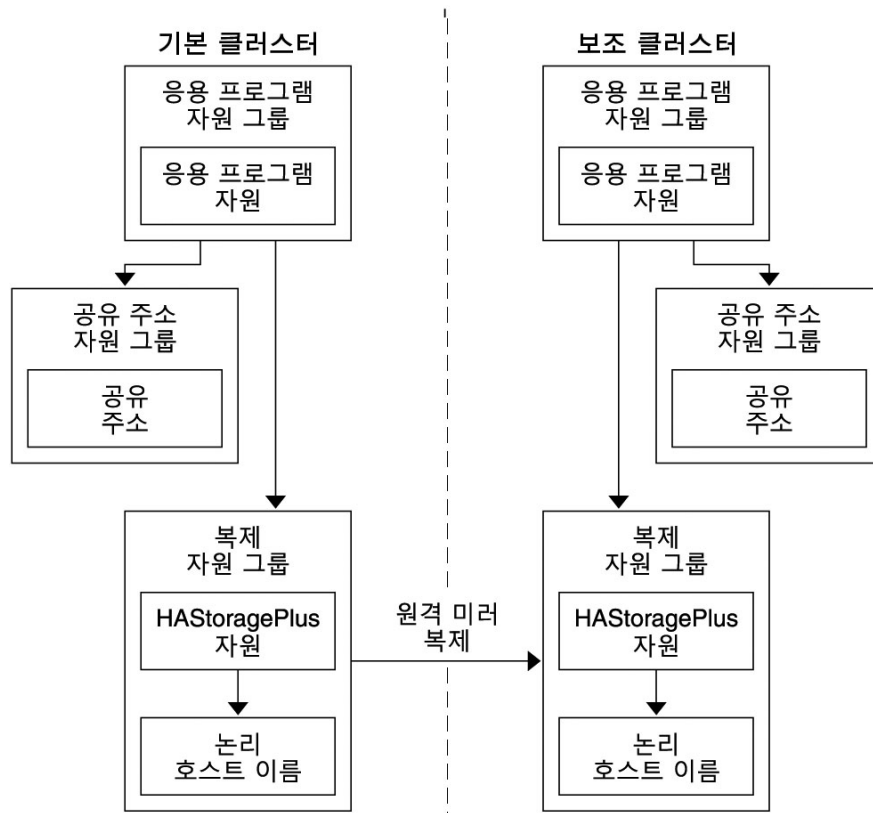
- 공유 주소 리소스 그룹에 대한 종속성 있음

확장 가능 응용 프로그램을 실행하는 노드는 들어오는 데이터를 분산하기 위해 공유 주소를 사용합니다.

- 기본 클러스터에서는 온라인, 보조 클러스터에서는 오프라인이어야 함

그림 6에서는 확장 가능 응용 프로그램에 있는 리소스 그룹의 구성을 보여 줍니다.

그림 6 확장 가능 응용 프로그램에서 리소스 그룹 구성

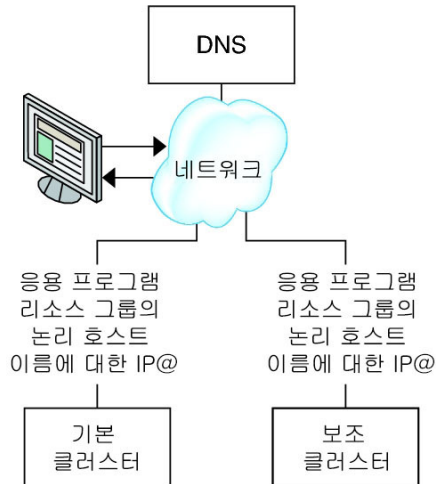


테이크오버 관리 지침

기본 클러스터가 실패하면 가능한 빨리 응용 프로그램이 보조 클러스터로 스위치오버되어야 합니다. 보조 클러스터로 스위치오버하려면 DNS를 업데이트해야 합니다.

클라이언트는 DNS를 사용하여 응용 프로그램의 논리 호스트 이름을 IP 주소와 매핑합니다. 응용 프로그램을 보조 클러스터로 이동하는 테이크오버를 마친 후에, 응용 프로그램의 논리 호스트 이름과 새 IP 주소 간의 매핑을 반영하도록 DNS 정보를 업데이트해야 합니다.

그림 7 클라이언트를 클러스터에 DNS 매핑



DNS를 업데이트하려면 `nsupdate` 명령을 사용합니다. 자세한 내용은 [nsupdate\(1M\)](#) 매뉴얼 페이지를 참조하십시오. 테이크오버를 관리하는 방법의 예는 [“테이크오버를 관리하는 방법의 예” \[319\]](#)를 참조하십시오.

복구 후에는 기본 클러스터를 다시 온라인으로 설정할 수 있습니다. 원래의 기본 클러스터로 다시 전환하려면 다음 작업을 수행합니다.

1. 기본 클러스터를 보조 클러스터와 동기화하여 기본 볼륨이 최신이 되게 합니다. 이를 위해 보조 노드에서 리소스 그룹을 중지하여 복제 데이터 스트림을 배출할 수 있도록 합니다.
2. 데이터 복제 방향을 반대로 합니다. 원래의 기본 클러스터가 현재가 되도록 다시 한번 원래의 보조 클러스터로 데이터를 복제합니다.
3. 기본 클러스터에서 리소스 그룹을 시작합니다.
4. 클라이언트가 기본 클러스터의 응용 프로그램에 액세스할 수 있도록 DNS를 업데이트합니다.

작업 맵: 데이터 복제 구성의 예

표 23. “작업 맵: 데이터 복제 구성의 예”는 이 예에서 NFS 응용 프로그램에 대해 Availability Suite 소프트웨어를 사용하여 데이터 복제를 구성하는 방법의 작업 목록입니다.

표 23 작업 맵: 데이터 복제 구성의 예

작업	지침
1. 클러스터를 연결하고 설치합니다.	“클러스터 연결 및 설치” [294]
2. 기본 클러스터 및 보조 클러스터에서 장치 그룹, NFS 응용 프로그램용 파일 시스템 및 리소스 그룹을 구성합니다.	“장치 그룹 및 리소스 그룹을 구성하는 방법의 예” [296]
3. 기본 클러스터 및 보조 클러스터에서 데이터 복제를 활성화합니다.	기본 클러스터에서 복제를 활성화하는 방법 [311] 보조 클러스터에서 복제를 활성화하는 방법 [312]
4. 데이터 복제를 수행합니다.	원격 미러 복제를 수행하는 방법 [314] 포인트 인 타임 스냅샷을 수행하는 방법 [315]
5. 데이터 복제 구성을 확인합니다.	복제가 올바르게 구성되었는지 확인하는 방법 [316]

클러스터 연결 및 설치

[그림 8](#)에서는 구성 예에서 사용하는 클러스터 구성을 보여 줍니다. 이 구성 예에서 보조 클러스터는 단일 노드를 포함하지만 다른 클러스터 구성도 사용할 수 있습니다.

그림 8 클러스터 구성 예

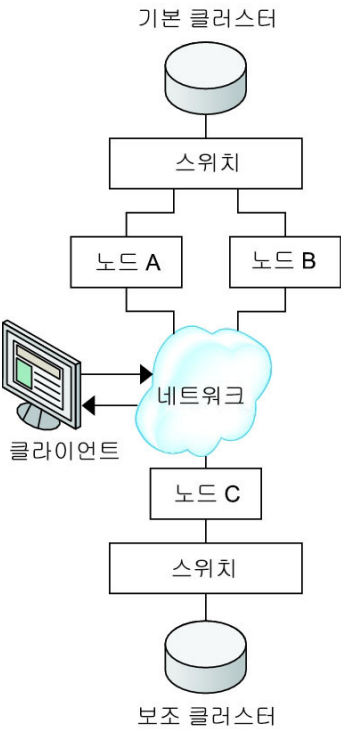


표 24. “필수 하드웨어 및 소프트웨어”에는 구성 예에 필요한 하드웨어 및 소프트웨어가 요약되어 있습니다. Availability Suite 소프트웨어 및 업데이트를 설치하기 전에 Oracle Solaris OS, Oracle Solaris Cluster 소프트웨어 및 볼륨 관리자 소프트웨어를 클러스터 노드에 설치해야 합니다.

표 24 필수 하드웨어 및 소프트웨어

하드웨어 또는 소프트웨어	요구사항
노드 하드웨어	Availability Suite 소프트웨어는 Oracle Solaris OS를 사용하는 모든 서버에서 지원됩니다. 사용할 하드웨어에 대한 자세한 내용은 Oracle Solaris Cluster Hardware Administration Manual 을 참조하십시오.
디스크 공간	약 15MB
Oracle Solaris OS	Oracle Solaris Cluster 소프트웨어에서 지원되는 Oracle Solaris OS 릴리스입니다. 모든 노드는 같은 Oracle Solaris OS 버전을 사용해야 합니다.

하드웨어 또는 소프트웨어	요구사항
	설치에 대한 자세한 내용은 Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서 를 참조하십시오.
Oracle Solaris Cluster 소프트웨어	최소 Oracle Solaris Cluster 4.1 소프트웨어 설치에 대한 자세한 내용은 Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서 를 참조하십시오.
볼륨 관리자 소프트웨어	Solaris Volume Manager 소프트웨어. 모든 노드에서 동일한 버전의 볼륨 관리자 소프트웨어를 사용해야 합니다. 설치에 대한 자세한 내용은 Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서 의 4 장, “Solaris Volume Manager 소프트웨어 구성”을 참조하십시오.
Availability Suite 소프트웨어	다양한 클러스터에서 다양한 버전의 Oracle Solaris OS 및 Oracle Solaris Cluster 소프트웨어를 사용할 수 있지만, 클러스터 간에 동일한 버전의 Availability Suite 소프트웨어를 사용해야 합니다. 소프트웨어 설치 방법에 대한 자세한 내용은 Availability Suite 소프트웨어 릴리스의 설치 설명서를 참조하십시오.
Availability Suite 소프트웨어 업데이트	최신 소프트웨어 업데이트 정보를 보려면 My Oracle Support 에 로그인하십시오.

장치 그룹 및 리소스 그룹을 구성하는 방법의 예

이 절에서는 NFS 응용 프로그램에서 장치 그룹 및 리소스 그룹을 구성하는 방법에 대해 설명합니다. 추가 정보는 “복제 리소스 그룹 구성” [289] 및 “응용 프로그램 리소스 그룹 구성” [289]을 참조하십시오.

이 절에서는 다음 절차에 대해 설명합니다.

- [기본 클러스터에 장치 그룹을 구성하는 방법](#) [297]
- [보조 클러스터에 장치 그룹을 구성하는 방법](#) [298]
- [NFS 응용 프로그램에서 기본 클러스터 파일 시스템을 구성하는 방법](#) [300]
- [NFS 응용 프로그램에서 보조 클러스터 파일 시스템을 구성하는 방법](#) [301]
- [기본 클러스터에서 복제 리소스 그룹을 만드는 방법](#) [302]
- [보조 클러스터에서 복제 리소스 그룹을 만드는 방법](#) [303]
- [기본 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법](#) [305]
- [보조 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법](#) [308]
- [복제가 올바르게 구성되었는지 확인하는 방법](#) [316]

다음 표에서는 구성 예에서 만든 그룹 및 리소스의 이름 목록을 표시합니다.

표 25 구성 예의 그룹 및 리소스 요약

그룹 또는 리소스	이름	설명
장치 그룹	devgrp	장치 그룹

그룹 또는 리소스	이름	설명
복제 리소스 그룹 및 리소스	devgrp-stor-rg	복제 리소스 그룹
	lhost-reprg-prim, lhost-reprg-sec	기본 및 보조 클러스터에서 복제 리소스 그룹의 논리 호스트 이름
	devgrp-stor	복제 리소스 그룹의 HASToragePlus 리소스
응용 프로그램 리소스 그룹 및 리소스	nfs-rg	응용 프로그램 리소스 그룹
	lhost-nfsrg-prim, lhost-nfsrg-sec	기본 및 보조 클러스터에서 응용 프로그램 리소스 그룹의 논리 호스트 이름
	nfs-dg-rs	응용 프로그램의 HASToragePlus 리소스
	nfs-rs	NFS 리소스

devgrp-stor-rg를 제외한 그룹 및 리소스의 이름은 예로 든 것이며 필요에 따라 변경할 수 있습니다. 복제 리소스 그룹의 이름은 *devicegroupname-stor-rg* 형식이어야 합니다.

Solaris Volume Manager 소프트웨어에 대한 자세한 내용은 [Oracle Solaris Cluster 4.3 소프트웨어 설치 설명서](#)의 4 장, “Solaris Volume Manager 소프트웨어 구성”을 참조하십시오.

▼ 기본 클러스터에 장치 그룹을 구성하는 방법

시작하기 전에 다음 작업을 완료했는지 확인합니다.

- 다음 절의 지침 및 요구사항을 읽습니다.
 - “클러스터의 Availability Suite 소프트웨어 이해” [285]
 - “클러스터 간 호스트 기반 데이터 복제 구성 지침” [288]
- “클러스터 연결 및 설치” [294]에 설명된 대로 기본 및 보조 클러스터를 설정합니다.

1. `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환하여 **nodeA**에 액세스합니다.

노드 nodeA는 기본 클러스터의 첫번째 노드입니다. 어떤 노드가 nodeA인지 미리 알리려면 [그림 8](#)을 참조하십시오.

2. NFS 데이터 및 연관된 복제를 포함하는 메타 세트를 만듭니다.

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

3. 메타 세트에 디스크를 추가합니다.

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

4. 메타 세트에 중재자를 추가합니다.

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

5. 필요한 볼륨(또는 메타 장치)을 만듭니다.

미러의 두 구성요소를 만듭니다.

```
nodeA# metainit -s nfsset d101 1 1 /dev/did/dsk/d6s2
nodeA# metainit -s nfsset d102 1 1 /dev/did/dsk/d7s2
```

구성요소 중 하나로 미러를 만듭니다.

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

나머지 구성요소를 미러에 연결하고 동기화하도록 허용합니다.

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

다음과 같이 미러에서 소프트 분할 영역을 만듭니다.

■ d200 - NFS 데이터(마스터 볼륨):

```
nodeA# metainit -s nfsset d200 -p d100 50G
```

■ d201 - NFS 데이터의 포인트 인 타임 복사본 볼륨:

```
nodeA# metainit -s nfsset d201 -p d100 50G
```

■ d202 - 포인트 인 타임 비트맵 볼륨:

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

■ d203 - 원격 새도우 비트맵 볼륨:

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

■ d204 - Oracle Solaris Cluster SUNW.NFS 구성 정보의 볼륨:

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

6. NFS 데이터 및 구성 볼륨에 대한 파일 시스템을 만듭니다.

```
nodeA# yes | newfs /dev/md/nfsset/rdisk/d200
nodeA# yes | newfs /dev/md/nfsset/rdisk/d204
```

다음 순서 [보조 클러스터에 장치 그룹을 구성하는 방법 \[298\]](#)으로 이동합니다.

▼ 보조 클러스터에 장치 그룹을 구성하는 방법

시작하기 전에 [기본 클러스터에 장치 그룹을 구성하는 방법 \[297\]](#) 절차를 완료합니다.

1. `solaris.cluster.modify` RBAC 권한 부여를 제공하는 역할로 전환하여 nodeC에 액세스합니다.
2. NFS 데이터 및 연관된 복제를 포함하는 메타 세트를 만듭니다.

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

3. 메타 세트에 디스크를 추가합니다.

아래 예에서 디스크 DID 번호가 서로 다르다고 가정합니다.

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

주 - 단일 노드 클러스터에는 중재자가 필요하지 않습니다.

4. 필요한 볼륨(또는 메타 장치)을 만듭니다.

미러의 두 구성요소를 만듭니다.

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

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

구성요소 중 하나로 미러를 만듭니다.

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

나머지 구성요소를 미러에 연결하고 동기화하도록 허용합니다.

```
metattach -s nfsset d100 d102
```

다음과 같이 미러에서 소프트 분할 영역을 만듭니다.

■ d200 - NFS 데이터 마스터 볼륨:

```
nodeC# metainit -s nfsset d200 -p d100 50G
```

■ d201 - NFS 데이터의 포인트 인 타임 복사본 볼륨:

```
nodeC# metainit -s nfsset d201 -p d100 50G
```

■ d202 - 포인트 인 타임 비트맵 볼륨:

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

■ d203 - 원격 새도우 비트맵 볼륨:

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

■ d204 - Oracle Solaris Cluster SUNW.NFS 구성 정보의 볼륨:

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

5. NFS 데이터 및 구성 볼륨에 대한 파일 시스템을 만듭니다.

```
nodeC# yes | newfs /dev/md/nfsset/rdsk/d200
```

```
nodeC# yes | newfs /dev/md/nfsset/rdsk/d204
```

다음 순서 [NFS 응용 프로그램에서 기본 클러스터 파일 시스템을 구성하는 방법 \[300\]](#)으로 이동합니다.

▼ NFS 응용 프로그램에서 기본 클러스터 파일 시스템을 구성하는 방법

시작하기 전에 [보조 클러스터에 장치 그룹을 구성하는 방법 \[298\]](#) 절차를 완료합니다.

1. **nodeA 및 nodeB에서 solaris.cluster.admin RBAC 권한 부여를 제공하는 역할로 전환합니다.**

2. **nodeA 및 nodeB에서 NFS 파일 시스템에 대한 마운트 지점 디렉토리를 만듭니다.**
예를 들면 다음과 같습니다.

```
nodeA# mkdir /global/mountpoint
```

3. **nodeA 및 nodeB에서 마운트 지점에 자동으로 마운트되지 않도록 마스터 볼륨을 구성합니다.**
nodeA 및 nodeB의 /etc/vfstab 파일에서 다음 텍스트를 추가 또는 대체합니다. 텍스트는 한 행이어야 합니다.

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \  
/global/mountpoint ufs 3 no global,logging
```

4. **nodeA 및 nodeB에서 메타 장치 d204의 마운트 지점을 만듭니다.**
다음 예에서는 마운트 지점 /global/etc를 만듭니다.

```
nodeA# mkdir /global/etc
```

5. **nodeA 및 nodeB에서 마운트 지점에 자동으로 마운트되도록 메타 장치 d204를 구성합니다.**
nodeA 및 nodeB의 /etc/vfstab 파일에서 다음 텍스트를 추가 또는 대체합니다. 텍스트는 한 행이어야 합니다.

```
/dev/md/nfsset/dsk/d204 /dev/md/nfsset/rdisk/d204 \  
/global/etc ufs 3 yes global,logging
```

6. **nodeA에 메타 장치 d204를 마운트합니다.**

```
nodeA# mount /global/etc
```

7. **Oracle Solaris Cluster HA for NFS 데이터 서비스에 대한 구성 파일 및 정보를 만듭니다.**

- a. **nodeA에 /global/etc/SUNW.nfs라는 디렉토리를 만듭니다.**

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

- b. **nodeA에 /global/etc/SUNW.nfs/dfstab.nfs-rs 파일을 만듭니다.**

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. **nodeA의 /global/etc/SUNW.nfs/dfstab.nfs-rs 파일에 다음 행을 추가합니다.**

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

다음 순서 [NFS 응용 프로그램에서 보조 클러스터 파일 시스템을 구성하는 방법 \[301\]](#)으로 이동합니다.

▼ NFS 응용 프로그램에서 보조 클러스터 파일 시스템을 구성하는 방법

시작하기 전에 [NFS 응용 프로그램에서 기본 클러스터 파일 시스템을 구성하는 방법 \[300\]](#) 절차를 완료합니다.

1. nodeC에서 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 전환합니다.

2. nodeC에서 NFS 파일 시스템의 마운트 지점 디렉토리를 만듭니다.

예를 들면 다음과 같습니다.

```
nodeC# mkdir /global/mountpoint
```

3. nodeC에서 마운트 지점에 자동으로 마운트되도록 마스터 볼륨을 구성합니다.

nodeC의 `/etc/vfstab` 파일에서 다음 텍스트를 추가 또는 대체합니다. 텍스트는 한 행이어야 합니다.

```
/dev/md/nfsset/dsk/d200 /dev/md/nfsset/rdisk/d200 \
/global/mountpoint ufs 3 yes global,logging
```

4. nodeA에 메타 장치 d204를 마운트합니다.

```
nodeC# mount /global/etc
```

5. Oracle Solaris Cluster HA for NFS 데이터 서비스에 대한 구성 파일 및 정보를 만듭니다.

- a. nodeA에 `/global/etc/SUNW.nfs`라는 디렉토리를 만듭니다.

```
nodeC# mkdir -p /global/etc/SUNW.nfs
```

- b. nodeA에 `/global/etc/SUNW.nfs/dfstab.nfs-rs` 파일을 만듭니다.

```
nodeC# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

- c. nodeA의 `/global/etc/SUNW.nfs/dfstab.nfs-rs` 파일에 다음 행을 추가합니다.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

다음 순서 [기본 클러스터에서 복제 리소스 그룹을 만드는 방법 \[302\]](#)으로 이동합니다.

▼ 기본 클러스터에서 복제 리소스 그룹을 만드는 방법

시작하기 전에 ■ [NFS 응용 프로그램에서 보조 클러스터 파일 시스템을 구성하는 방법 \[301\]](#) 절차를 완료합니다.

■ `/etc/netmasks` 파일에 모든 논리적 호스트 이름에 대한 IP 주소 서브넷과 넷마스크 항목이 있는지 확인합니다. 필요하다면 `/etc/netmasks` 파일을 편집하여 누락된 항목을 추가합니다.

1. `solaris.cluster.modify`, `solaris.cluster.admin`, `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 `nodeA`에 액세스합니다.

2. `SUNW.HAStoragePlus` 리소스 유형을 등록합니다.

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

3. 장치 그룹의 복제 리소스 그룹을 만듭니다.

```
nodeA# clresourcegroup create -n nodeA,nodeB devgrp-stor-rg
```

```
-n nodeA,nodeB
```

클러스터 노드 `nodeA` 및 `nodeB`에서 복제 리소스 그룹을 마스터할 수 있도록 지정합니다.

```
devgrp-stor-rg
```

복제 리소스 그룹의 이름입니다. 이 이름에서 `devgrp`는 장치 그룹의 이름을 지정합니다.

4. `SUNW.HAStoragePlus` 리소스를 복제 리소스 그룹에 추가합니다.

```
nodeA# clresource create -g devgrp-stor-rg -t SUNW.HAStoragePlus \
-p GlobalDevicePaths=nfsset \
-p AffinityOn=True \
devgrp-stor
```

```
-g
```

리소스를 추가할 리소스 그룹을 지정합니다.

```
-p GlobalDevicePaths=
```

Availability Suite 소프트웨어에서 사용하는 장치 그룹을 지정합니다.

```
-p AffinityOn=True
```

`SUNW.HAStoragePlus` 리소스가 `-p GlobalDevicePaths=`에 정의된 전역 장치 및 클러스터 파일 시스템에 대해 유사성 스위치오버를 수행하도록 지정합니다. 따라서 복제 리소스 그룹이 페일오버하거나 스위치오버되면 관련 장치 그룹도 스위치오버됩니다.

이러한 확장 등록 정보에 대한 자세한 내용은 [SUNW.HAStoragePlus\(5\)](#) 매뉴얼 페이지를 참조하십시오.

5. 복제 리소스 그룹에 논리 호스트 이름 리소스를 추가합니다.

```
nodeA# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-prim
```

기본 클러스터에 있는 복제 리소스 그룹의 논리 호스트 이름은 lhost-reprg-prim입니다.

6. 리소스를 활성화하고 리소스 그룹을 관리 및 온라인화합니다.

```
nodeA# clresourcegroup online -emM -n nodeA devgrp-stor-rg
```

```
-e
```

연관된 리소스를 사용으로 설정합니다.

```
-M
```

리소스 그룹을 관리합니다.

```
-n
```

리소스 그룹을 온라인으로 전환할 노드를 지정합니다.

7. 리소스 그룹이 온라인 상태인지 확인합니다.

```
nodeA# clresourcegroup status devgrp-stor-rg
```

리소스 그룹 상태 필드를 검사하여 복제 리소스 그룹이 nodeA에서 온라인 상태인지 확인합니다.

다음 순서 [보조 클러스터에서 복제 리소스 그룹을 만드는 방법 \[303\]](#)으로 이동합니다.

▼ 보조 클러스터에서 복제 리소스 그룹을 만드는 방법

- 시작하기 전에
- [기본 클러스터에서 복제 리소스 그룹을 만드는 방법 \[302\]](#) 절차를 완료합니다.
 - /etc/netmasks 파일에 모든 논리적 호스트 이름에 대한 IP 주소 서브넷과 넷마스크 항목이 있는지 확인합니다. 필요하다면 /etc/netmasks 파일을 편집하여 누락된 항목을 추가합니다.

1. solaris.cluster.modify, solaris.cluster.admin, solaris.cluster.read RBAC 권한 부여를 제공하는 역할로 nodeC에 액세스합니다.

2. SUNW.HAStoragePlus를 리소스 유형으로 등록합니다.

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

3. 장치 그룹의 복제 리소스 그룹을 만듭니다.

```
nodeC# clresourcegroup create -n nodeC devgrp-stor-rg
```

```
create
```

리소스 그룹을 만듭니다.

-n
리소스 그룹의 노드 목록을 지정합니다.

devgrp
장치 그룹의 이름입니다.

devgrp-stor-rg
복제 리소스 그룹의 이름입니다.

4. 복제 리소스 그룹에 **SUNW.HAStoragePlus** 리소스를 추가합니다.

```
nodeC# clresource create \  
-t SUNW.HAStoragePlus \  
-p GlobalDevicePaths=nfsset \  
-p AffinityOn=True \  
devgrp-stor
```

create
리소스를 만듭니다.

-t
리소스 유형을 지정합니다.

-p **GlobalDevicePaths=**
Availability Suite 소프트웨어에서 사용하는 장치 그룹을 지정합니다.

-p **AffinityOn=True**
SUNW.HAStoragePlus 리소스가 -p GlobalDevicePaths=에 정의된 전역 장치 및 클러스터 파일 시스템에 대해 유사성 스위치오버를 수행하도록 지정합니다. 따라서 복제 리소스 그룹이 페일오버하거나 스위치오버되면 관련 장치 그룹도 스위치오버됩니다.

devgrp-stor
복제 리소스 그룹의 HAStoragePlus 리소스

이러한 확장 등록 정보에 대한 자세한 내용은 [SUNW.HAStoragePlus\(5\)](#) 매뉴얼 페이지를 참조하십시오.

5. 복제 리소스 그룹에 논리 호스트 이름 리소스를 추가합니다.

```
nodeC# clreslogicalhostname create -g devgrp-stor-rg lhost-reprg-sec
```

보조 클러스터에 있는 복제 리소스 그룹의 논리적 호스트 이름은 lhost-reprg-sec입니다.

6. 리소스를 활성화하고 리소스 그룹을 관리 및 온라인화합니다.

```
nodeC# clresourcegroup online -eM -n nodeC devgrp-stor-rg
```

online

온라인으로 전환합니다.

-e

연관된 리소스를 사용으로 설정합니다.

-M

리소스 그룹을 관리합니다.

-n

리소스 그룹을 온라인으로 전환할 노드를 지정합니다.

7. 리소스 그룹이 온라인 상태인지 확인합니다.

```
nodeC# clresourcegroup status devgrp-stor-rg
```

리소스 그룹 상태 필드를 검사하여 복제 리소스 그룹이 nodeC에서 온라인 상태인지 확인합니다.

다음 순서 [기본 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법 \[305\]](#)으로 이동합니다.

▼ 기본 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법

이 절차에서는 NFS에서 응용 프로그램 리소스 그룹을 만드는 방법에 대해 설명합니다. 이 절차는 이 응용 프로그램에만 해당되며, 다른 유형의 응용 프로그램에서는 사용할 수 없습니다.

시작하기 전에 ■ [보조 클러스터에서 복제 리소스 그룹을 만드는 방법 \[303\]](#) 절차를 완료합니다.

■ /etc/netmasks 파일에 모든 논리적 호스트 이름에 대한 IP 주소 서브넷과 넷마스크 항목이 있는지 확인합니다. 필요하다면 /etc/netmasks 파일을 편집하여 누락된 항목을 추가합니다.

1. `solaris.cluster.modify`, `solaris.cluster.admin`, `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 nodeA에 액세스합니다.
2. `SUNW.nfs`를 리소스 유형으로 등록합니다.

```
nodeA# clresourcetype register SUNW.nfs
```

3. `SUNW.HAStoragePlus`가 리소스 유형으로 등록되지 않은 경우 등록합니다.

```
nodeA# clresourcetype register SUNW.HAStoragePlus
```

4. NFS 서비스의 응용 프로그램 리소스 그룹을 만듭니다.

```
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
```

그룹의 리소스에서 관리 파일을 기록할 수 있는 디렉토리를 지정합니다.

```
Auto_start_on_new_cluster=False
```

응용 프로그램 리소스 그룹이 자동으로 시작되지 않도록 지정합니다.

```
RG_affinities=+++devgrp-stor-rg
```

응용 프로그램 리소스 그룹과 함께 배치할 리소스 그룹을 지정합니다. 이 예에서 응용 프로그램 리소스 그룹은 복제 리소스 그룹 devgrp-stor-rg와 함께 배치해야 합니다.

복제 리소스 그룹이 새로운 기본 노드로 스위치오버될 경우 응용 프로그램 리소스 그룹은 자동으로 스위치오버됩니다. 그러나 응용 프로그램 리소스 그룹을 새로운 기본 노드로 스위치오버하려고 시도하면 해당 동작이 배치 요구사항을 위반하므로 스위치오버가 차단됩니다.

```
nfs-rg
```

응용 프로그램 리소스 그룹의 이름

5. 응용 프로그램 리소스 그룹에 SUNW.HAStoragePlus 리소스를 추가합니다.

```
nodeA# clresource create -g nfs-rg \  
-t SUNW.HAStoragePlus \  
-p FileSystemMountPoints=/global/mountpoint \  
-p AffinityOn=True \  
nfs-dg-rs
```

```
create
```

리소스를 만듭니다.

```
-g
```

리소스가 추가되는 리소스 그룹을 지정합니다.

```
-t SUNW.HAStoragePlus
```

리소스 유형을 SUNW.HAStoragePlus로 지정합니다.

```
-p FileSystemMountPoints=/global/mountpoint
```

파일 시스템의 마운트 지점을 전역으로 지정합니다.

-p AffinityOn=True

응용 프로그램 리소스가 -p FileSystemMountPoints에 정의된 전역 장치 및 클러스터 파일 시스템에 대한 유사성 스위치오버를 수행하도록 지정합니다. 따라서 응용 프로그램 리소스 그룹이 페일오버되거나 스위치오버되면 관련 장치 그룹도 스위치오버됩니다.

nfs-dg-rs

NFS 응용 프로그램의 HASToragePlus 리소스 이름입니다.

이러한 확장 등록 정보에 대한 자세한 내용은 [SUNW.HASToragePlus\(5\)](#) 매뉴얼 페이지를 참조하십시오.

6. 응용 프로그램 리소스 그룹에 논리 호스트 이름 리소스를 추가합니다.

```
nodeA# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-prim
```

기본 클러스터에 있는 응용 프로그램 리소스 그룹의 논리 호스트 이름은 lhost-nfsrg-prim입니다.

7. dfstab.resource-name 구성 파일을 만들어 포함하는 리소스 그룹의 Pathprefix 디렉토리 아래에 있는 SUNW.nfs 하위 디렉토리에 배치합니다.

a. nodeA에 SUNW.nfs라는 디렉토리를 만듭니다.

```
nodeA# mkdir -p /global/etc/SUNW.nfs
```

b. nodeA에 dfstab.resource-name 파일을 만듭니다.

```
nodeA# touch /global/etc/SUNW.nfs/dfstab.nfs-rs
```

c. nodeA의 /global/etc/SUNW.nfs/dfstab.nfs-rs 파일에 다음 행을 추가합니다.

```
share -F nfs -o rw -d "HA NFS" /global/mountpoint
```

8. 응용 프로그램 리소스 그룹을 온라인으로 전환합니다.

```
nodeA# clresourcegroup online -M -n nodeA nfs-rg
```

online

리소스 그룹을 온라인 상태로 전환합니다.

-e

연결된 리소스를 활성화합니다.

-M

리소스 그룹을 관리합니다.

```
-n
```

리소스 그룹을 온라인으로 전환할 노드를 지정합니다.

```
nfs-rg
```

리소스 그룹의 이름입니다.

9. 응용 프로그램 리소스 그룹이 온라인 상태인지 확인합니다.

```
nodeA# clresourcegroup status
```

리소스 그룹 상태 필드를 검사하여 응용 프로그램 리소스 그룹이 nodeA 및 nodeB에서 온라인 상태인지 확인합니다.

다음 순서 [보조 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법 \[308\]](#)으로 이동합니다.

▼ 보조 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법

- 시작하기 전에
- [기본 클러스터에서 NFS 응용 프로그램 리소스 그룹을 만드는 방법 \[305\]](#) 절차를 완료합니다.
 - /etc/netmasks 파일에 모든 논리적 호스트 이름에 대한 IP 주소 서브넷과 넷마스크 항목이 있는지 확인합니다. 필요하다면 /etc/netmasks 파일을 편집하여 누락된 항목을 추가합니다.

1. `solaris.cluster.modify`, `solaris.cluster.admin`, `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 nodeC에 액세스합니다.

2. `SUNW.nfs`를 리소스 유형으로 등록합니다.

```
nodeC# clresourcetype register SUNW.nfs
```

3. `SUNW.HAStoragePlus`가 리소스 유형으로 등록되지 않은 경우 등록합니다.

```
nodeC# clresourcetype register SUNW.HAStoragePlus
```

4. 장치 그룹의 응용 프로그램 리소스 그룹을 만듭니다.

```
nodeC# clresourcegroup create \  
-p Pathprefix=/global/etc \  
-p Auto_start_on_new_cluster=False \  
-p RG_affinities=+++devgrp-stor-rg \  
nfs-rg
```

create

리소스 그룹을 만듭니다.

-p

리소스 그룹의 등록 정보를 지정합니다.

Pathprefix=/global/etc

그룹의 리소스에서 관리 파일을 기록할 수 있는 디렉토리를 지정합니다.

Auto_start_on_new_cluster=False

응용 프로그램 리소스 그룹이 자동으로 시작되지 않도록 지정합니다.

RG_affinities=+++devgrp-stor-rg

응용 프로그램 리소스 그룹과 함께 배치할 리소스 그룹을 지정합니다. 이 예에서 응용 프로그램 리소스 그룹은 복제 리소스 그룹 devgrp-stor-rg와 함께 배치해야 합니다.

복제 리소스 그룹이 새로운 기본 노드로 스위치오버될 경우 응용 프로그램 리소스 그룹은 자동으로 스위치오버됩니다. 그러나 응용 프로그램 리소스 그룹을 새로운 기본 노드로 스위치오버하려고 시도하면 해당 동작이 배치 요구사항을 위반하므로 스위치오버가 차단됩니다.

nfs-rg

응용 프로그램 리소스 그룹의 이름

5. 응용 프로그램 리소스 그룹에 SUNW.HAStoragePlus 리소스를 추가합니다.

```
nodeC# clresource create -g nfs-rg \
-t SUNW.HAStoragePlus \
-p FileSystemMountPoints=/global/mountpoint \
-p AffinityOn=True \
nfs-dg-rs
```

create

리소스를 만듭니다.

-g

리소스가 추가되는 리소스 그룹을 지정합니다.

-t SUNW.HAStoragePlus

리소스 유형을 SUNW.HAStoragePlus로 지정합니다.

-p

리소스의 등록 정보를 지정합니다.

```
FileSystemMountPoints=/global/mountpoint
```

파일 시스템의 마운트 지점을 전역으로 지정합니다.

```
AffinityOn=True
```

응용 프로그램 리소스가 -p FileSystemMountPoints에 정의된 전역 장치 및 클러스터 파일 시스템에 대한 유사성 스위치오버를 수행하도록 지정합니다. 따라서 응용 프로그램 리소스 그룹이 페일오버되거나 스위치오버되면 관련 장치 그룹도 스위치오버됩니다.

```
nfs-dg-rs
```

NFS 응용 프로그램의 HAStoragePlus 리소스 이름입니다.

6. 응용 프로그램 리소스 그룹에 논리 호스트 이름 리소스를 추가합니다.

```
nodeC# clreslogicalhostname create -g nfs-rg \
lhost-nfsrg-sec
```

보조 클러스터에 있는 응용 프로그램 리소스 그룹의 논리 호스트 이름은 lhost-nfsrg-sec입니다.

7. 응용 프로그램 리소스 그룹에 NFS 리소스를 추가합니다.

```
nodeC# clresource create -g nfs-rg \
-t SUNW.nfs -p Resource_dependencies=nfs-dg-rs nfs-rg
```

8. 전역 볼륨이 기본 클러스터에 마운트되는 경우 보조 클러스터에서 전역 볼륨을 마운트 해제합니다.

```
nodeC# umount /global/mountpoint
```

보조 클러스터에 볼륨이 마운트되는 경우 동기화는 실패합니다.

다음 순서 [“데이터 복제 활성화 방법의 예” \[310\]](#)로 이동합니다.

데이터 복제 활성화 방법의 예

이 절에서는 구성 예에서 데이터 복제가 활성화되는 방법에 대해 설명합니다. 이 절에서는 Availability Suite 소프트웨어 명령 sndradm 및 iiadm을 사용합니다. 이러한 명령에 대한 자세한 내용은 Availability Suite 설명서를 참조하십시오.

이 절에서는 다음 절차에 대해 설명합니다.

- [기본 클러스터에서 복제를 활성화하는 방법 \[311\]](#)
- [보조 클러스터에서 복제를 활성화하는 방법 \[312\]](#)

▼ 기본 클러스터에서 복제를 활성화하는 방법

1. `solaris.cluster.read` RBAC 권한 부여를 제공하는 역할로 `nodeA`에 액세스합니다.
2. 모든 트랜잭션을 비웁니다.

```
nodeA# lockfs -a -f
```

3. 논리 호스트 이름 `lhost-reprg-prim` 및 `lhost-reprg-sec`가 온라인 상태인지 확인합니다.

```
nodeA# clresourcegroup status
nodeC# clresourcegroup status
```

리소스 그룹의 상태 필드를 검사합니다.

4. 기본 클러스터에서 보조 클러스터로의 원격 미러 복제를 활성화합니다.

이 단계에서는 기본 클러스터에서 보조 클러스터로의 복제를 활성화합니다. 이 단계에서는 기본 클러스터 마스터 볼륨(`d200`)에서 보조 클러스터 마스터 볼륨(`d200`)으로의 복제를 활성화합니다. 또한 이 단계에서는 `d203`의 원격 미러 비트맵에 대한 복제를 활성화합니다.

- 기본 클러스터와 보조 클러스터가 비동기화되는 경우 Availability Suite 소프트웨어에 다음 명령을 실행합니다.

```
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
```

- 기본 클러스터와 보조 클러스터가 동기화되는 경우 Availability Suite 소프트웨어에 다음 명령을 실행합니다.

```
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. 자동 동기화를 사용으로 설정합니다.

Availability Suite 소프트웨어에 다음 명령을 실행합니다.

```
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
```

이 단계에서는 자동 동기화를 활성화합니다. 자동 동기화의 활성 상태가 `on`으로 설정된 경우 시스템이 재부트되거나 장애가 발생하면 볼륨 세트가 재동기화됩니다.

6. 클러스터가 로깅 모드인지 확인합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeA# /usr/sbin/sndradm -P
```

출력 내용이 다음과 같이 표시됩니다.

```
/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
```

로깅 모드에서 상태는 logging이고 자동 동기화의 활성화 상태는 off입니다. 디스크의 데이터 볼륨에 기록될 때 동일한 디스크의 비트맵 파일이 업데이트됩니다.

7. 포인트 인 타임 스냅샷을 사용으로 설정합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeA# /usr/sbin/iidm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeA# /usr/sbin/iidm -w \
/dev/md/nfsset/rdisk/d201
```

이 단계에서는 기본 클러스터 마스터 볼륨을 동일한 클러스터의 새도우 볼륨으로 복사할 수 있습니다. 마스터 볼륨, 새도우 볼륨 및 포인트 인 타임 비트맵 볼륨은 동일한 장치 그룹에 있어야 합니다. 이 예에서 마스터 볼륨은 d200이고 새도우 볼륨은 d201이며 포인트 인 타임 비트맵 볼륨은 d203입니다.

8. 포인트 인 타임 스냅샷을 원격 미러 세트에 추가합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeA# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

이 단계에서는 포인트 인 타임 스냅샷을 원격 미러 볼륨 세트와 연결합니다. Availability Suite 소프트웨어를 사용하면 원격 미러 복제가 발생하기 전에 포인트 인 타임 스냅샷이 수행됩니다.

다음 순서 [보조 클러스터에서 복제를 활성화하는 방법 \[312\]](#)으로 이동합니다.

▼ 보조 클러스터에서 복제를 활성화하는 방법

시작하기 전에 [기본 클러스터에서 복제를 활성화하는 방법 \[311\]](#) 절차를 완료합니다.

1. nodeC에 root 역할로 액세스합니다.

2. 모든 트랜잭션을 비웁니다.

```
nodeC# lockfs -a -f
```

3. 기본 클러스터에서 보조 클러스터로의 원격 미러 복제를 활성화합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
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
```

기본 클러스터는 보조 클러스터를 감지하고 동기화를 시작합니다. 클러스터의 상태에 대한 자세한 내용은 Availability Suite에 대한 시스템 로그 파일 /var/adm을 참조하십시오.

4. 독립 포인트 인 타임 스냅샷을 활성화합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeC# /usr/sbin/iadm -e ind \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
nodeC# /usr/sbin/iadm -w \
/dev/md/nfsset/rdisk/d201
```

5. 포인트 인 타임 스냅샷을 원격 미러 세트에 추가합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeC# /usr/sbin/sndradm -I a \
/dev/md/nfsset/rdisk/d200 \
/dev/md/nfsset/rdisk/d201 \
/dev/md/nfsset/rdisk/d202
```

다음 순서 [“데이터 복제 수행 방법의 예” \[313\]](#)로 이동합니다.

데이터 복제 수행 방법의 예

이 절에서는 구성 예에서 데이터 복제가 수행되는 방법에 대해 설명합니다. 이 절에서는 Availability Suite 소프트웨어 명령 sndradm 및 iadm을 사용합니다. 이러한 명령에 대한 자세한 내용은 Availability Suite 설명서를 참조하십시오.

이 절에서는 다음 절차에 대해 설명합니다.

■ 원격 미러 복제를 수행하는 방법 [314]

- [포인트 인 타임 스냅샷을 수행하는 방법 \[315\]](#)
- [복제가 올바르게 구성되었는지 확인하는 방법 \[316\]](#)

▼ 원격 미러 복제를 수행하는 방법

이 절차에서는 기본 디스크의 마스터 볼륨이 보조 디스크의 마스터 볼륨으로 복제됩니다. 마스터 볼륨은 d200이고 원격 미러 비트맵 볼륨은 d203입니다.

1. **nodeA에 root 역할로 액세스합니다.**
2. **클러스터가 로깅 모드인지 확인합니다.**
Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeA# /usr/sbin/sndradm -P
```

출력 내용이 다음과 같이 표시됩니다.

```
/dev/md/nfsset/rdsk/d200 ->
lhost-reprg-sec:/dev/md/nfsset/rdsk/d200
autosync: off, max q writes:4194304, max q fbas:16384, mode:sync,ctag:
devgrp, state: logging
```

로깅 모드에서 상태는 logging이고 자동 동기화의 활성 상태는 off입니다. 디스크의 데이터 볼륨에 기록될 때 동일한 디스크의 비트맵 파일이 업데이트됩니다.

3. **모든 트랜잭션을 비웁니다.**

nodeA# **lockfs -a -f**
4. **nodeC에서 1단계부터 3단계까지 반복합니다.**
5. **nodeA의 마스터 볼륨을 nodeC의 마스터 볼륨으로 복사합니다.**
Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeA# /usr/sbin/sndradm -n -m 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
```

6. **복제가 완료되고 볼륨이 동기화될 때까지 기다립니다.**
Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeA# /usr/sbin/sndradm -n -w lhost-reprg-prim \
/dev/md/nfsset/rdsk/d200 \
/dev/md/nfsset/rdsk/d203 lhost-reprg-sec \
```

```
/dev/md/nfsset/rdisk/d200 \  
/dev/md/nfsset/rdisk/d203 ip sync
```

7. 클러스터가 복제 모드에 있는지 확인합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeA# /usr/sbin/sndradm -P
```

출력 내용이 다음과 같이 표시됩니다.

```
/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
```

복제 모드에서 상태는 replicating,이고 자동 동기화의 활성 상태는 on입니다. 기본 볼륨을 쓸 경우 Availability Suite 소프트웨어에서 보조 볼륨이 업데이트됩니다.

다음 순서 [포인트 인 타임 스냅샷을 수행하는 방법 \[315\]](#)으로 이동합니다.

▼ 포인트 인 타임 스냅샷을 수행하는 방법

이 절차에서는 기본 클러스터의 새도우 볼륨을 기본 클러스터의 마스터 볼륨으로 동기화하기 위해 포인트 인 타임 스냅샷이 사용됩니다. 마스터 볼륨은 d200이고 비트맵 볼륨은 d203이며 새도우 볼륨은 d201입니다.

시작하기 전에 [원격 미리 복제를 수행하는 방법 \[314\]](#) 절차를 완료합니다.

1. `solaris.cluster.modify` 및 `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 `nodeA`에 액세스합니다.
2. `nodeA`에서 실행 중인 리소스를 사용 안함으로 설정합니다.

```
nodeA# clresource disable nfs-rs
```

3. 기본 클러스터를 로깅 모드로 변경합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
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
```

디스크의 데이터 볼륨에 기록될 때 동일한 디스크의 비트맵 파일이 업데이트됩니다. 복제가 발생하지 않습니다.

4. 기본 클러스터의 새도우 볼륨을 기본 클러스터의 마스터 볼륨과 동기화합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeA# /usr/sbin/iidm -u s /dev/md/nfsset/rdisk/d201
nodeA# /usr/sbin/iidm -w /dev/md/nfsset/rdisk/d201
```

5. 보조 클러스터의 새도우 볼륨을 보조 클러스터의 마스터 볼륨과 동기화합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
nodeC# /usr/sbin/iidm -u s /dev/md/nfsset/rdisk/d201
nodeC# /usr/sbin/iidm -w /dev/md/nfsset/rdisk/d201
```

6. nodeA에서 응용 프로그램을 다시 시작합니다.

```
nodeA# clresource enable nfs-rs
```

7. 보조 볼륨을 기본 볼륨과 재동기화합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
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
```

다음 순서 [복제가 올바르게 구성되었는지 확인하는 방법 \[316\]](#)으로 이동합니다.

▼ 복제가 올바르게 구성되었는지 확인하는 방법

시작하기 전에 [포인트 인 타임 스냅샷을 수행하는 방법 \[315\]](#) 절차를 완료합니다.

1. `solaris.cluster.admin` RBAC 권한 부여를 제공하는 역할로 nodeA 및 nodeC에 액세스합니다.

2. 기본 클러스터가 복제 모드에 있고 자동 동기화가 켜져 있는지 확인합니다.

Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
nodeA# /usr/sbin/sndradm -P
```

출력 내용이 다음과 같이 표시됩니다.

```
/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
```

복제 모드에서 상태는 `replicating`,이고 자동 동기화의 활성 상태는 `on`입니다. 기본 볼륨을 쓸 경우 Availability Suite 소프트웨어에서 보조 볼륨이 업데이트됩니다.

3. 기본 클러스터가 복제 모드에 있지 않으면 복제 모드로 변경합니다.
Availability Suite 소프트웨어에 다음 명령을 사용합니다.

```
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. 클라이언트 시스템에 디렉토리를 만듭니다.

- a. 클라이언트 시스템에 `root` 역할로 로그인합니다.
다음과 같은 메시지가 표시됩니다.

```
client-machine#
```

- b. 클라이언트 시스템에 디렉토리를 만듭니다.

```
client-machine# mkdir /dir
```

5. 응용 프로그램 디렉토리에 기본 볼륨을 마운트하고 마운트된 디렉토리를 표시합니다.

- a. 응용 프로그램 디렉토리에 기본 볼륨을 마운트합니다.

```
client-machine# mount -o rw lhost-nfsrg-prim:/global/mountpoint /dir
```

- b. 마운트된 디렉토리를 표시합니다.

```
client-machine# ls /dir
```

6. 응용 프로그램 디렉토리에서 기본 볼륨을 마운트 해제합니다.

- a. 응용 프로그램 디렉토리에서 기본 볼륨을 마운트 해제합니다.

```
client-machine# umount /dir
```

- b. 기본 클러스터에서 응용 프로그램 리소스 그룹을 오프라인화합니다.

```
nodeA# clresource disable -g nfs-rg +
nodeA# clresourcegroup offline nfs-rg
```

- c. 기본 클러스터를 로깅 모드로 변경합니다.
Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
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
```

디스크의 데이터 볼륨에 기록될 때 동일한 디스크의 비트맵 파일이 업데이트됩니다. 복제가 발생하지 않습니다.

- d. **PathPrefix** 디렉토리를 사용할 수 있는지 확인합니다.

```
nodeC# mount | grep /global/etc
```

- e. 파일 시스템이 보조 클러스터에 마운트되기에 적합한지 확인합니다.

```
nodeC# fsck -y /dev/md/nfsset/rdisk/d200
```

- f. 응용 프로그램을 관리되는 상태로 가져와서 보조 클러스터에서 온라인으로 전환합니다.

```
nodeC# clresourcegroup online -eM nodeC nfs-rg
```

- g. 클라이언트 시스템에 **root** 역할로 액세스합니다.

다음과 같은 메시지가 표시됩니다.

```
client-machine#
```

- h. **4단계**에서 만든 응용 프로그램 디렉토리를 보조 볼륨의 응용 프로그램 디렉토리에 마운트합니다.

```
client-machine# mount -o rw lhost-nfsrg-sec:/global/mountpoint /dir
```

- i. 마운트된 디렉토리를 표시합니다.

```
client-machine# ls /dir
```

7. **5단계**에 표시된 디렉토리가 **6단계**에 표시된 디렉토리와 동일한지 확인합니다.

8. 기본 볼륨의 응용 프로그램을 마운트된 응용 프로그램 디렉토리로 반환합니다.

- a. 보조 볼륨에서 응용 프로그램 리소스 그룹을 오프라인화합니다.

```
nodeC# clresource disable -g nfs-rg +  
nodeC# clresourcegroup offline nfs-rg
```

- b. 전역 볼륨이 보조 볼륨에서 마운트 해제되도록 합니다.

```
nodeC# umount /global/mountpoint
```

- c. 응용 프로그램 리소스 그룹을 관리되는 상태로 가져와서 기본 클러스터에서 온라인으로 전환합니다.

```
nodeA# clresourcegroup online -eM nodeA nfs-rg
```

d. 기본 볼륨을 복제 모드로 변경합니다.

Availability Suite 소프트웨어에 대해 다음 명령을 실행합니다.

```
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
```

기본 볼륨을 쓸 경우 Availability Suite 소프트웨어에서 보조 볼륨이 업데이트됩니다.

참조 “테이크오버를 관리하는 방법의 예” [319]

테이크오버를 관리하는 방법의 예

이 절에서는 DNS 항목을 업데이트하는 방법에 대해 설명합니다. 추가 정보는 “테이크오버 관리 지침” [292]을 참조하십시오.

▼ DNS 항목 업데이트 방법

DNS에서 클라이언트를 클러스터에 매핑하는 방법은 [그림 7](#)을 참조하십시오.

1. **nsupdate** 명령을 시작합니다.

자세한 내용은 [nsupdate\(1M\)](#) 매뉴얼 페이지를 참조하십시오.

2. 두 클러스터 모두에서 응용 프로그램 리소스 그룹의 논리적 호스트 이름과 클러스터 IP 주소 간의 현재 DNS 매핑을 제거합니다.

```
> 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

기본 클러스터의 IP 주소(역방향)

ipaddress2rev

보조 클러스터의 IP 주소(역방향)

ttl

지속 시간(초) 기본값은 3600입니다.

3. 두 클러스터 모두에서 응용 프로그램 리소스 그룹의 논리 호스트 이름과 클러스터 IP 주소 간에 새로운 DNS 매핑을 만듭니다.

기본 논리 호스트 이름을 보조 클러스터의 IP 주소로 매핑하고 보조 논리 호스트 이름을 기본 클러스터의 IP 주소로 매핑합니다.

```
> 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

보조 클러스터의 IP 주소(전방향)

ipaddress1fwd

기본 클러스터의 IP 주소(전방향)

색인

번호와 기호

- /etc/vfstab 파일, 55
 - 구성 확인, 145
 - 마운트 지점 추가, 144
- /kernel/drv/
 - md.conf 파일, 120
- /var/adm/messages 파일, 92
- AffinityOn 등록 정보
 - 데이터 복제를 위한 확장 등록 정보, 289
- Availability Suite
 - 데이터 복제에 사용, 285
- boot 명령, 69
- BUI 살펴볼 내용 Oracle Solaris Cluster Manager
- cconsole 명령 살펴볼 내용 pconsole 유틸리티
- claccess 명령, 27
- cldevice 명령, 27
- cldevicegroup 명령, 27
- clinterconnect 명령, 27
- clnasdevice 명령, 27
- clnode 명령, 240, 241
- clquorum 명령, 27
- clreslogicalhostname 명령, 27
- clresource 명령, 27
 - 리소스 그룹 삭제, 248
 - 리소스 및 리소스 그룹 삭제, 248
- clresourcegroup 명령, 27, 241
- clresourcetype 명령, 27
- clressharedaddress 명령, 27
- clsetup, 26, 27, 32
 - 영역 클러스터 관리, 243
 - 영역 클러스터 만들기, 25
 - 영역 클러스터에 네트워크 주소 추가, 246
 - 장치 그룹 관리, 113
 - 전송 스위치 관리, 181
 - 쿼럼 장치 관리, 157
- clsetup 유틸리티
 - 영역 클러스터 만들기, 23
- clsnmphot 명령, 27
- clsnmpmib 명령, 27
- clsnmpuser 명령, 27
- cltelemattribute 명령, 27
- cluster check
 - 명령, 50
- cluster check 명령, 27
 - vfstab 파일 검사, 145
- cluster shutdown 명령, 65
- clzonecluster 명령, 27
 - solaris10 브랜드 영역 클러스터 패치 적용, 267
- 부트, 69
- 설명, 32
- 정지, 65
- consoles
 - 연결 대상, 31
- CPU
 - 구성, 260
- CPU 공유
 - 구성, 259
 - 전역 클러스터 노드, 260
 - 제어, 259
- DID 정보
 - 수동 업데이트, 152
- DID 정보 수동 업데이트, 152
- DLPI, 185
- DNS(Domain Name System)
 - 데이터 복제에서 업데이트, 319
 - 업데이트 지침, 293
- EMC SRDF
 - DID 장치 구성, 104
 - Domino 모드, 96
 - 관리, 101
 - 구성 예, 106
 - 구성 확인, 105

- 복제 그룹 구성, 102
- 요구사항, 97
- 제한 사항, 97
- 최고 사례, 98
- 캠퍼스 클러스터의 기본 공간이 전체 실패한 후 복구, 111
- failback 등록 정보, 130
- fence_level 살펴볼 내용 복제 도중
- GlobalDevicePaths
 - 데이터 복제를 위한 확장 등록 정보, 289
- GUI 살펴볼 내용 Oracle Solaris Cluster Manager IPMP
 - 관리, 195
 - 상태 확인, 40
- IPMP 그룹에서 네트워크 인터페이스 연결 해제, 196
- iSCSI 저장소
 - 쿼럼 장치로 사용
 - 링크 기반 IPMP 사용, 196
 - 프로브 기반 IPMP 사용, 196
- labeled 브랜드 영역 클러스터, 23
- lofi 파일
 - 제거, 233
- md.tab 파일, 29
- metaset 명령, 99
- MIB
 - SNMP 이벤트 사용 및 사용 안함, 235, 236
 - SNMP 이벤트 프로토콜 변경, 236
- network-bind-address
 - 확인, 280
- NFS(네트워크 파일 시스템)
 - 기본 클러스터의 리소스, 307
 - 데이터 복제를 위해 응용 프로그램 파일 시스템 구성, 300
- ntp.conf.sc 파일, 224
- numsecondaries 등록 정보, 132
- OBP(OpenBoot PROM), 221
- Oracle GlassFish
 - 업데이트 제한 사항, 277
- Oracle Solaris Cluster 쿼럼 서버
 - 쿼럼 장치로 지원, 159
- Oracle Solaris Cluster Manager, 26, 277
 - 로그온 방법, 277
 - 문제 해결, 277
 - 사용, 277
 - 수행 가능한 작업
- Oracle ZFS Storage Appliance NAS 장치 추가, 162
- 개인 어댑터 제거, 186
- 개인 어댑터 추가, 184
- 공유 저장소 추가, 216
- 노드 등록 정보 편집, 154
- 노드 비우기, 65, 227
- 노드 시스템 메시지 보기, 31, 32
- 노드를 유지 보수 상태로 설정, 227
- 노드의 상태 보기, 40
- 디스크 경로의 모니터링 사용 안함으로 설정, 151
- 디스크 경로의 모니터링 사용으로 설정, 149
- 리소스 및 리소스 그룹 보기, 35
- 영역 구성 보기, 25
- 영역 클러스터 노드 부트, 83
- 영역 클러스터 노드 재부트, 87
- 영역 클러스터 노드 종료, 80, 207
- 영역 클러스터 노드에서 로드 한계 만들기, 241
- 영역 클러스터 노드에서 소프트웨어 제거, 25
- 영역 클러스터 만들기, 25, 243
- 영역 클러스터 삭제, 247
- 영역 클러스터 종료, 25
- 영역 클러스터에 공유 저장소 추가, 200
- 영역 클러스터에 네트워크 주소 추가, 246
- 영역 클러스터에 저장 장치 추가, 243
- 영역 클러스터에 파일 시스템 추가, 243
- 영역 클러스터에서 저장 장치 제거, 251
- 영역 클러스터에서 파일 시스템 제거, 249
- 영역 클러스터의 리소스 보안 등록 정보 편집, 243
- 장치 그룹을 오프라인으로 전환, 113, 137
- 장치 그룹을 온라인으로 전환, 113, 135
- 전송 어댑터 제거, 186
- 전송 어댑터 추가, 184
- 전역 클러스터 노드 종료, 80
- 전역 클러스터 노드에서 로드 한계 만들기, 241
- 케이블 사용 안함으로 설정, 190
- 케이블 사용으로 설정, 189
- 케이블 제거, 186
- 케이블 추가, 184
- 쿼럼 서버 만들기, 163
- 쿼럼 장치 만들기, 160
- 쿼럼 장치 사용 안함으로 설정, 170
- 쿼럼 장치 사용으로 설정, 171
- 쿼럼 장치 재설정, 170

- 궤럼 장치 제거, 166
 - 궤럼 정보 보기, 173
 - 클러스터 구성 보기, 41
 - 클러스터 구성 유틸리티 액세스, 32
 - 클러스터 구성요소 상태 확인, 37
 - 클러스터 상호 연결 사용으로 설정, 194
 - 클러스터 상호 연결의 상태 확인, 183
 - 파일 시스템 제거, 146
 - 파일 시스템 추가, 143
 - 액세스, 278
 - 업데이트 제한 사항, 277
 - 접근성 컨트롤 설정, 281
 - 토폴로지, 277
 - 토폴로지 뷰, 281
 - Oracle Solaris OS
 - CPU 제어, 259
 - svcadm 명령, 222
 - 영역 클러스터 정의, 23
 - 저장소 기반 복제, 94
 - 전역 클러스터 정의, 23
 - 전역 클러스터에 대한 관리 작업, 24
 - 특별 지침
 - 노드 부트, 83
 - 노드 재부트, 87
 - 호스트 기반 복제, 94
 - Oracle ZFS Storage Appliance
 - 궤럼 장치로 지원, 159
 - 궤럼 장치로 추가, 161
 - patchadd
 - solaris10 브랜드 영역 클러스터 패치 적용, 268
 - pconsole 유틸리티, 31
 - 보안 연결, 31
 - 사용, 204
 - RBAC, 59
 - 권한 프로파일(설명), 59
 - 작업
 - 사용, 59
 - 사용자 수정, 62
 - 사용자정의 역할 추가, 61
 - 설정, 59
 - 역할 추가, 61
 - SATA 저장소, 160
 - 궤럼 장치로 지원, 159
 - scinstall 명령
 - 클러스터 노드 복원, 202
 - showrev -p 명령, 33
 - SNMP
 - 사용자 제거, 240
 - 사용자 추가, 239
 - 이벤트 MIB 사용 및 사용 안함, 235, 236
 - 프로토콜 변경, 236
 - 호스트 사용, 237
 - 호스트 사용 안함, 238
 - SNMP 이벤트 MIB 사용 및 사용 안함, 235, 236
 - solaris 브랜드 영역 클러스터, 23
 - 업데이트, 267
 - Solaris OS 살펴볼 내용 Oracle Solaris OS
 - Solaris Volume Manager
 - 원시 디스크 장치 이름, 144
 - solaris10 브랜드 영역 클러스터, 23
 - clzonecluster로 패치 적용, 267
 - patchadd로 패치 적용, 268
 - 대체 BE 패치 적용, 268
 - SRDF 살펴볼 내용 EMC SRDF
 - ssh, 31
 - Unified Archive
 - 클러스터 노드 복원, 202
 - vfstab 파일
 - 구성 확인, 145
 - 마운트 지점 추가, 144
 - ZFS
 - 구성
 - HAStoragePlus 없이 로컬 저장소 풀, 124
 - 복제된 장치 그룹, 123
 - 루트 파일 시스템 제한, 113
 - 복제, 123
 - 추가
 - HAStoragePlus 없이 저장소 풀, 124
 - 복제된 장치 그룹, 123
 - 파일 시스템 삭제, 248
 - ZFS Storage Appliance 살펴볼 내용 궤럼 장치
 - zpool
 - HAStoragePlus 없이 로컬 구성, 123, 124
 - ZPoolsSearchDir
 - 데이터 복제를 위한 확장 등록 정보, 289
- ㄱ
- 개요
- 궤럼, 157
- 개인 호스트 이름
- 변경, 222

- 검사
 - 전역 마운트 지점, 55
- 검증
 - 영역 클러스터 구성, 50
 - 전역 클러스터 구성, 50
- 게스트 도메인, 83
- 공용 네트워크
 - 관리, 181, 195
 - 동적 재구성, 197
- 공유 디스크 경로
 - 모니터링, 148
 - 자동 재부트 사용 안함, 155
 - 자동 재부트 사용으로 설정, 154
- 공유 SCSI 디스크
 - 쿼럼 장치로 지원, 159
- 공통 에이전트 컨테이너
 - 보안 키 구성, 280
- 관리
 - EMC SRDF 복제된 장치, 101
 - IPMP, 195
 - PNM, 181
 - 영역 클러스터, 24, 243
 - 저장소 기반의 복제된 장치, 101
 - 전역 클러스터, 24
 - 전역 클러스터 설정, 215
 - 쿼럼, 157
 - 클러스터 상호 연결 및 공용 네트워크, 181
 - 클러스터 파일 시스템, 113
- 관리 도구
 - Oracle Solaris Cluster Manager, 277
- 관리 콘솔, 29
- 구성
 - HAStoragePlus 없이 로컬 ZFS 저장소 풀, 124
 - 노드에서 로드 한계, 241
 - 데이터 복제, 285
 - 보안 키, 280
 - 복제된 ZFS 장치 그룹, 123
 - 저장소 기반의 복제된 장치, 101
- 구성 예(캠퍼스 클러스터)
 - 2공간, 저장소 기반 데이터 복제, 95
 - 2공간, 저장소 기반 복제, 95
- 권한 프로파일
 - RBAC, 59
- 권한, 전역 장치, 100
- 꼭 찬 /var/adm/messages 파일 복구, 92
- 나열
 - 장치 그룹 구성, 134
 - 쿼럼 구성, 173
 - 네트워크 주소
 - 영역 클러스터에 추가, 246
 - 노드
 - ID 찾기, 217
 - 기본, 100, 130
 - 로드 한계 구성, 241
 - 보조, 130
 - 부트, 79
 - 연결 대상, 31
 - 영역 클러스터에서 이름 바꾸기, 225
 - 유지 보수 상태로 전환, 227
 - 인증, 218
 - 전역 클러스터에서 이름 바꾸기, 225
 - 제거
 - 영역 클러스터에서, 207
 - 오류 메시지, 213
 - 장치 그룹에서, 126
 - 전역 클러스터에서, 208
 - 종료, 79
 - 추가, 199
 - 노드 이름 바꾸기
 - 영역 클러스터, 225
 - 전역 클러스터, 225
 - 논리 호스트 이름 리소스
 - 데이터 복제 테이크오버에서의 역할, 289
 - 다시 시작
 - 영역 클러스터 노드, 87
 - 전역 클러스터 노드, 87
 - 대체 BE
 - solaris10 브랜드 영역 클러스터 패치 적용, 268
 - 데이터 복제, 93
 - DNS 항목 업데이트, 319
 - 구성
 - NFS 응용 프로그램 리소스 그룹, 305
 - NFS 응용 프로그램의 파일 시스템, 300
 - 유사성 스위치오버, 289, 302
 - 장치 그룹, 297
 - 구성 예, 293
 - 구성 확인, 316

- 동기식, 287
 - 리소스 그룹
 - 공유 주소, 291
 - 구성, 289
 - 만들기, 302
 - 응용 프로그램, 289
 - 이름 지정 규칙, 289
 - 페일오버 응용 프로그램, 290
 - 확장 가능 응용 프로그램, 291
 - 비동기식, 287
 - 사용으로 설정, 310
 - 소개, 285
 - 예, 313
 - 원격 미리, 286, 314
 - 저장소 기반, 94, 95
 - 정의, 93
 - 지침
 - 리소스 그룹 구성, 288
 - 스위치오버 관리, 292
 - 테이크오버 관리, 292
 - 테이크오버 관리, 319
 - 포인트 인 타임 스냅샷, 287, 315
 - 필요한 하드웨어 및 소프트웨어, 295
 - 호스트 기반, 93, 285
 - 데이터 복제를 위한 공유 주소 리소스 그룹, 291
 - 데이터 복제를 위한 스위치오버
 - 수행, 319
 - 유사성 스위치오버, 289
 - 데이터 복제를 위한 페일오버 응용 프로그램
 - AffinityOn 등록 정보, 289
 - GlobalDevicePaths, 289
 - ZPoolsSearchDir, 289
 - 관리, 319
 - 지침
 - 리소스 그룹, 290
 - 테이크오버 관리, 292
 - 데이터 복제를 위한 확장 가능 응용 프로그램, 291
 - 데이터 복제를 위한 확장 등록 정보
 - 복제 리소스, 304
 - 응용 프로그램 리소스, 306, 309
 - 동기식 데이터 복제, 96, 287
 - 동적 재구성, 100, 100
 - 공용 네트워크 인터페이스, 197
 - 쿼럼 장치, 158
 - 클러스터 상호 연결, 182
 - 등록 정보
 - failback, 130
 - numsecondaries, 132
 - preferenced, 130
 - 등록 해제
 - Solaris Volume Manager 장치 그룹, 126
 - 디스크 경로
 - 모니터링, 99, 148, 149
 - 오류 디스크 경로 인쇄, 151
 - 모니터링 해제, 150
 - 상태 오류 해결, 152
- ㄹ**
- 로그인
 - 원격, 31
 - 로드 한계
 - concentrate_load 등록 정보, 240
 - preemption_mode 등록 정보, 240
 - 노드에서 구성, 240, 241
 - 로컬 미러링 살펴볼 내용 저장소 기반 복제
 - 루프백 마운트
 - 영역 클러스터로 파일 시스템 내보내기, 248
 - 리소스
 - 구성 정보 표시, 35
 - 삭제, 248
 - 리소스 그룹
 - 데이터 복제
 - 구성, 289
 - 구성 지침, 288
 - 페일오버에서 역할, 289
 - 릴리스 정보, 33
- ㄴ**
- 마운트 지점
 - /etc/vfstab 파일 수정, 144
 - 전역, 55
 - 마이그레이션
 - 전역 장치 이름 공간, 117
 - 마지막 쿼럼 장치
 - 제거, 166
 - 매니페스트
 - 자동 설치 프로그램, 204
 - 명령, 27
 - boot, 69
 - cconsole, 31

- claccess, 27
 - cldevice, 27
 - cldevicegroup, 27
 - clinterconnect, 27
 - clnasdevice, 27
 - clquorum, 27
 - clreslogicalhostname, 27
 - clresource, 27
 - clresourcegroup, 27
 - clresourcetype, 27
 - clressharedaddress, 27
 - clsetup, 27
 - clsnmp host, 27
 - clsnmpmib, 27
 - clsnmpuser, 27
 - cltelemetryattribute, 27
 - cluster check, 27, 30, 50, 55
 - cluster shutdown, 65
 - clzonecluster, 27, 65
 - clzonecluster boot, 69
 - clzonecluster verify, 50
 - metaset, 99
 - 명령줄 관리 도구, 26
 - 모니터링
 - 공유 디스크 경로, 154
 - 디스크 경로, 149
 - 모니터링 해제
 - 디스크 경로, 150
 - 문제 해결
 - network-bind-address, 280
 - Oracle Solaris Cluster Manager, 279
 - 전송 케이블, 어댑터, 스위치, 194
 - 미러, 온라인 백업, 271
- ㅂ**
- 백업
 - 미러 온라인, 271
 - 클러스터, 29, 271
 - 변경
 - numsecondaries 등록 정보, 132
 - SNMP 이벤트 MIB 프로토콜, 236
 - 개인 호스트 이름, 222
 - 기본 노드, 135
 - 등록 정보, 130
 - 클러스터 이름, 216
- 보안 키
 - 구성, 280
 - 보조
 - 기본 번호, 130
 - 원하는 수 설정, 132
 - 복구
 - 저장소 기반 데이터 복제를 사용한 클러스터, 98
 - 쿼럼 장치, 174
 - 복원
 - 루트 파일 시스템, 274
 - 클러스터 노드
 - scinstall 사용, 202
 - Unified Archive에서, 202
 - 클러스터 파일, 273
 - 복제 살펴볼 내용 데이터 복제
 - 복제, 저장소 기반, 95
 - 볼륨 살펴볼 내용 저장소 기반 복제
 - 부트
 - 비클러스터 모드, 90
 - 영역 클러스터, 65
 - 영역 클러스터 노드, 79
 - 전역 클러스터, 65
 - 전역 클러스터 노드, 79
 - 비동기식 데이터 복제, 96, 287
 - 비클러스터 모드 부트, 90
 - 비트맵
 - 원격 미러 복제, 286
 - 포인트 인 타임 스냅샷, 287
- ㅅ**
- 사용
 - 역할(RBAC), 59
 - 사용자
 - SNMP 제거, 240
 - SNMP 추가, 239
 - 등록 정보 수정, 62
 - 상태
 - 영역 클러스터 구성요소, 37
 - 전역 클러스터 구성요소, 37
 - 상호 연결
 - 문제 해결, 194
 - 사용으로 설정, 194
 - 서비스 중단
 - 쿼럼 장치, 170
 - 설정

- 역할(RBAC), 59
- 소프트웨어 업데이트
 - 개요, 263
- 속성 살펴볼 내용 등록 정보
- 수정
 - 사용자(RBAC), 62
 - 쿼럼 장치 노드 목록, 169
- 스냅샷 살펴볼 내용 저장소 기반 복제
 - 포인트 인 타임, 287
- 스위치, 전송, 186
- 스위치백
 - 데이터 복제 수행 지침, 293
- 시간 초과
 - 쿼럼 장치 기본값 변경, 174
- 시작
 - pconsole 유틸리티, 204
 - 영역 클러스터, 69
 - 영역 클러스터 노드, 79
 - 전역 클러스터, 69
 - 전역 클러스터 노드, 79
-
- 액세스
 - Oracle Solaris Cluster Manager, 278
- 어댑터, 전송, 186
- 업그레이드
 - solaris 브랜드 유형의 페일오버 영역, 263
 - solaris10 브랜드 유형의 페일오버 영역, 263
 - 개요, 263
- 업데이트, 267
 - 살펴볼 다른 내용 패치 적용
 - AI 서버, 269
 - Oracle GlassFish에 대한 제한 사항, 277
 - Oracle Solaris Cluster Manager에 대한 제한 사항, 277
 - solaris 브랜드 영역 클러스터, 267
 - 개요, 263
 - 쿼럼 서버, 269
 - 특정 패키지, 266
 - 팁, 265
- 역할
 - 사용자정의 역할 추가, 61
 - 설정, 59
 - 역할 추가, 61
- 역할 기반 액세스 제어 살펴볼 내용 RBAC

- 영역 경로
 - 이동, 243
- 영역 클러스터
 - solaris 브랜드
 - 업데이트, 267
 - solaris10 브랜드
 - clzonecluster로 패치 적용, 267
 - patchadd로 패치 적용, 268
 - 대체 BE 패치 적용, 268
- 관리, 215
- 구성 검증, 50
- 구성 정보 보기, 48
- 구성요소 상태, 37
- 네트워크 주소 추가, 246
- 만들기, 24
- 복제, 243
- 부트, 65
- 브랜드, 23
- 영역 경로 이동, 243
- 응용 프로그램에 대해 준비, 243
- 재부트, 73
- 정의, 23
- 종료, 65
- 지원되는 직접 마운트, 248
- 통합 아카이브에서 구성, 244
- 통합 아카이브에서 설치, 245
- 파일 시스템 제거, 243
- 영역 클러스터 노드
 - IP 주소 및 NIC 지정, 199
 - 부트, 79
 - 재부트, 87
 - 종료, 79
- 예제
 - 기능 유효성 검사 실행, 53
 - 대화식 유효성 검사 나열, 53
- 오류 메시지
 - /var/adm/messages 파일, 92
- 노드 제거, 213
- 원격 로그인, 31
- 원격 미러 복제
 - 수행, 314
 - 정의, 286
- 원격 미러링 살펴볼 내용 저장소 기반 복제
- 원격 복제 살펴볼 내용 저장소 기반 복제
- 원시 디스크 장치
 - 이름 지정 규약, 144

- 원시 디스크 장치 그룹
 - 추가, 122
 - 유사성 스위치오버
 - 데이터 복제를 위한 구성, 302
 - 유지 보수
 - 웨어 장치, 170
 - 유지 보수 상태
 - 노드, 227
 - 웨어 장치 상태 해제, 171
 - 웨어 장치 전환, 170
 - 응용 프로그램 리소스 그룹
 - 데이터 복제를 위한 구성, 305
 - 지침, 289
 - 이름 공간
 - 마이그레이션, 117
 - 전역, 99
 - 이름 지정 규약
 - 원시 디스크 장치, 144
 - 이름 지정 규칙
 - 복제 리소스 그룹, 289
 - 이벤트 MIB
 - log_number 변경, 236
 - min_severity 변경, 236
 - SNMP 사용 및 사용 안함, 235, 236
 - SNMP 프로토콜 변경, 236
 - 인쇄
 - 오류 디스크 경로, 151
- ㅉ**
- 자동 설치 프로그램
 - 매니페스트, 204
 - 장치
 - 전역, 99
 - 장치 그룹
 - Solaris Volume Manager
 - 추가, 120
 - 관리 개요, 113
 - 구성 나열, 134
 - 기본 소유권, 130
 - 데이터 복제를 위한 구성, 297
 - 등록 정보 변경, 130
 - 원시 디스크
 - 추가, 122
 - 유지 보수 상태, 136
 - 제거
 - 및 등록 해제, 126
 - 추가, 122
 - 장치 그룹에 대한 기본 노드 전환, 135
 - 장치 그룹의 기본 소유권, 130
 - 재부트
 - 영역 클러스터, 73
 - 영역 클러스터 노드, 87
 - 전역 클러스터, 73
 - 전역 클러스터 노드, 87
 - 재해 허용 한계
 - 정의, 285
 - 저장소 기반 데이터 복제, 95
 - 및 웨어 장치, 98
 - 복구, 98
 - 요구사항, 97
 - 정의, 94
 - 제한 사항, 97
 - 최고 사례, 98
 - 저장소 기반의 복제된 장치
 - 관리, 101
 - 구성, 101
 - 저장소 어레이
 - 제거, 211
 - 전송 스위치, 추가, 37, 41, 181, 186
 - 전송 어댑터, 추가, 37, 41, 181, 186
 - 전송 케이블
 - 사용 안함으로 설정, 189
 - 사용으로 설정, 188
 - 추가, 37, 41, 181, 186
 - 전송 케이블 사용 안함으로 설정, 189
 - 전송 케이블을 사용으로 설정, 188
 - 전역
 - 마운트 지점, 검사, 55
 - 마운트 지점, 확인, 148
 - 이름 공간, 99, 115
 - 장치, 99
 - 권한 설정, 100
 - 동적 재구성, 100
 - 전역 이름 공간 업데이트, 115
 - 전역 장치 이름 공간
 - 마이그레이션, 117
 - 전역 클러스터
 - 관리, 215
 - 구성 검증, 50
 - 구성 정보 보기, 41
 - 구성요소 상태, 37

- 노드 제거, 208
 - 부트, 65
 - 재부트, 73
 - 정의, 23
 - 종료, 65
 - 전역 클러스터 노드
 - CPU 공유, 260
 - 부트, 79
 - 재부트, 87
 - 종료, 79
 - 확인
 - 상태, 206
 - 전원 관리, 215
 - 전환
 - 장치 그룹에 대한 기본 노드, 135
 - 접근성
 - Oracle Solaris Cluster Manager, 281
 - 제거
 - lofi 장치 파일, 233
 - Oracle Solaris Cluster 소프트웨어, 231
 - SNMP
 - 사용자, 240
 - 호스트, 238
 - Solaris Volume Manager 장치 그룹, 126
 - 노드, 206, 208
 - 마지막 쿼럼 장치, 166
 - 모든 장치 그룹의 노드, 126
 - 영역 클러스터에서, 207
 - 영역 클러스터에서 리소스 및 리소스 그룹, 248
 - 저장소 어레이, 211
 - 전송 케이블, 어댑터, 스위치, 186
 - 쿼럼 장치, 159, 165
 - 클러스터 파일 시스템, 146
 - 패키지, 269
 - 종료
 - 노드, 79
 - 영역 클러스터, 65
 - 영역 클러스터 노드, 79
 - 전역 클러스터, 65
 - 전역 클러스터 노드, 79
 - 중지
 - 영역 클러스터, 73
 - 영역 클러스터 노드, 79
 - 전역 클러스터, 73
 - 전역 클러스터 노드, 79
 - 지원되는 쿼럼 장치 유형, 159
 - 직접 마운트
 - 영역 클러스터로 파일 시스템 내보내기, 248
 - 직접 연결된 공유 디스크 쿼럼 장치
 - 추가, 160
- ㄷ
- 찾기
 - 영역 클러스터의 노드 ID, 217
 - 전역 클러스터의 노드 ID, 217
 - 최고 사례, 98
 - EMC SRDF, 98
 - 저장소 기반 데이터 복제, 98
 - 추가
 - SNMP
 - 사용자, 239
 - 호스트, 237
 - Solaris Volume Manager 장치 그룹, 122
 - ZFS
 - HAStoragePlus 없이 저장소 풀, 124
 - 복제된 장치 그룹, 123
 - 노드, 199
 - 영역 클러스터에, 200
 - 전역 클러스터에, 200
 - 사용자정의 역할(RBAC), 61
 - 역할(RBAC), 61
 - 영역 클러스터에 네트워크 주소, 246
 - 장치 그룹, 120, 122
 - 전송 케이블, 어댑터, 스위치, 37, 41, 181
 - 직접 연결된 공유 디스크 쿼럼 장치, 160
 - 쿼럼 장치, 160
 - Oracle ZFS Storage Appliance NAS, 161
 - 쿼럼 서버, 163
 - 클러스터 파일 시스템, 143
- ㄴ
- 캠퍼스 클러스터
 - 저장소 기반 데이터 복제, 95
 - 저장소 기반 데이터 복제를 사용한 복구, 98
 - 케이블, 전송, 186
 - 쿼럼
 - 개요, 157
 - 관리, 157
 - 쿼럼 서버 살펴볼 내용 쿼럼 장치
 - 쿼럼 서버 쿼럼 장치

- 설치 요구사항, 163
- 제거 문제 해결, 166
- 추가, 163
- 쿼럼 장치
 - 교체, 168
 - 구성 나열, 173
 - 기본 시간 초과 변경, 174
 - 노드 목록 수정, 169
 - 및 저장소 기반 복제, 98
 - 복구, 174
 - 유지 보수 상태, 장치 상태 해제, 171
 - 유지 보수 상태, 장치 전환, 170
 - 장치의 동적 재구성, 158
 - 제거, 159, 165
 - 마지막 쿼럼 장치, 166
 - 지원되는 유형, 159
 - 추가, 160
 - Oracle ZFS Storage Appliance NAS, 161
 - 직접 연결된 공유 디스크 쿼럼 장치, 160
 - 쿼럼 서버, 163
- 쿼럼 장치 교체, 168
- 클러스터
 - 노드 인증, 218
 - 백업, 29, 271
 - 범위, 32
 - 시간 설정, 220
 - 이름 변경, 216
 - 파일 복원, 273
- 클러스터 모니터
 - Oracle Solaris Cluster Manager 토폴로지 사용, 282
- 클러스터 상호 연결
 - 관리, 181
 - 동적 재구성, 182
- 클러스터 시간 설정, 220
- 클러스터 콘솔에 보안 연결, 31
- 클러스터 파일 시스템, 99
 - 관리, 113
 - 구성 확인, 145
 - 마운트 옵션, 145
 - 제거, 146
 - 추가, 143
- 클러스터 파일 시스템의 마운트 옵션
 - 요구사항, 145

E

- 토폴로지
 - Oracle Solaris Cluster Manager에서 사용, 282
- 통합 아카이브
 - 영역 클러스터 구성, 244
 - 영역 클러스터 설치, 245

F

- 파일
 - /etc/vfstab, 55
 - md.conf, 120
 - md.tab, 29
 - ntp.conf.sc, 224
- 파일 시스템
 - NFS 응용 프로그램
 - 데이터 복제를 위해 구성, 300
 - 루트 복원
 - 설명, 274
 - 영역 클러스터에서 제거, 243
- 패치 적용, 267
 - 살펴볼 다른 내용 업데이트
 - solaris10 브랜드 영역 클러스터
 - clzonecluster 사용, 267
 - patchadd 사용, 268
 - 대체 BE, 268
- 패키지
 - 제거, 269
 - 특정 패키지 업데이트, 266
- 페어 웨어 스케줄러
 - CPU 공유 구성, 260
- 포인트 인 타임 스냅샷
 - 수행, 315
 - 정의, 287
- 프로파일
 - RBAC 권한, 59

H

- 호스트
 - SNMP 추가 및 제거, 237, 238
- 호스트 기반 데이터 복제
 - 예, 285
 - 정의, 93
- 확인

vfstab 구성, 145
데이터 복제 구성, 316
전역 마운트 지점, 148
클러스터 노드 상태, 206

