



Verwenden Sie „System Controller replace“-Befehle, um ein Upgrade der Controller-Hardware mit ONTAP 9.5 auf 9.7 durchzuführen

Upgrade controllers

NetApp
February 19, 2026

Inhalt

Verwenden Sie „System Controller replace“-Befehle, um ein Upgrade der Controller-Hardware mit ONTAP 9.5 auf 9.7 durchzuführen	1
Erfahren Sie mehr über dieses ARL-Upgradeverfahren	1
Automatisierung des Controller-Upgrades	2
Entscheiden Sie, ob Sie dieses Verfahren zur Aggregatverlagerung verwenden möchten	2
Unterstützte System-Upgrade-Kombinationen	3
Wählen Sie ein anderes Verfahren zum Hardware-Upgrade	4
Die erforderlichen Tools und Dokumentationen	4
Richtlinien für das Controller-Upgrade mit ARL	4
Unterstützte Upgrades für ARL	4
2-Node-Cluster ohne Switches	5
Upgrades werden für ARL nicht unterstützt	5
MetroCluster FC-Konfiguration	5
Fehlerbehebung	5
Überprüfen Sie den Systemzustand der MetroCluster-Konfiguration	6
Prüfen Sie auf MetroCluster-Konfigurationsfehler	7
Überprüfung von UmschalttaFunktionen, Healing und Switchback	7
Erfahren Sie mehr über die ARL-Upgradesequenz	7
Aktualisieren Sie das Node-Paar	7
Übersicht über die ARL-Upgrade-Sequenz	7
Stufe 1: Upgrade vorbereiten	9
Bereiten Sie die Knoten für ein Upgrade vor	9
Management der Storage-Verschlüsselung mit dem Onboard Key Manager	14
Stufe 2: Knoten1 verschieben und ausmustern	14
Verschieben von Aggregaten ohne Root-Wurzeln und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, auf Knoten 2	14
Verschiebung ausgefallener oder Vetos von Aggregaten	17
Node1 ausmustern	18
Vorbereitungen für den Netzboot	18
Phase 3: Installieren und booten Sie node3	19
Installieren und booten Sie node3	19
Legen Sie die FC- oder UTA/UTA2-Konfiguration auf node3 fest	25
Weisen Sie node1-Festplatten Knoten 3 neu zu	29
Ports von node1 nach node3 zuordnen	35
Fügen Sie dem Quorum bei, wenn ein Node über einen anderen Satz an Netzwerkports verfügt	40
Überprüfen Sie die Installation von node3	42
Verschieben Sie Aggregate ohne Root-Root-Fehler und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, von node2 auf node3	43
Phase 4: Knoten2 verschieben und ausmustern	47
Verschieben von Aggregaten und NAS-Daten-LIFs ohne Root-Wurzeln von Knoten 2 auf Knoten 3	47
Node2 ausmustern	49
Phase 5: installieren und booten sie node4	49
installieren und booten sie node4	49

Legen Sie die FC- oder UTA/UTA2-Konfiguration auf node4 fest	55
Weisen Sie Node2-Festplatten node4 neu zu	60
Weisen Sie Ports von node2 nach node4 zu	65
Fügen Sie dem Quorum bei, wenn ein Node über einen anderen Satz an Netzwerkports verfügt	69
Überprüfen Sie die installation von node4	71
Verschieben Sie Aggregate ohne Root-Root-Fehler und NAS-Daten-LIFs, die sich im Besitz von node2 befinden, von node3 auf node4	73
Phase 6: Schließen Sie das Upgrade ab	76
Authentifizierungsmanagement mit KMIP-Servern	76
Vergewissern Sie sich, dass die neuen Controller ordnungsgemäß eingerichtet sind	77
Richten Sie Storage Encryption auf dem neuen Controller-Modul ein.	79
Einrichtung von NetApp Volume oder Aggregate Encryption auf dem neuen Controller-Modul	80
Ausmustern des alten Systems	83
Setzen Sie den SnapMirror Betrieb fort	83
Fehlerbehebung	83
Fehler bei der Aggregatverschiebung	83
Neustarts, Panikspiele oder Energiezyklen	85
Probleme, die in mehreren Phasen des Verfahrens auftreten können	89
Fehler bei der LIF-Migration	90
Quellen	90
Referenzinhalt	90
Referenzstandorte	92

Verwenden Sie „System Controller replace“-Befehle, um ein Upgrade der Controller-Hardware mit ONTAP 9.5 auf 9.7 durchzuführen

Erfahren Sie mehr über dieses ARL-Upgradeverfahren

Es gibt verschiedene Methoden zur Aggregate Relocation (ARL) für die Aktualisierung der Controller-Hardware. Dieses Verfahren beschreibt die Aktualisierung der Controller-Hardware mithilfe von ARL und „System Controller Replace Commands“ auf Systemen mit ONTAP 9.7, 9.6 oder 9.5.

Während des Verfahrens führen Sie ein Upgrade der ursprünglichen Controller Hardware mit der Ersatz-Controller-Hardware durch. Hierbei werden die Eigentumsrechte an Aggregaten verschoben, die nicht mit Root-Berechtigungen verbunden sind. Sie migrieren Aggregate mehrmals von Node zu Node, um zu bestätigen, dass mindestens ein Node während des Upgrades Daten von den Aggregaten bereitstellt. Außerdem migrieren Sie Daten-logische Schnittstellen (LIFs) und weisen Sie die Netzwerk-Ports auf dem neuen Controller den Schnittstellengruppen zu, während Sie fortfahren.

In diesen Informationen verwendete Terminologie

In dieser Information werden die ursprünglichen Knoten „node1“ und „node2“ genannt und die neuen Knoten „node3“ und „node4“ genannt. Während des beschriebenen Verfahrens wird "node1" durch "node3" ersetzt und "node2" durch "node4" ersetzt.

Die Begriffe "node1", "node2", "node3" und "node4" werden nur verwendet, um zwischen den ursprünglichen und den neuen Knoten zu unterscheiden. Wenn Sie das Verfahren befolgen, müssen Sie die richtigen Namen Ihrer ursprünglichen und neuen Knoten ersetzen. In der Realität ändern sich die Namen der Knoten jedoch nicht: „node3“ hat den gleichen Namen wie „node1“ und „node4“ hat nach dem Upgrade der Controller-Hardware den gleichen Namen wie „node2“.

Wichtige Informationen:

- Diese Vorgehensweise ist komplex und setzt voraus, dass Sie über erweiterte ONTAP-Administrationsfähigkeiten verfügen. Sie müssen außerdem die ["Richtlinien zum Upgraden von Controllern mit ARL"](#) und die ["ARL-Upgradesequenz"](#) bevor Sie mit dem Upgrade beginnen.
- Bei dieser Vorgehensweise wird vorausgesetzt, dass die Ersatz-Controller-Hardware neu ist und nicht verwendet wurde. Die erforderlichen Schritte zur Vorbereitung gebrauchter Controller mit dem `wipeconfig` Befehl ist in dieser Prozedur nicht enthalten. Wenn bereits die Ersatz-Controller-Hardware verwendet wurde, müssen Sie sich an den technischen Support wenden, insbesondere wenn auf den Controllern Data ONTAP in 7-Mode ausgeführt wurde.
- Sie können dieses Verfahren zum Upgrade der Controller Hardware in Clustern mit mehr als zwei Nodes verwenden. Sie müssen jedoch die Verfahren für jedes HA-Paar im Cluster separat durchführen.
- Dieses Verfahren gilt für FAS -Systeme und AFF -Systeme.
- Ab ONTAP 9.6 gilt dieses Verfahren für Systeme mit MetroCluster-Konfiguration mit 4 Nodes oder höher. Da sich die MetroCluster-Konfigurationsstandorte an zwei physisch unterschiedlichen Standorten befinden können, muss das automatisierte Controller-Upgrade für ein HA-Paar individuell an jedem MetroCluster Standort durchgeführt werden.
- Wenn Sie ein Upgrade von einem AFF A320 System durchführen, können Sie das Upgrade der Controller-Hardware durch Volume-Verschiebung durchführen oder den technischen Support kontaktieren. Wenn Sie bereit sind, die Lautstärke zu verschieben, lesen Sie ["Quellen"](#) Link zu *Upgrade durch Verschieben von*

Automatisierung des Controller-Upgrades

Während eines Controller-Upgrades wird der Controller durch einen anderen Controller ersetzt, auf dem eine neuere oder leistungsstärkere Plattform läuft.

In früheren Versionen dieses Inhalts enthielten Anweisungen für einen unterbrechungsfreien Controller-Update, der vollständig manuell ausgeführt wurde. Dieser Inhalt enthält die Schritte für das neue automatisierte Verfahren.

Der manuelle Prozess war langwierig und komplex, aber in diesem vereinfachten Verfahren können Sie ein Controller-Update mithilfe von Aggregatverschiebung implementieren, sodass effizientere, unterbrechungsfreie Upgrades für HA-Paare möglich sind. Vor allem in Bezug auf Validierung, Sammlung von Informationen und Nachprüfungen sind deutlich weniger manuelle Schritte erforderlich.

Entscheiden Sie, ob Sie dieses Verfahren zur Aggregatverlagerung verwenden möchten

Es gibt verschiedene Methoden zur Aggregate Relocation (ARL) für die Aktualisierung der Controller-Hardware. Dieses Verfahren beschreibt die Aktualisierung der Controller-Hardware mittels ARL und „System Controller Replace Commands“ auf Systemen mit ONTAP 9.7, 9.6 oder 9.5. Sie sollten dieses komplexe Verfahren nur verwenden, wenn Sie ein erfahrener ONTAP-Administrator sind.

Um zu entscheiden, ob dieses ARL-Verfahren für Ihr Controller-Hardware-Upgrade geeignet ist, sollten Sie alle folgenden Umstände für unterstützte Upgrades überprüfen:

- Sie aktualisieren gerade NetApp Controller mit ONTAP 9.5, 9.6 oder 9.7. Dieses Dokument gilt nicht für ein Upgrade auf ONTAP 9.8.
- Sie möchten die neuen Controller nicht als neues HA-Paar zum Cluster hinzufügen und die Daten mithilfe von Volume-Verschiebungen migrieren.
- Sie sind in der Verwaltung von ONTAP erfahren und sind mit den Risiken der Arbeit im Diagnose-Privilege-Modus vertraut.
- Ihre Hardware-Upgrade-Kombination finden Sie in der [Unterstützte Modellmatrix](#).
- Wenn Sie eine MetroCluster Konfiguration aktualisieren, handelt es sich um eine FC-Konfiguration mit 4 oder mehr Nodes und auf allen Nodes wird ONTAP 9.6 oder 9.7 ausgeführt.



- Wenn Sie ein System aktualisieren, indem Sie Controller-Module im selben Gehäuse austauschen, wie z. B. AFF A800 oder AFF C800, empfiehlt NetApp dringend, das Upgrade-Verfahren zu verwenden, das ["Upgrades von Controller-Modellen mit ARL, wobei das vorhandene Systemgehäuse und die Festplatten erhalten bleiben"](#). Dieses ARL-Verfahren umfasst die Schritte, die sicherstellen, dass die internen Festplatten sicher im Gehäuse bleiben, wenn Sie die Controller während des Upgrade-Vorgangs entfernen und installieren.

["Erfahren Sie mehr über die unterstützten System-Upgrade-Kombinationen mit ARL, wobei das vorhandene Systemgehäuse und die Festplatten erhalten bleiben"](#).

- Dabei können Sie NetApp Storage Encryption (NSE), NetApp Volume Encryption (NVE) und NetApp Aggregate Encryption (NAE) verwenden.

Unterstützte System-Upgrade-Kombinationen

Die folgende Tabelle zeigt die unterstützte Modellmatrix für das Controller-Upgrade mit diesem ARL-Verfahren.

Alter Controller	Ersatz-Controller
FAS8020, FAS8040, FAS8060, FAS8080	FAS8200, FAS8300, FAS8700, FAS9000
AFF8020, AFF8040, AFF8060, AFF8080	AFF A300, AFF A400, AFF A700 ¹ , AFF A800 ²
FAS8200	FAS8700, FAS9000, FAS8300 ^{4, 5}
AFF A300	AFF A700 ¹ , AFF A800 ^{2, 3} , AFF A400 ^{4, 5}



Wenn die Kombination aus dem Controller-Upgrade-Modell nicht in der oben stehenden Tabelle aufgeführt ist, wenden Sie sich an den technischen Support.

Das automatisierte Upgrade ¹ARL für das AFF A700 System wird von ONTAP 9.7P2 unterstützt.

²Wenn Sie auf ein AFF A800 oder ein System aktualisieren, das interne und externe Festplatten unterstützt, müssen Sie die spezifischen Anweisungen für das Root-Aggregat auf internen NVMe-Festplatten befolgen. Siehe ["Weisen Sie node1-Festplatten Knoten 3, Schritt 9, neu zu"](#) Und ["Weisen Sie Node2-Festplatten node4, Schritt 9, neu zu"](#) Die

Das automatisierte Upgrade ³ARL von einer AFF A300 auf ein AFF A800 System wird von ONTAP 9.7P5 unterstützt.

Das automatisierte Upgrade ⁴ARL von einer AFF A300 auf eine AFF A400 und eine FAS8200 zu einem FAS8300 System wird von ONTAP 9.7P8 unterstützt.

⁵Wenn Sie in einer 2-Node-Cluster-Konfiguration ein Upgrade von einer AFF A300 auf eine AFF A400 oder ein FAS8200 auf ein FAS8300 System durchführen, müssen Sie für das Controller-Upgrade temporäre Cluster-Ports auswählen. Die AFF A400- und FAS8300-Systeme sind in zwei Konfigurationen erhältlich – als Ethernet-Bundle, bei dem die Ports der Mezzanine-Karte Ethernet-Typ und als FC-Bundle enthalten sind. Dort befinden sich die Mezzanine-Ports vom FC-Typ.

- Bei einer AFF A400 oder einer FAS8300 mit Ethernet-Konfiguration können Sie jeden der beiden Mezzanine-Ports als temporäre Cluster-Ports verwenden.
- Bei einer AFF A400 oder einem FAS8300 mit FC-Typ-Konfiguration müssen Sie eine 10-GbE-Netzwerkschnittstellenkarte mit vier Ports (Teilenummer X1147A) hinzufügen, um temporäre Cluster Ports bereitstellen zu können.

- Nach Abschluss eines Controller-Upgrades mithilfe von temporären Cluster-Ports können Sie Cluster-LIFs unterbrechungsfrei zu e3a und e3b migrieren, 100-GbE-Ports auf einem AFF A400 System sowie e0c und e0d, 100-GbE-Ports auf einem FAS8300 System.

Wählen Sie ein anderes Verfahren zum Hardware-Upgrade

- ["Überprüfen Sie die verfügbaren alternativen ARL-Methoden zum Aktualisieren der Controller-Hardware"](#).
- Wenn Sie eine andere Methode zum Upgrade der Controller-Hardware bevorzugen und bereit sind, Volume-Verschiebungen durchzuführen, lesen Sie ["Quellen"](#) Link zu *Upgrade durch Verschieben von Volumes oder Storage*.

Verwandte Informationen

Siehe ["Quellen"](#) um auf die *ONTAP 9-Dokumentation* zu verlinken.

Die erforderlichen Tools und Dokumentationen

Sie müssen über spezielle Tools verfügen, um die neue Hardware zu installieren, und Sie müssen während des Upgrade-Prozesses andere Dokumente referenzieren.

Für die Durchführung des Upgrades benötigen Sie die folgenden Tools:

- Erdungsband
- #2 Kreuzschlitzschraubendreher

Wechseln Sie zum ["Quellen"](#) Abschnitt für den Zugriff auf die Liste der für dieses Upgrade erforderlichen Referenzdokumente und Referenzsites

Richtlinien für das Controller-Upgrade mit ARL

Um zu verstehen, ob Sie bei einem Controller-Upgrade von ONTAP 9.5 auf ONTAP 9.7 mit Aggregate Relocation (ARL) arbeiten können, hängt von der Plattform und der Konfiguration der ursprünglichen Controller sowie von dem Ersatz-Controller ab.

Unterstützte Upgrades für ARL

Wenn Sie ein Node-Paar mit diesem ARL-Verfahren für ONTAP 9.5 auf ONTAP 9.7 aktualisieren, müssen Sie sicherstellen, dass ARL an den Original- und Ersatz-Controllern ausgeführt werden kann.

Sie sollten die Größe aller definierten Aggregate und die Anzahl der Festplatten überprüfen, die vom ursprünglichen System unterstützt werden. Dann müssen Sie die Aggregatgröße und Anzahl der unterstützten Festplatten mit der Aggregatgröße und der Anzahl der vom neuen System unterstützten Festplatten vergleichen. Siehe ["Quellen"](#) Link zum *Hardware Universe*, wo diese Information verfügbar ist. Die Aggregatgröße und die Anzahl der vom neuen System unterstützten Festplatten müssen gleich oder größer sein als die Aggregatgröße und Anzahl der vom ursprünglichen System unterstützten Festplatten.

Sie sollten in den Cluster-Mischregeln validieren, ob neue Nodes zusammen mit den vorhandenen Nodes in das Cluster integriert werden können, wenn der ursprüngliche Controller ersetzt wird. Weitere Informationen zu Regeln für die Kombination von Clustern finden Sie unter ["Quellen"](#) Zum Verknüpfen mit der *Hardware Universe*.



Bevor Sie ein AFF-Systemupgrade durchführen, müssen Sie ONTAP auf Version 9.5P1 oder neuer aktualisieren. Diese Versionsebenen sind für ein erfolgreiches Upgrade erforderlich.



Informationen zum Upgrade eines Systems, das interne Laufwerke unterstützt (z. B. eine FAS2700 oder AFF A250), aber KEINE internen Laufwerke enthält, finden Sie unter ["Quellen"](#). Und verwenden Sie das Verfahren in der *Aggregate Relocation*, um den für Ihre Version von ONTAP korrekten Controller-Hardware-Inhalt manuell zu aktualisieren.

Wenn Sie ONTAP 9.6P11, 9.7P8 oder neuere Versionen verwenden, wird empfohlen, die Aktivierung von Connectivity, Lebendigkeit und Availability Monitor (CLAM)-Übernahme zu aktivieren, um das Cluster bei bestimmten Node-Ausfällen in Quorum zurückzugeben. Der `kernel-service` Für Befehl ist der erweiterte Zugriff auf die Berechtigungsebene erforderlich. Weitere Informationen finden Sie unter: ["NetApp KB-Artikel SU436: DIE CLAM-Übernahme hat sich die Standardkonfiguration geändert"](#).

Das Controller-Upgrade mit ARL wird auf Systemen unterstützt, die mit SnapLock Enterprise und SnapLock Compliance Volumes konfiguriert sind.

2-Node-Cluster ohne Switches

Wenn Sie Nodes in einem 2-Node-Cluster ohne Switches aktualisieren, können Sie die Nodes im Cluster ohne Switches während des Upgrades belassen. Sie müssen sie nicht in ein Switch-Cluster konvertieren.

Upgrades werden für ARL nicht unterstützt

Sie können die folgenden Aktualisierungen nicht ausführen:

- Zum Austausch von Controllern, die die mit den ursprünglichen Controllern verbundenen Platten-Shelves nicht unterstützen

Siehe ["Quellen"](#) Um Informationen zur Hardware Universe Festplattenunterstützung zu erhalten.

- Um Controller der Einstiegsklasse mit internen Laufwerken zu erhalten, beispielsweise eine FAS 2500.

Informationen zum Upgrade von Controllern der Einstiegsklasse mit internen Laufwerken finden Sie unter ["Quellen"](#) Link zu *Upgrade durch Verschiebung von Volumes oder Storage* und Vorgang *Upgrade eines Node-Paares, auf dem Clustered Data ONTAP durch Verschieben von Volumes* ausgeführt wird.

MetroCluster FC-Konfiguration

In einer MetroCluster FC-Konfiguration müssen die Knoten für Disaster Recovery/Failover-Standort so schnell wie möglich ersetzt werden. Nichtübereinstimmungen in Controller-Modellen innerhalb einer MetroCluster-Konfiguration werden nicht unterstützt, da dies dazu führen kann, dass die Disaster Recovery-Spiegelung offline geht. Verwenden Sie die `-skip-metrocluster-check true` Befehl zum Umgehen von MetroCluster-Prüfungen, wenn Sie Knoten am zweiten Standort ersetzen.

Fehlerbehebung

Möglicherweise ist beim Upgrade des Node-Paares ein Fehler auftritt. Der Node kann abstürzen, Aggregate werden möglicherweise nicht verschoben oder LIFs werden nicht migriert. Die Ursache des Fehlers und seiner Lösung hängt davon ab, wann der Fehler während des Aktualisierungsvorgangs aufgetreten ist.

Eine Beschreibung der verschiedenen Phasen des Verfahrens finden Sie in der Tabelle im Abschnitt „Übersicht

über das ARL-Upgrade“. Informationen über mögliche Ausfälle werden in der Phase des Verfahrens aufgelistet.

Sollten beim Aktualisieren der Controller Probleme auftreten, konsultieren Sie bitte die ["Fehlerbehebung"](#) Abschnitt. Die Informationen über mögliche Fehler sind nach den einzelnen Phasen des Verfahrens aufgelistet. ["ARL-Upgradesequenz"](#) Die

Wenn Sie keine Lösung für das Problem finden, wenden Sie sich an den technischen Support.

Überprüfen Sie den Systemzustand der MetroCluster-Konfiguration

Bevor Sie ein Upgrade auf einer Fabric-MetroCluster-Konfiguration starten, müssen Sie den Zustand der MetroCluster-Konfiguration überprüfen, um den ordnungsgemäßen Betrieb sicherzustellen.

Schritte

1. Vergewissern Sie sich, dass die MetroCluster-Komponenten ordnungsgemäß sind:

```
metrocluster check run
```

```
dpgqa-mcc-funct-8040-0403_siteA::*> metrocluster check run
```

Der Vorgang wird im Hintergrund ausgeführt.

2. Nach dem `metrocluster check run` Vorgang abgeschlossen, Ergebnisse anzeigen:

```
metrocluster check show
```

Nach etwa fünf Minuten werden die folgenden Ergebnisse angezeigt:

```
metrocluster_siteA::*> metrocluster check show
Last Checked On: 4/7/2019 21:15:05
Component          Result
-----
nodes              ok
lifs               ok
config-replication ok
aggregates         warning
clusters           ok
connections        not-applicable
volumes            ok
7 entries were displayed.
```

3. Überprüfen Sie den Status des laufenden MetroCluster-Prüfvorgangs:

```
metrocluster operation history show -job-id 38
```

4. Vergewissern Sie sich, dass es keine Systemzustandsmeldungen gibt:

```
system health alert show
```

Prüfen Sie auf MetroCluster-Konfigurationsfehler

Sie können das Active IQ Config Advisor Tool auf der NetApp Support-Website verwenden, um häufige Konfigurationsfehler zu überprüfen.

Wenn Sie keine MetroCluster-Konfiguration haben, können Sie diesen Abschnitt überspringen.

Über diese Aufgabe

Active IQ Config Advisor ist ein Tool zur Konfigurationsvalidierung und Statusüberprüfung. Sie können die Lösung sowohl an sicheren Standorten als auch an nicht sicheren Standorten zur Datenerfassung und Systemanalyse einsetzen.



Der Support für Config Advisor ist begrenzt und steht nur online zur Verfügung.

1. Laden Sie die herunter ["Active IQ Config Advisor"](#) Werkzeug.
2. Führen Sie Active IQ Config Advisor aus, überprüfen Sie die Ausgabe und folgen Sie seinen Empfehlungen, um eventuelle Probleme zu beheben.

Überprüfung von UmschalttaFunktionen, Healing und Switchback

Sie sollten die Umschalttavorgänge, die Reparatur und den Wechsel der MetroCluster Konfiguration überprüfen.

Siehe ["Quellen"](#) Verbinden mit Inhalten für *MetroCluster Management and Disaster Recovery* und Verwenden der genannten Verfahren für ausgehandelte Umschaltung, Heilung und Umschalten.

Erfahren Sie mehr über die ARL-Upgradesequenz

Bevor Sie die Nodes mit ARL aktualisieren, sollten Sie unbedingt verstehen, wie das Verfahren funktioniert. In diesem Inhalt wird das Verfahren in mehrere Phasen unterteilt.

Aktualisieren Sie das Node-Paar

Zum Upgrade des Node-Paars müssen Sie die ursprünglichen Nodes vorbereiten und dann für die ursprünglichen und die neuen Nodes eine Reihe von Schritten ausführen. Anschließend können Sie die ursprünglichen Knoten außer Betrieb nehmen.

Übersicht über die ARL-Upgrade-Sequenz

Während des Verfahrens aktualisieren Sie die ursprüngliche Controller Hardware mit der Ersatz-Controller-Hardware, einem Controller gleichzeitig. Nutzen Sie die HA-Paar-Konfiguration, um das Eigentum von Aggregaten ohne Root-Berechtigungen zu verschieben. Alle Aggregate außerhalb der Root-Ebene müssen zwei Umlagerungen durchlaufen, um das endgültige Ziel zu erreichen, nämlich den korrekten aktualisierten

Node.

Jedes Aggregat hat einen Hausbesitzer und aktuellen Eigentümer. Der Hausbesitzer ist der eigentliche Eigentümer des Aggregats, und der aktuelle Eigentümer ist der temporäre Eigentümer.

Die folgende Tabelle beschreibt die grundlegenden Aufgaben, die Sie in den einzelnen Phasen ausführen, und den Zustand der Aggregateigentümer am Ende der Phase. Detaillierte Schritte sind im weiteren Verlauf des Verfahrens aufgeführt:

Stufe	Schritte
"Stufe 1: Bereiten Sie sich auf das Upgrade vor"	In Phase 1 führen Sie Vorabprüfungen durch und korrigieren, falls erforderlich, die Eigentümerschaft für die Aggregate. Sie müssen bestimmte Informationen aufzeichnen, wenn Sie Storage-Verschlüsselung mit dem Onboard Key Manager managen und die SnapMirror Beziehungen stilllegen möchten. Gesamteigentum am Ende von Phase 1: • Node1 ist der Hausbesitzer und der aktuelle Besitzer der node1 Aggregate. • Node2 ist der Hausbesitzer und der aktuelle Besitzer der node2 Aggregate.
"Stufe 2: Knoten1 verschieben und ausmustern"	Während Phase 2 werden node1-Aggregate und NAS-Daten-LIFs in Knoten 2 verschoben. Dieser Prozess ist weitgehend automatisiert; der Vorgang hält an, damit Sie seinen Status überprüfen können. Sie müssen den Vorgang manuell fortsetzen. Bei Bedarf verschieben Sie fehlerhafte oder Vetos Aggregate. Sie müssen die erforderlichen Node1-Informationen für die spätere Verwendung im Verfahren notieren und dann Node1 ausmustern. Sie können sich auch später beim Verfahren auf den Netzboot node3 und node4 vorbereiten. Gesamteigentum am Ende von Phase 2: • Node2 ist der aktuelle Besitzer von node1 Aggregaten. • Node2 ist der Hausbesitzer und der aktuelle Besitzer von node2 Aggregaten.
"Phase 3: Installieren und booten Sie node3"	In Phase 3 installieren und starten Sie Node 3, ordnen Sie die Cluster- und Node-Management-Ports von Node 1 zu Node 3 zu, weisen Sie die Node 1-Festplatten Node 3 neu zu und überprüfen Sie die Node 3-Installation. Wenn erforderlich, legen Sie die FC- oder UTA/UTA2-Konfiguration auf node3 fest und bestätigen, dass node3 Quorum beigetreten ist. Außerdem werden die LIFs für NAS-Daten-LIFs und nicht-Root-Aggregate von node2 auf node3 verschoben und Sie überprüfen, ob die SAN-LIFs auf node3 vorhanden sind. Gesamteigentum am Ende von Stufe 3: • Node3 ist der Hausbesitzer und der aktuelle Besitzer von node1 Aggregaten. • Node2 ist der Hausbesitzer und der aktuelle Besitzer von node2 Aggregaten.

Stufe	Schritte
"Phase 4: Knoten2 verschieben und ausmustern"	Während Phase 4 werden node2-Aggregate ohne Root-Root-Fehler und logische Daten-LIFs außerhalb des SAN in Knoten 3 verschoben. Sie notieren auch die notwendigen node2 Informationen und setzen dann node2 in den Ruhestand. Gesamteigentum am Ende von Stufe 4: • Node3 ist der Hausbesitzer und aktuelle Besitzer von Aggregaten, die ursprünglich zu node1 gehörten. • Node2 ist der Hausbesitzer von node2 Aggregaten. • Node3 ist der aktuelle Besitzer von node2 Aggregaten.
"Phase 5: installieren und booten sie node4"	Während Phase 5 installieren und starten Sie node4, ordnen Sie die Cluster- und Node-Management-Ports von node2 zu node4 zu, weisen Sie die Node-2-Laufwerke node4 zu und überprüfen Sie die node-4-Installation. Wenn erforderlich, legen Sie die FC- oder UTA/UTA2-Konfiguration auf node4 fest und bestätigen, dass node4 Quorum beigetreten ist. Außerdem werden Knoten2-NAS-Daten-LIFs und nicht-Root-Aggregate von node3 auf node4 verschoben und überprüft, ob die SAN-LIFs auf node4 vorhanden sind. Gesamteigentum am Ende von Stufe 5: • Node3 ist der Hausbesitzer und aktuelle Besitzer der Aggregate, die ursprünglich zu node1 gehörten. • Node4 ist der Hausbesitzer und aktuelle Besitzer von Aggregaten, die ursprünglich zu node2 gehörten.
"Phase 6: Schließen Sie das Upgrade ab"	In Phase 6 bestätigen Sie, dass die neuen Nodes ordnungsgemäß eingerichtet wurden. Bei aktivierter Verschlüsselung konfigurieren und einrichten Sie Storage Encryption bzw. NetApp Volume Encryption. Zudem sollten die alten Nodes außer Betrieb gesetzt und der SnapMirror Betrieb fortgesetzt werden.

Stufe 1: Upgrade vorbereiten

Bereiten Sie die Knoten für ein Upgrade vor

Der Prozess des Controller-Austauschs beginnt mit einer Reihe von Vorabprüfungen. Sie sammeln auch Informationen über die ursprünglichen Nodes, die Sie später verwenden können. Falls erforderlich, ermitteln Sie den Typ der verwendeten Self-Encrypting Drives.

Schritte

1. Starten Sie den Controller-Ersatzprozess, indem Sie den folgenden Befehl in die ONTAP-Befehlszeile eingeben:

```
system controller replace start -nodes node_names
```



Sie können diesen Befehl nur auf der erweiterten Berechtigungsebene ausführen:
`set -privilege advanced`

Sie sehen die folgende Ausgabe:

Warning:

1. Current ONTAP version is 9.x

Before starting controller replacement operation, ensure that the new controllers are running the version 9.x

2. Verify that NVMEM or NVRAM batteries of the new nodes are charged, and charge them if they are not. You need to physically check the new nodes to see if the NVMEM or NVRAM batteries are charged. You can check the battery status either by connecting to a serial console or using SSH, logging into the Service Processor (SP) or Baseboard Management Controller (BMC) for your system, and use the system sensors to see if the battery has a sufficient charge.

Attention: Do not try to clear the NVRAM contents. If there is a need to clear the contents of NVRAM, contact NetApp technical support.

3. If a controller was previously part of a different cluster, run wipeconfig before using it as the replacement controller.

Do you want to continue? {y|n}: y

2. Drücken Sie y, Sie sehen die folgende Ausgabe:

Controller replacement operation: Prechecks in progress.

Controller replacement operation has been paused for user intervention.

Das System führt die folgenden Vorabprüfungen durch. Notieren Sie die Ausgabe jeder Vorabprüfung zur Verwendung im weiteren Verlauf des Verfahrens:

Pre-Check	Beschreibung
Cluster-Integritätsprüfung	Überprüft alle Nodes im Cluster, um sicherzustellen, dass sie sich in einem ordnungsgemäßen Zustand befinden.
MCC Cluster Check	Überprüft, ob es sich bei dem System um eine MetroCluster-Konfiguration handelt. Der Vorgang erkennt automatisch, ob es sich um eine MetroCluster Konfiguration handelt oder nicht, und führt die spezifischen Vorabprüfungen und Verifizierungsüberprüfungen durch. Es wird nur eine MetroCluster FC-Konfiguration mit 4 Nodes unterstützt. Bei 2-Node-MetroCluster-Konfiguration und 4-Node-MetroCluster IP-Konfiguration schlägt die Prüfung fehl. Wenn die MetroCluster-Konfiguration im Umschaltzustand ist, schlägt die Prüfung fehl.
Statusprüfung Der Aggregatverschiebung	Überprüft, ob eine Aggregatverschiebung bereits erfolgt. Wenn eine weitere Aggregatverschiebung erfolgt, schlägt die Prüfung fehl.

Pre-Check	Beschreibung
Modellname Prüfen	Überprüft, ob die Controller-Modelle bei diesem Verfahren unterstützt werden. Wenn die Modelle nicht unterstützt werden, schlägt die Aufgabe fehl.
Cluster-Quorum-Prüfung	Überprüft, ob die zu ersetzenden Nodes sich in Quorum befinden. Wenn sich die Knoten nicht im Quorum befinden, schlägt die Aufgabe fehl.
Überprüfung Der Bildversion	Überprüft, ob die zu ersetzenden Nodes dieselbe Version von ONTAP ausführen. Wenn sich die ONTAP-Image-Versionen unterscheiden, schlägt die Aufgabe fehl. Die neuen Knoten müssen auf ihren ursprünglichen Knoten dieselbe Version von ONTAP 9.x installiert sein. Wenn die neuen Nodes über eine andere Version von ONTAP installiert sind, müssen Sie die neuen Controller nach der Installation als Netzboot einsetzen. Anweisungen zum Upgrade von ONTAP finden Sie unter "Quellen" Link zu <i>Upgrade ONTAP</i> .
HA-Statusüberprüfung	Überprüft, ob beide Nodes, die ersetzt werden, in einer HA-Paar-Konfiguration mit Hochverfügbarkeit vorhanden sind. Wenn das Speicher-Failover für die Controller nicht aktiviert ist, schlägt die Aufgabe fehl.
Aggregatstatus-Prüfung	Wenn die Nodes ersetzt werden, eigene Aggregate, für die sie nicht der Home-Inhaber sind, schlägt die Aufgabe fehl. Die Nodes sollten nicht im Besitz von nicht lokalen Aggregaten sein.
Überprüfung Des Festplattenstatus	Wenn zu ersetzende Knoten keine oder fehlerhafte Festplatten haben, schlägt die Aufgabe fehl. Falls Festplatten fehlen, lesen Sie "Quellen" den Link zu <i>Festplatten- und Aggregatmanagement mit der CLI, logisches Storage Management mit der CLI</i> und <i>HA-Paar-Management</i> , um den Storage für das HA-Paar zu konfigurieren.
LIF-Statusüberprüfung von Daten	Überprüft, ob für einen der zu ersetzenden Nodes keine lokalen Daten-LIFs vorhanden sind. Die Nodes sollten keine Daten-LIFs enthalten, für die sie nicht der Home-Inhaber sind. Wenn einer der Nodes nicht-lokale Daten-LIFs enthält, schlägt die Aufgabe fehl.
LIF-Status des Clusters	Überprüft, ob die Cluster-LIFs für beide Nodes aktiv sind. Wenn die Cluster-LIFs ausgefallen sind, schlägt die Aufgabe fehl.
ASUP-Statusprüfung	Wenn ASUP Benachrichtigungen nicht konfiguriert sind, schlägt die Aufgabe fehl. Sie müssen AutoSupport aktivieren, bevor Sie mit dem Austausch des Controllers beginnen.
CPU-Auslastungs-Prüfung	Überprüft, ob die CPU-Auslastung bei allen zu ersetzenden Nodes mehr als 50 % beträgt. Wenn die CPU-Nutzung über einen erheblichen Zeitraum mehr als 50 % beträgt, schlägt die Aufgabe fehl.
Aggregatrekonstruktion	Überprüft, ob bei beliebigen Datenaggregaten eine Rekonstruktion durchgeführt wird. Wenn die Aggregatrekonstruktion ausgeführt wird, schlägt die Aufgabe fehl.
Knoten Affinität Job Überprüfung	Überprüft, ob Jobs mit Knotenorientierung ausgeführt werden. Wenn Knotenaffinitätsjobs ausgeführt werden, schlägt die Prüfung fehl.

3. Wenn der Controller-Ersatzvorgang gestartet und die Vorabprüfungen abgeschlossen sind, hält der

Vorgang die Aktivierung ein, damit Sie die Ausgabeinformationen, die Sie später bei der Konfiguration von node3 benötigen könnten, sammeln können.

4. Führen Sie den folgenden Befehlssatz aus, wie durch das Verfahren zum Austausch des Controllers auf der Systemkonsole gesteuert.

Führen Sie von dem seriellen Port aus, der mit jedem Node verbunden ist, und speichern Sie die Ausgabe der folgenden Befehle einzeln:

```
° vsriver services name-service dns show
° network interface show -curr-node local -role cluster,intercluster,node-
  mgmt,clustermgmt, data
° network port show -node local -type physical
° service-processor show -node local -instance
° network fcp adapter show -node local
° network port ifgrp show -node local
° network port vlan show
° system node show -instance -node local
° run -node local sysconfig
° storage aggregate show -node local
° volume show -node local
° network interface failover-groups show
° storage array config show -switch switch_name
° system license show -owner local
° storage encryption disk show
```



Wenn NetApp Volume Encryption (NVE) oder NetApp Aggregate Encryption (NAE) mit Onboard Key Manager verwendet wird, halten Sie die Schlüsselmanager-Passphrase bereit, um die Resynchronisierung des Schlüsselmanagers später im Verfahren durchzuführen.

5. Wenn Ihr System Self-Encrypting Drives verwendet, lesen Sie den Artikel der Knowledge Base ["Wie erkennen Sie, ob ein Laufwerk FIPS-zertifiziert ist"](#) Ermitteln der Art der Self-Encrypting Drives, die auf dem HA-Paar verwendet werden, das Sie aktualisieren. ONTAP unterstützt zwei Arten von Self-Encrypting Drives:

- ° FIPS-zertifizierte NetApp Storage Encryption (NSE) SAS- oder NVMe-Laufwerke
- ° Self-Encrypting-NVMe-Laufwerke (SED) ohne FIPS



FIPS-Laufwerke können nicht mit anderen Laufwerkstypen auf demselben Node oder HA-Paar kombiniert werden.

SEDs können mit Laufwerken ohne Verschlüsselung auf demselben Node oder HA-Paar kombiniert werden.

["Weitere Informationen zu unterstützten Self-Encrypting Drives"](#).

Korrigieren Sie die Aggregateigentümer bei Ausfall einer ARL-Vorabprüfung

Wenn die aggregierte Statusprüfung fehlschlägt, müssen Sie Aggregate des Partner-Node an den Node „Home-Owner“ zurückgeben und den Vorabprüfvorgang erneut initiieren.

Schritte

1. Gibt die Aggregate zurück, die derzeit dem Partner-Node gehören, an den Home-Owner-Node:

```
storage aggregate relocation start -node source_node -destination destination_node -aggregate-list *
```

2. Überprüfen Sie, dass weder node1 noch node2 noch Eigentümer von Aggregaten ist, für die es der aktuelle Eigentümer ist (aber nicht der Hausbesitzer):

```
storage aggregate show -nodes node_name -is-home false -fields owner-name, home-name, state
```

Das folgende Beispiel zeigt die Ausgabe des Befehls, wenn ein Node sowohl der aktuelle Eigentümer als auch der Home-Inhaber von Aggregaten ist:

```
cluster::> storage aggregate show -nodes node1 -is-home true -fields
owner-name,home-name,state
aggregate    home-name  owner-name  state
-----
aggr1        node1        node1        online
aggr2        node1        node1        online
aggr3        node1        node1        online
aggr4        node1        node1        online

4 entries were displayed.
```

Nachdem Sie fertig sind

Sie müssen den Controller-Ersatzprozess neu starten:

```
system controller replace start -nodes node_names
```

Lizenz

Einige Funktionen erfordern Lizenzen, die als *Packages* ausgegeben werden, die eine oder mehrere Funktionen enthalten. Jeder Node im Cluster muss über seinen eigenen Schlüssel für jede Funktion im Cluster verfügen.

Wenn Sie keine neuen Lizenzschlüssel haben, stehen dem neuen Controller derzeit lizenzierte Funktionen im Cluster zur Verfügung. Durch die Verwendung nicht lizenzierter Funktionen auf dem Controller können Sie jedoch möglicherweise die Einhaltung Ihrer Lizenzvereinbarung verschließen. Sie sollten daher nach Abschluss des Upgrades den neuen Lizenzschlüssel oder die neuen Schlüssel für den neuen Controller installieren.

Siehe "[Quellen](#)" Link zur *NetApp-Support-Website*, auf der Sie neue 28-stellige Lizenzschlüssel für ONTAP erhalten können. Die Schlüssel sind im Abschnitt „*My Support*“ unter „*Software licenses*“ verfügbar. Wenn auf

der Website nicht die erforderlichen Lizenzschlüssel vorhanden sind, können Sie sich an Ihren NetApp Ansprechpartner wenden.

Ausführliche Informationen zur Lizenzierung finden Sie unter ["Quellen"](#) Verknüpfen mit der Referenz *Systemadministration*.

Management der Storage-Verschlüsselung mit dem Onboard Key Manager

Sie können den Onboard Key Manager (OKM) zur Verwaltung der Schlüssel verwenden. Wenn Sie das OKM eingerichtet haben, müssen Sie die Passphrase und das Sicherungsmaterial aufzeichnen, bevor Sie mit dem Upgrade beginnen.

Schritte

1. Notieren Sie die Cluster-weite Passphrase.

Dies ist die Passphrase, die eingegeben wurde, als das OKM mit der CLI oder REST-API konfiguriert oder aktualisiert wurde.

2. Sichern Sie die Key-Manager-Informationen, indem Sie den ausführen `security key-manager onboard show-backup` Befehl.

Stilllegen der SnapMirror Beziehungen (optional)

Bevor Sie mit dem Verfahren fortfahren, müssen Sie bestätigen, dass alle SnapMirror Beziehungen stillgelegt werden. Wenn eine SnapMirror Beziehung stillgelegt wird, bleibt es bei einem Neustart und einem Failover stillgelegt.

Schritte

1. Überprüfen Sie den SnapMirror Beziehungsstatus auf dem Ziel-Cluster:

```
snapmirror show
```



Wenn der Status „Übertragen“ lautet, müssen Sie diese Transfers abbrechen:
`snapmirror abort -destination-vserver vserver_name`

Der Abbruch schlägt fehl, wenn sich die SnapMirror-Beziehung nicht im Zustand „Übertragen“ befindet.

2. Alle Beziehungen zwischen dem Cluster stilllegen:

```
snapmirror quiesce -destination-vserver *
```

Stufe 2: Knoten1 verschieben und ausmustern

Verschieben von Aggregaten ohne Root-Wurzeln und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, auf Knoten 2

Bevor Sie node1 durch Node3 ersetzen können, müssen Sie die nicht-Root-Aggregate und NAS-Daten-LIFs von node1 auf node2 verschieben, bevor Sie die Ressourcen von node1 schließlich in node3 verschieben.

Bevor Sie beginnen

Der Vorgang muss bereits angehalten werden, wenn Sie mit der Aufgabe beginnen. Sie müssen den Vorgang manuell fortsetzen.

Über diese Aufgabe

Remote-LIFs verarbeiten den Datenverkehr zu SAN-LUNs während des Upgrades. Das Verschieben von SAN-LIFs ist für den Zustand des Clusters oder des Service während des Upgrades nicht erforderlich. Sie müssen überprüfen, ob die LIFs sich in einem ordnungsgemäßen Zustand befinden und sich auf den entsprechenden Ports befinden, nachdem Sie node3 in den Online-Modus versetzt haben.



Der Home-Inhaber für die Aggregate und LIFs wird nicht geändert, nur der aktuelle Besitzer wird geändert.

Schritte

1. Wiederaufnahme der Vorgänge für die Aggregatverschiebung und die LIF-Verschiebung von NAS-Daten:

```
system controller replace resume
```

Alle Aggregate ohne Root-Root-Root-Root-Daten und LIFs werden von node1 auf node2 migriert.

Der Vorgang angehalten, damit Sie überprüfen können, ob alle node1-Aggregate und LIFs für nicht-SAN-Daten in node2 migriert wurden.

2. Überprüfen Sie den Status der Aggregatverschiebung und der LIF-Verschiebung von NAS-Daten:

```
system controller replace show-details
```

3. Wenn der Vorgang noch angehalten wird, vergewissern Sie sich, dass alle nicht-Root-Aggregate online sind, damit ihr Status bei node2 lautet:

```
storage aggregate show -node <node2> -state online -root false
```

Das folgende Beispiel zeigt, dass die nicht-Root-Aggregate auf node2 online sind:

```
cluster::> storage aggregate show -node node2 -state online -root false

Aggregate  Size      Available  Used%   State  #Vols  Nodes  RAID Status
-----
aggr_1     744.9GB   744.8GB    0%      online  5      node2
raid_dp,normal
aggr_2     825.0GB   825.0GB    0%      online  1      node2
raid_dp,normal
2 entries were displayed.
```

Wenn die Aggregate offline gegangen sind oder in node2 fremd geworden sind, bringen Sie sie mit dem folgenden Befehl auf node2, einmal für jedes Aggregat online:

```
storage aggregate online -aggregate <aggregate_name>
```

4. Überprüfen Sie, ob alle Volumes auf node2 online sind, indem Sie den folgenden Befehl auf node2 verwenden und seine Ausgabe überprüfen:

```
volume show -node <node2> -state offline
```

Wenn ein Volume auf node2 offline ist, bringen Sie sie mit dem folgenden Befehl auf node2 für jedes Volume online:

```
volume online -vserver <vserver_name> -volume <volume_name>
```

Der `vserver_name` Die Verwendung mit diesem Befehl finden Sie in der Ausgabe des vorherigen `volume show` Befehl.

5. Wenn die Ports, die derzeit Daten-LIFs hosten, nicht auf der neuen Hardware vorhanden sind, entfernen Sie sie aus der Broadcast-Domäne:

```
network port broadcast-domain remove-ports
```

6. Wenn irgendwelche LIFs ausgefallen sind, setzen Sie den Administratorstatus der LIFs auf `up`. Geben Sie den folgenden Befehl ein, einmal für jede LIF:

```
network interface modify -vserver vserver_name -lif LIF_name-home-node  
nodename -status-admin up
```

7. Wenn Schnittstellengruppen oder VLANs konfiguriert sind, führen Sie die folgenden Teilschritte aus:

- Wenn Sie sie noch nicht gespeichert haben, notieren Sie die VLAN- und Schnittstellengruppen-Informationen, damit Sie die VLANs und Schnittstellengruppen auf node3 neu erstellen können, nachdem node3 gestartet wurde.
- Entfernen Sie die VLANs aus den Schnittstellengruppen:

```
network port vlan delete -node nodename -port ifgrp -vlan-id VLAN_ID
```



Befolgen Sie die Korrekturmaßnahme, um alle Fehler zu beheben, die vom Befehl `vlan delete` vorgeschlagen werden.

- Geben Sie den folgenden Befehl ein und überprüfen Sie seine Ausgabe, um zu sehen, ob Schnittstellengruppen auf dem Node konfiguriert sind:

```
network port ifgrp show -node nodename -ifgrp ifgrp_name -instance
```

Das System zeigt Schnittstellengruppeninformationen für den Node an, wie im folgenden Beispiel gezeigt:

```
cluster::> network port ifgrp show -node node1 -ifgrp a0a -instance
Node: node1
Interface Group Name: a0a
Distribution Function: ip
Create Policy: multimode_lacp
MAC Address: 02:a0:98:17:dc:d4
Port Participation: partial
Network Ports: e2c, e2d
Up Ports: e2c
Down Ports: e2d
```

- a. Wenn Schnittstellengruppen auf dem Node konfiguriert sind, notieren Sie die Namen dieser Gruppen und der ihnen zugewiesenen Ports. Löschen Sie dann die Ports, indem Sie den folgenden Befehl eingeben, und zwar einmal für jeden Port:

```
network port ifgrp remove-port -node nodename -ifgrp ifgrp_name -port
netport
```

Verschiebung ausgefallener oder Vetos von Aggregaten

Falls Aggregate nicht verschoben oder ein Veto eingesetzt werden kann, müssen sie die Aggregate manuell verschieben oder – falls erforderlich – entweder die Vetos oder Zielprüfungen überschreiben.

Über diese Aufgabe

Das System unterbricht den Umzugsvorgang aufgrund des Fehlers.

Schritte

1. Überprüfen Sie die EMS-Protokolle, um festzustellen, warum das Aggregat nicht verschoben oder ein Veto eingelegt hat.
2. Verschiebung ausgefallener oder Vetos von Aggregaten:

```
storage aggregate relocation start -node node1 -destination node2 aggregate-
list * -ndocontroller-upgrade true
```

3. Geben Sie bei der entsprechenden Aufforderung ein *y*.
4. Sie können die Verschiebung mit einer der folgenden Methoden erzwingen:

Option	Beschreibung
Veto-Prüfungen werden überschrieben	Geben Sie Folgendes ein: storage aggregate relocation start -override -vetoes * -ndocontroller-upgrade true
Zielprüfungen überschreiben	Geben Sie Folgendes ein: storage aggregate relocation start -overridedestination-checks * -ndo -controllerupgrade true

Node1 ausmustern

Um „node1“ außer Betrieb zu nehmen, setzen Sie den automatischen Vorgang fort, um das HA-Paar mit node2 zu deaktivieren und node1 ordnungsgemäß herunterzufahren. Später im Verfahren entfernen Sie Knoten 1 aus dem Rack oder Gehäuse.

Schritte

1. Vorgang fortsetzen:

```
system controller replace resume
```

2. Vergewissern Sie sich, dass node1 angehalten wurde:

```
system controller replace show-details
```

Nachdem Sie fertig sind

Sie können Node1 nach Abschluss des Upgrades außer Betrieb nehmen. Siehe ["Ausmustern des alten Systems"](#).

Vorbereitungen für den Netzboot

Nachdem Sie später noch Node3 und node4 physisch gerast haben, müssen Sie sie eventuell als Netzboot Netboot eingesetzt werden. Der Begriff „Netzboot“ bedeutet, dass Sie über ein ONTAP Image, das auf einem Remote Server gespeichert ist, booten. Bei der Vorbereitung auf den Netzboot legen Sie eine Kopie des ONTAP 9-Startabbilds auf einen Webserver, auf den das System zugreifen kann.

Bevor Sie beginnen

- Vergewissern Sie sich, dass Sie mit dem System auf einen HTTP-Server zugreifen können.
- Siehe ["Quellen"](#) Um eine Verknüpfung zur NetApp Support-Website zu erhalten und die erforderlichen Systemdateien für Ihre Plattform und die richtige Version von ONTAP herunterzuladen.

Über diese Aufgabe

Sie müssen die neuen Controller als Netzboot ansehen, wenn sie nicht die gleiche Version von ONTAP 9 auf ihnen installiert sind, die auf den ursprünglichen Controllern installiert ist. Nachdem Sie jeden neuen Controller installiert haben, starten Sie das System über das auf dem Webserver gespeicherte ONTAP 9-Image. Anschließend können Sie die richtigen Dateien auf das Boot-Medium herunterladen, um später das System zu booten.

Sie müssen die Controller jedoch nicht per Netzboot fahren, wenn auf den Original-Controllern die gleiche Version von ONTAP 9 installiert ist. Wenn ja, können Sie diesen Abschnitt überspringen und mit fortfahren ["Stufe 3 Installieren und Booten von Knoten3"](#)

Schritte

1. Rufen Sie die NetApp Support Site auf, um die Dateien zum Netzboot des Systems herunterzuladen.
2. Laden Sie die entsprechende ONTAP Software im Bereich Software Downloads auf der NetApp Support Website herunter und speichern Sie die `<ontap_version>_image.tgz` Datei in einem webbasierten Verzeichnis.
3. Wechseln Sie in das Verzeichnis für den Zugriff über das Internet, und stellen Sie sicher, dass die

benötigten Dateien verfügbar sind.

Für...	Dann...
Systeme der FAS/AFF8000 Serie	Extrahieren Sie den Inhalt des <code><ontap_version>_image.tgz</code> Datei zum Zielverzeichnis: <code>tar -zxvf <ontap_version>_image.tgz</code>  Wenn Sie die Inhalte unter Windows extrahieren, verwenden Sie 7-Zip oder WinRAR, um das Netzboot-Bild zu extrahieren. Ihre Verzeichnisliste sollte einen Netzboot-Ordner mit einer Kernel-Datei enthalten: <code>netboot/kernel</code>
Alle anderen Systeme	Ihre Verzeichnisliste sollte die folgende Datei enthalten: <code><ontap_version>_image.tgz</code>  Sie müssen den Inhalt des nicht extrahieren <code><ontap_version>_image.tgz</code> Datei:

Sie verwenden die Informationen in den Verzeichnissen in ["Phase 3"](#).

Phase 3: Installieren und booten Sie node3

Installieren und booten Sie node3

Sie müssen node3 im Rack installieren, Verbindungen von node1 zu node3, Boot node3 übertragen und ONTAP installieren. Sie müssen dann eine der freien Festplatten von node1, alle Festplatten, die zum Root-Volume gehören, und alle nicht-Root-Aggregate, die zuvor nicht in node2 verschoben wurden, wie in diesem Abschnitt beschrieben neu zuweisen.

Über diese Aufgabe

Der Umzugsvorgang wird zu Beginn dieser Phase angehalten. Dieser Prozess ist weitgehend automatisiert; der Vorgang hält an, damit Sie seinen Status überprüfen können. Sie müssen den Vorgang manuell fortsetzen. Außerdem müssen Sie überprüfen, ob die SAN-LIFs erfolgreich in Knoten 3 verschoben wurden.

Sie müssen als Netzboot node3 wechseln, wenn nicht die gleiche Version von ONTAP 9 installiert ist auf node1. Nachdem Sie node3 installiert haben, starten Sie es vom ONTAP 9-Image, das auf dem Webserver gespeichert ist. Anschließend können Sie die richtigen Dateien auf das Boot-Medium für nachfolgende Systemstarts herunterladen, indem Sie den Anweisungen in folgen ["Vorbereitungen für den Netzboot"](#).



- Bei einem AFF A800 oder AFF C800 -Controller-Upgrade müssen Sie sicherstellen, dass alle Laufwerke im Gehäuse fest an der Mittelplatine sitzen, bevor Sie Knoten 1 entfernen. Weitere Informationen finden Sie unter "[Ersetzen Sie die AFF A800- oder AFF C800-Controller-Module](#)".
- Wenn Sie ein System mit Speicherfestplatten aktualisieren, müssen Sie diesen gesamten Abschnitt abschließen und anschließend mit den fortfahren "[Konfigurieren Sie FC-Ports auf node3](#)" Und "[UTA/UTA2-Ports in node3 prüfen und konfigurieren](#)" Geben Sie Abschnitte ein, und geben Sie Befehle an der Cluster-Eingabeaufforderung ein.

Schritte

1. stellen Sie sicher, dass Sie Platz im Rack für node3 haben.

Wenn sich Node1 und Node2 in einem separaten Chassis befanden, können Sie Node3 in denselben Rack-Standort wie node1 platzieren. Wenn sich Node1 jedoch im selben Chassis mit node2 befand, müssen Sie den Node3 in seinen eigenen Regalbereich legen, vorzugsweise in der Nähe der Position von node1.

2. Installieren Sie node3 im Rack und befolgen Sie die Anweisungen *Installation und Setup* für Ihr Node-Modell.



Wenn Sie auf ein System upgraden, bei dem sich beide Knoten im selben Chassis befinden, installieren Sie Knoten 4 und Knoten 3 im Chassis. Wenn Sie nicht beide Knoten im selben Chassis installieren, verhält sich Knoten 3 beim Starten so, als ob er sich in einer Dual-Chassis-Konfiguration befände, und beim Starten von Knoten 4 wird die Verbindung zwischen den Knoten nicht hergestellt.

3. Kabelnode3, Verschieben der Verbindungen von node1 nach node3.

Verkabeln Sie die folgenden Verbindungen mithilfe der *Installations- und Einrichtungsanweisungen* für die Node3-Plattform, des entsprechenden Disk-Shelf-Dokuments und der *HA-Paarverwaltung*-Dokumentation.

Siehe "[Quellen](#)" zur Verknüpfung mit *HA-Paarverwaltung*.

- Konsole (Remote-Management-Port)
- Cluster-Ports
- Datenports
- Cluster- und Node-Management-Ports
- Storage
- SAN-Konfigurationen: iSCSI Ethernet und FC Switch Ports



Möglicherweise müssen Sie die Interconnect-Karte oder die Cluster Interconnect-Kabelverbindung von node1 zu node3 nicht verschieben, da die meisten Plattform-Modelle über ein einzigartiges Interconnect-Kartenmodell verfügen. Für die MetroCluster Konfiguration müssen Sie die FC-VI-Kabelverbindungen von node1 auf node3 verschieben. Wenn der neue Host keine FC-VI-Karte besitzt, müssen Sie möglicherweise die FC-VI-Karte verschieben.

4. Einschalten Sie den Netzstrom auf node3, und unterbrechen Sie dann den Bootvorgang, indem Sie an der Konsole Strg-C drücken, um auf die Eingabeaufforderung der Boot-Umgebung zuzugreifen.

Wenn Sie ein Upgrade auf ein System mit beiden Nodes im gleichen Chassis durchführen, wird node4 auch neu gebootet. Allerdings kann man den node4-Stiefel bis später ignorieren.



Wenn Sie node3 booten, wird möglicherweise die folgende Warnmeldung angezeigt:

WARNING: The battery is unfit to retain data during a power outage. This is likely because the battery is discharged but could be due to other temporary conditions.
When the battery is ready, the boot process will complete and services will be engaged.
To override this delay, press 'c' followed by 'Enter'

5. Wenn die Warnmeldung in angezeigt wird [Schritt 4](#), Nehmen Sie die folgenden Aktionen:
 - a. Überprüfen Sie auf Meldungen der Konsole, die auf ein anderes Problem als eine schwache NVRAM-Batterie hinweisen und ergreifen Sie gegebenenfalls erforderliche Korrekturmaßnahmen.
 - b. Warten Sie, bis der Akku geladen ist und der Bootvorgang abgeschlossen ist.



Überschreiben Sie die Verzögerung nicht. Wenn Sie den Akku nicht aufladen, kann dies zu Datenverlust führen.



Siehe "[Vorbereitungen für den Netzboot](#)".

6. Konfigurieren Sie die Netzboot-Verbindung, indem Sie eine der folgenden Aktionen auswählen.



Sie müssen den Management-Port und die IP als Netzboot-Verbindung verwenden. Verwenden Sie keine Daten-LIF-IP, oder es kann während des Upgrades ein Datenausfall auftreten.

Wenn DHCP (Dynamic Host Configuration Protocol) lautet...	Dann...
Wird Ausgeführt	Konfigurieren Sie die Verbindung automatisch, indem Sie an der Eingabeaufforderung der Boot-Umgebung den folgenden Befehl eingeben: <code>ifconfig e0M -auto</code>

Wenn DHCP (Dynamic Host Configuration Protocol) lautet...	Dann...
Nicht ausgeführt	Konfigurieren Sie die Verbindung manuell, indem Sie an der Eingabeaufforderung der Boot-Umgebung den folgenden Befehl eingeben: <pre>ifconfig e0M -addr=<i>filer_addr</i> -mask=<i>netmask</i> -gw=<i>gateway</i> -dns=<i>dns_addr</i> -domain=<i>dns_domain</i></pre> <i>filer_addr</i> Ist die IP-Adresse des Speichersystems (obligatorisch). <i>netmask</i> Ist die Netzwerkmaske des Storage-Systems (erforderlich). <i>gateway</i> Ist das Gateway für das Storage-System. (Pflichtfeld). <i>dns_addr</i> Ist die IP-Adresse eines Namensservers in Ihrem Netzwerk (optional). <i>dns_domain</i> Der Domain Name (DNS) ist der Domain-Name. Wenn Sie diesen optionalen Parameter verwenden, benötigen Sie in der Netzboot-Server-URL keinen vollqualifizierten Domännennamen. Sie benötigen nur den Host-Namen des Servers.  Andere Parameter können für Ihre Schnittstelle erforderlich sein. Eingabe <code>help ifconfig</code> Details finden Sie in der Firmware-Eingabeaufforderung.

7. Netzboot auf Node3 durchführen:

Für...	Dann...
Systeme der FAS/AFF8000 Serie	<code>netboot http://<web_server_ip/path_to_web-accessible_directory>/netboot/kernel</code>
Alle anderen Systeme	<code>netboot http://<web_server_ip/path_to_web-accessible_directory>/<ontap_version>_image.tgz</code>

Der `<path_to_the_web-accessible_directory>` Sollten Sie dazu führen, wo Sie das heruntergeladen haben `<ontap_version>_image.tgz` Im Abschnitt ["Vorbereitungen für den Netzboot"](#).



Unterbrechen Sie den Startvorgang nicht.

8. im Startmenü Option wählen (7) Install new software first.

Mit dieser Menüoption wird das neue ONTAP-Image auf das Startgerät heruntergeladen und installiert.

Ignorieren Sie die folgende Meldung:

This procedure is not supported for Non-Disruptive Upgrade on an HA pair

Der Hinweis gilt für unterbrechungsfreie Upgrades der ONTAP und keine Upgrades von Controllern.



Aktualisieren Sie den neuen Node immer als Netzboot auf das gewünschte Image. Wenn Sie eine andere Methode zur Installation des Images auf dem neuen Controller verwenden, wird möglicherweise das falsche Image installiert. Dieses Problem gilt für alle ONTAP Versionen. Das Netzboot wird mit der Option kombiniert (7) `Install new software` Entfernt das Boot-Medium und platziert dieselbe ONTAP-Version auf beiden Image-Partitionen.

9. Wenn Sie aufgefordert werden, den Vorgang fortzusetzen, geben Sie ein `y`, Und wenn Sie zur Eingabe des Pakets aufgefordert werden, geben Sie die URL ein:

```
http://<web_server_ip/path_to_web-  
accessible_directory>/<ontap_version>_image.tgz
```

10. Vervollständigen Sie die folgenden Teilschritte, um das Controller-Modul neu zu starten:
 - a. Eingabe `n` So überspringen Sie die Backup-Recovery, wenn folgende Eingabeaufforderung angezeigt wird:

```
Do you want to restore the backup configuration now? {y|n}
```

- b. Eingabe `y` Um den Neustart zu starten, wenn die folgende Eingabeaufforderung angezeigt wird:

```
The node must be rebooted to start using the newly installed software. Do  
you want to reboot now? {y|n}
```

Das Controller-Modul wird neu gestartet, stoppt aber im Startmenü, da das Boot-Gerät neu formatiert wurde und die Konfigurationsdaten wiederhergestellt werden müssen.

11. Wählen Sie den Wartungsmodus aus 5 Öffnen Sie das Startmenü, und geben Sie ein `y` Wenn Sie aufgefordert werden, den Startvorgang fortzusetzen.
12.] Überprüfen Sie, ob Controller und Chassis als `ha` konfiguriert sind:

```
ha-config show
```

Das folgende Beispiel zeigt die Ausgabe von `ha-config show` Befehl:

```
Chassis HA configuration: ha  
Controller HA configuration: ha
```



Das System zeichnet in einem PROM auf, ob es sich um ein HA-Paar oder eine eigenständige Konfiguration handelt. Der Status muss auf allen Komponenten im Standalone-System oder im HA-Paar der gleiche sein.

13. Wenn Controller und Chassis nicht als `ha` konfiguriert sind, verwenden Sie zum Korrigieren der Konfiguration die folgenden Befehle:

```
ha-config modify controller ha
```

```
ha-config modify chassis ha
```

Wenn Sie eine MetroCluster-Konfiguration haben, verwenden Sie die folgenden Befehle, um den Controller und das Chassis zu ändern:

```
ha-config modify controller mcc
```

```
ha-config modify chassis mcc
```

14. Wartungsmodus beenden:

```
halt
```

Unterbrechen Sie DAS AUTOBOOT, indem Sie an der Eingabeaufforderung der Boot-Umgebung Strg-C drücken.

15. auf node2 überprüfen Sie Datum, Uhrzeit und Zeitzone des Systems:

```
date
```

16. prüfen Sie das Datum in node3 mithilfe des folgenden Befehls an der Eingabeaufforderung der Boot-Umgebung:

```
show date
```

17. Geben Sie bei Bedarf das Datum auf node3 ein:

```
set date mm/dd/yyyy
```

18. auf node3 überprüfen Sie die Zeit mit dem folgenden Befehl an der Eingabeaufforderung der Boot-Umgebung:

```
show time
```

19. Ggf. Die Zeit auf node3 einstellen:

```
set time hh:mm:ss
```

20. legen Sie im Boot-Loader die Partner-System-ID auf node3 fest:

```
setenv partner-sysid node2_sysid
```

Für Knoten 3, partner-sysid Muss der von node2 sein.

- a. Einstellungen speichern:

```
saveenv
```

21. Überprüfen Sie den partner-sysid Für Knoten 3:

```
printenv partner-sysid
```

1. Wenn NetApp Storage Encryption (NSE) Laufwerke installiert sind, führen Sie die folgenden Schritte durch.



Falls Sie dies noch nicht bereits in der Prozedur getan haben, lesen Sie den Artikel in der Knowledge Base "[Wie erkennen Sie, ob ein Laufwerk FIPS-zertifiziert ist](#)" Ermitteln der Art der verwendeten Self-Encrypting Drives.

- a. Einstellen `bootarg.storageencryption.support` Bis `true` Oder `false`:

Wenn die folgenden Laufwerke verwendet werden...	Dann...
NSE-Laufwerke, die den Self-Encryption-Anforderungen von FIPS 140-2 Level 2 entsprechen	<code>setenv bootarg.storageencryption.support true</code>
NetApp ohne FIPS SEDs	<code>setenv bootarg.storageencryption.support false</code>



FIPS-Laufwerke können nicht mit anderen Laufwerkstypen auf demselben Node oder HA-Paar kombiniert werden. SEDs können mit Laufwerken ohne Verschlüsselung auf demselben Node oder HA-Paar kombiniert werden.

- b. Wenden Sie sich an den NetApp Support, um Hilfe beim Wiederherstellen der integrierten Schlüsselmanagementinformationen zu erhalten.

2. Starten Sie den Node im Boot-Menü:

```
boot_ontap menu
```

Was kommt als Nächstes?

- Wenn Sie ein System mit einer FC- oder UTA/UTA2-Konfiguration haben, "[Legen Sie die FC- oder UTA/UTA2-Konfiguration auf Knoten3 fest](#)".
- Wenn Sie keine FC- oder UTA/UTA2-Konfiguration haben, "[weisen Sie die Datenträger von Knoten1 Knoten3 neu zu](#)" damit Knoten3 die Festplatten von Knoten1 erkennen kann.
- Wenn Sie eine MetroCluster Konfiguration haben, "[weisen Sie die Datenträger von Knoten1 Knoten3 neu zu](#)".

Legen Sie die FC- oder UTA/UTA2-Konfiguration auf node3 fest

Wenn node3 integrierte FC-Ports, Onboard Unified Target Adapter (UTA/UTA2)-Ports oder eine UTA/UTA2-Karte hat, müssen Sie die Einstellungen konfigurieren, bevor Sie den Rest des Verfahrens abschließen.

Über diese Aufgabe

Möglicherweise müssen Sie den Abschnitt ausfüllen [Konfigurieren Sie FC-Ports auf node3](#), Der Abschnitt [UTA/UTA2-Ports in node3 prüfen und konfigurieren](#), Oder beide Abschnitte.



Unter Umständen bezieht sich bei den Marketingmaterialien von NetApp der Begriff UTA2 auf Adapter und Ports des konvergierten Netzwerkadapters (CNA). Allerdings verwendet die CLI den Begriff CNA.

Wenn Knoten 3 nicht über Onboard FC-Ports, Onboard UTA/UTA2-Ports oder eine UTA/UTA2-Karte verfügt und Sie ein System mit Storage-Festplatten aktualisieren, können Sie zu überspringen "[Weisen Sie node1-Festplatten Knoten 3 neu zu](#)".

Konfigurieren Sie FC-Ports auf node3

Wenn Knoten3 über FC-Ports verfügt (entweder integriert oder auf einem zusätzlichen FC-Adapter), müssen Sie die Portkonfigurationen auf dem Knoten festlegen, bevor Sie ihn in Betrieb nehmen, da die Ports bei der Auslieferung der Systeme nicht vorkonfiguriert sind. Wenn Sie die Ports nicht konfigurieren, kann es zu einer Dienstunterbrechung kommen.

Bevor Sie beginnen

Sie müssen die Werte der FC-Port-Einstellungen von node1 haben, die Sie im Abschnitt gespeichert haben ["Bereiten Sie die Knoten für ein Upgrade vor"](#).

Über diese Aufgabe

Sie können diesen Abschnitt überspringen, wenn Ihr System über keine FC-Konfigurationen verfügt. Wenn Ihr System über integrierte UTA/UTA2-Ports oder eine UTA/UTA2-Karte verfügt, konfigurieren Sie sie in [UTA/UTA2-Ports in node3 prüfen und konfigurieren](#).



Geben Sie die Befehle in diesem Abschnitt an der Shell-Eingabeaufforderung im Wartungsmodus ein.

Schritte

1. Vergleichen Sie die FC-Einstellungen auf Knoten3 mit den Einstellungen, die Sie zuvor von Knoten1 erfasst haben.
2. Führen Sie eine der folgenden Aktionen aus, um die FC-Ports auf Knoten3 nach Bedarf zu ändern:

Im Wartungsmodus (Option 5 im Bootmenü):

- So programmieren Sie als Zielports:

```
ucadmin modify -m fc -t target <adapter>
```

Zum Beispiel: `ucadmin modify -m fc -t target 2a`

- So programmieren Sie Initiator-Ports:

```
ucadmin modify -m fc -t initiator <adapter>
```

Zum Beispiel: `ucadmin modify -m fc -t initiator 2b`

3. Überprüfen Sie die neuen Einstellungen, indem Sie den folgenden Befehl verwenden und die Ausgabe untersuchen:

```
ucadmin show
```

4. Stoppen Sie den Knoten:

```
halt
```

5. Booten Sie das System über die LOADER-Eingabeaufforderung:

```
boot_ontap menu
```

6. Nachdem Sie den Befehl eingegeben haben, warten Sie, bis das System an der Eingabeaufforderung der Boot-Umgebung angehalten wird.

7. Wählen Sie die Option 5 Wählen Sie im Bootmenü für den Wartungsmodus aus.

8. Führen Sie eine der folgenden Aktionen aus:

- Wenn node3 eine UTA/UTA2-Karte oder Onboard-Ports zu UTA/UTA2 hat, gehen Sie zu [UTA/UTA2-Ports in node3 prüfen und konfigurieren](#).
- Wenn node3 keine UTA/UTA2-Karte oder UTA/UTA2 Onboard-Ports hat, überspringen sie [UTA/UTA2-Ports in node3 prüfen und konfigurieren](#) den "Weisen Sie node1-Festplatten Knoten 3 neu zu"Bus und gehen Sie zu .

UTA/UTA2-Ports in node3 prüfen und konfigurieren

Wenn node3 Onboard UTA/UTA2-Ports oder eine UTA/UTA2-Karte hat, müssen Sie die Konfiguration der Ports überprüfen und sie möglicherweise neu konfigurieren, je nachdem, wie Sie das aktualisierte System verwenden möchten.

Bevor Sie beginnen

Sie müssen die richtigen SFP+ Module für die UTA/UTA2-Ports besitzen.

Über diese Aufgabe

Wenn Sie einen Unified Target Adapter (UTA/UTA2)-Port für FC verwenden möchten, müssen Sie zuerst überprüfen, wie der Port konfiguriert ist.



Bei NetApp Marketingmaterialien wird möglicherweise der Begriff UTA2 verwendet, um sich auf CNA-Adapter und Ports zu beziehen. Allerdings verwendet die CLI den Begriff CNA.

Sie können die `ucadmin show` Befehl zum Anzeigen oder Überprüfen der aktuellen Portkonfiguration, wie in der folgenden Beispielausgabe gezeigt:

```
*> ucadmin show
```

Adapter	Current Mode	Current Type	Pending Mode	Pending Type	Admin Status
0e	fc	target	-	initiator	offline
0f	fc	target	-	initiator	offline
0g	fc	target	-	initiator	offline
0h	fc	target	-	initiator	offline
1a	fc	target	-	-	online
1b	fc	target	-	-	online

6 entries were displayed.

DIE UTA2-Ports können im nativen FC-Modus oder im UTA/UTA2-Modus konfiguriert werden. Der FC-Modus unterstützt FC Initiator und FC Target. Der UTA-/UTA2-Modus ermöglicht gleichzeitige NIC- und FCoE-Traffic über die gleiche 10-GbE-SFP+-Schnittstelle und unterstützt FC-Ziele.

Möglicherweise finden Sie UTA/UTA2-Anschlüsse auf einem Zusatzadapter oder auf dem Controller-Motherboard und diese weisen die folgenden Konfigurationen auf. Sie sollten jedoch die Konfiguration der UTA/UTA2-Anschlüsse auf dem Node3 überprüfen und diese gegebenenfalls ändern:

- UTA-/UTA2-Karten, die bestellt werden, werden vor dem Versand konfiguriert, um die von Ihnen geforderte

Persönlichkeit zu erhalten.

- DIE UTA2-Karten, die separat vom Controller bestellt werden, werden mit der standardmäßigen FC-Zielgruppe ausgeliefert.
- Onboard UTA/UTA2-Ports auf neuen Controllern werden vor dem Versand konfiguriert, um die Persönlichkeit zu erhalten, die Sie anfordern.



Sie müssen sich im Wartungsmodus befinden, um UTA/UTA2-Ports zu konfigurieren. Geben Sie die Befehle in diesem Abschnitt an der Shell-Eingabeaufforderung im Wartungsmodus ein.

Schritte

1. Wenn das aktuelle SFP+-Modul nicht mit der gewünschten Verwendung übereinstimmt, ersetzen Sie es durch das richtige SFP+-Modul.

Wenden Sie sich an Ihren NetApp Ansprechpartner, um das richtige SFP+ Modul zu erhalten.

2. Überprüfen Sie die UTA/UTA2-Porteinstellungen:

```
ucadmin show
```

Untersuchen Sie die Ausgabe und stellen Sie fest, ob die UTA/UTA2-Ports die gewünschte Persönlichkeit haben.

Die Ausgabe im folgenden Beispiel zeigt, dass sich der Typ des Adapters „1b“ in „Initiator“ ändert und dass sich der Modus der Adapter „2a“ und „2b“ in „cna“ ändert. Der CNA-Modus ermöglicht Ihnen, die Karte als Netzwerkadapter zu verwenden.

```
*> ucadmin show

      Current      Current      Pending      Pending      Admin
Adapter  Mode      Type      Mode      Type      Status
-----  -
1a       fc         initiator -         -         online
1b       fc         target   -         initiator online
2a       fc         target   cna        -         online
2b       fc         target   cna        -         online
*>
```

3. Führen Sie eine der folgenden Aktionen durch:

Wenn die UTA/UTA2-Ports...	Dann...
Haben Sie nicht die Persönlichkeit, die Sie wollen	Gehe zu Schritt 4 .
Haben Sie die Persönlichkeit, die Sie wollen	Überspringen Sie Schritt 4 bis Schritt 8 und gehen Sie zu Schritt 9 .

4. Führen Sie eine der folgenden Aktionen aus:

Wenn Sie konfigurieren...	Dann...
Ports auf einer UTA/UTA2-Karte	Gehe zu Schritt 5
Onboard UTA/UTA2-Ports	Überspringen Sie Schritt 5 und gehen Sie zu Schritt 6 .

5. Wenn sich der Adapter im Initiatormodus befindet und der UTA/UTA2-Port online ist, schalten Sie den UTA/UTA2-Port offline:

```
storage disable adapter <adapter_name>
```

Adapter im Zielmodus sind im Wartungsmodus automatisch offline.

6. Wenn die aktuelle Konfiguration nicht der gewünschten Verwendung entspricht, ändern Sie die Konfiguration nach Bedarf:

```
ucadmin modify -m fc|cna -t initiator|target <adapter_name>
```

- -m Ist der Persönlichkeitsmodus, fc Oder cna.
- -t Ist der Typ FC4, target Oder initiator.



Sie müssen den FC-Initiator für Bandlaufwerke und MetroCluster -Konfigurationen verwenden. Sie müssen das FC-Ziel für SAN-Clients verwenden.

7. Schalten Sie alle Zielports online, indem Sie für jeden Port einmal den folgenden Befehl eingeben:

```
storage enable adapter <adapter_name>
```

8. Verkabeln Sie den Port.

1. Beenden des Wartungsmodus:

```
halt
```

2. Starten Sie den Knoten im Startmenü, indem Sie Folgendes ausführen: `boot_ontap menu` .

Was kommt als Nächstes?

- Wenn Sie auf ein AFF A800 -System upgraden, gehen Sie zu "[Weisen Sie node1-Festplatten Knoten 3, Schritt 9, neu zu](#)" .
- Für alle anderen System-Upgrades gehen Sie zu "[Weisen Sie node1-Festplatten Knoten 3, Schritt 1, neu zu](#)" .

Weisen Sie node1-Festplatten Knoten 3 neu zu

Sie müssen die Festplatten, die zu node1 gehörten, zu node3 neu zuweisen, bevor Sie die Installation von node3 überprüfen.

Schritte

1. Überprüfen Sie, ob Knoten1 im Bootmenü angehalten wurde, und weisen Sie die Festplatten von Knoten1 Knoten3 neu zu:

`boot_after_controller_replacement`

Nach einer kurzen Verzögerung werden Sie aufgefordert, den Namen des Node einzugeben, der ersetzt wird. Wenn gemeinsam genutzte Festplatten vorhanden sind (auch Advanced Disk Partitioning (ADP) oder partitionierte Festplatten), werden Sie aufgefordert, den Node-Namen des HA-Partners einzugeben.

Diese Eingabeaufforderungen sind möglicherweise in den Konsolenmeldungen verborgen. Wenn Sie keinen Node-Namen eingeben oder einen falschen Namen eingeben, werden Sie aufgefordert, den Namen erneut einzugeben.

Erweitern Sie das Ausgabebeispiel der Konsole

```
LOADER-A> boot_ontap menu

...
*****
*
* Press Ctrl-C for Boot Menu. *
*
*****

.
.
Please choose one of the following:
(1) Normal Boot.
(2) Boot without /etc/rc.
(3) Change password.
(4) Clean configuration and initialize all disks.
(5) Maintenance mode boot.
(6) Update flash from backup config.
(7) Install new software first.
(8) Reboot node.
(9) Configure Advanced Drive Partitioning.
Selection (1-9)? 22/7

.
.
(boot_after_controller_replacement)  Boot after controller upgrade
(9a)                                  Unpartition all disks and
remove their ownership information.
(9b)                                  Clean configuration and
initialize node with partitioned disks.
(9c)                                  Clean configuration and
initialize node with whole disks.
(9d)                                  Reboot the node.
(9e)                                  Return to main boot menu.

Please choose one of the following:

(1) Normal Boot.
(2) Boot without /etc/rc.
(3) Change password.
(4) Clean configuration and initialize all disks.
(5) Maintenance mode boot.
(6) Update flash from backup config.
(7) Install new software first.
(8) Reboot node.
(9) Configure Advanced Drive Partitioning.
Selection (1-9)? boot_after_controller_replacement
```

```

.
This will replace all flash-based configuration with the last backup
to
disks. Are you sure you want to continue?: yes
.
.
Controller Replacement: Provide name of the node you would like to
replace: <name of the node being replaced>
Controller Replacement: Provide High Availability partner of node1:
<nodename of the partner of the node being replaced>
Changing sysid of node <node being replaced> disks.
Fetched sanown old_owner_sysid = 536953334 and calculated old sys id
= 536953334
Partner sysid = 4294967295, owner sysid = 536953334
.
.
.
Terminated
<node reboots>
.
.
System rebooting...
.
Restoring env file from boot media...
copy_env_file:scenario = head upgrade
Successfully restored env file from boot media...
.
.
System rebooting...
.
.
.
WARNING: System ID mismatch. This usually occurs when replacing a
boot device or NVRAM cards!
Override system ID? {y|n} y
Login:
...

```

2. Wenn das System in eine Neustartschleife gerät und die Meldung `no disks found`, liegt dies daran, dass die Ports wieder auf den Zielmodus zurückgesetzt wurden und daher keine Festplatten erkannt werden können. Ausführen [Schritt 3](#) Zu [Schritt 8](#) auf Knoten3, um dieses Problem zu beheben.
3. Drücken Sie während des AUTOBOOTS Strg-C, um den Knoten an der Eingabeaufforderung `Loader>` anzuhalten.
4. wechseln Sie an der LOADER-Eingabeaufforderung in den Wartungsmodus:

```
boot_ontap maint
```

5.] im Wartungsmodus werden alle zuvor festgelegten Initiator-Ports angezeigt, die sich jetzt im Zielmodus befinden:

```
ucadmin show
```

Ändern Sie die Ports zurück in den Initiatormodus:

```
ucadmin modify -m fc -t initiator -f adapter name
```

6. Überprüfen Sie, ob die Ports in den Initiatormodus geändert wurden:

```
ucadmin show
```

7. Wartungsmodus beenden:

```
halt
```



Wenn Sie ein Upgrade von einem System durchführen, das externe Festplatten unterstützt, auf ein System, das auch externe Festplatten unterstützt, gehen Sie zu [Schritt 8](#).

Wenn Sie ein Upgrade von einem System durchführen, das externe Festplatten auf ein System unterstützt, das sowohl interne als auch externe Festplatten unterstützt, [Schritt 9z](#).
B. ein AFF A800-System, gehen Sie zu .

8. Starten Sie an der Loader-Eingabeaufforderung:

```
boot_ontap menu
```

Beim Booten erkennt der Node jetzt alle Festplatten, die zuvor ihm zugewiesen waren, und kann wie erwartet gebootet werden.

Wenn die Clusterknoten, die Sie ersetzen, die Root-Volume-Verschlüsselung verwenden, kann ONTAP die Volume-Informationen von den Festplatten nicht lesen. Stellen Sie die Schlüssel für das Root-Volume wieder her:

- a. Zurück zum speziellen Startmenü:

```
LOADER> boot_ontap menu
```

```
Please choose one of the following:
(1) Normal Boot.
(2) Boot without /etc/rc.
(3) Change password.
(4) Clean configuration and initialize all disks.
(5) Maintenance mode boot.
(6) Update flash from backup config.
(7) Install new software first.
(8) Reboot node.
(9) Configure Advanced Drive Partitioning.
(10) Set Onboard Key Manager recovery secrets.
(11) Configure node for external key management.
```

```
Selection (1-11)? 10
```

a. Wählen Sie **(10) Set Onboard Key Manager Recovery Secrets**

b. Eingabe `y` An der folgenden Eingabeaufforderung:

```
This option must be used only in disaster recovery procedures. Are you sure?
(y or n): y
```

c. Geben Sie an der Eingabeaufforderung die Passphrase für das Schlüsselmanagement ein.

d. Geben Sie bei Aufforderung die Backup-Daten ein.



Sie müssen die Passphrase und Sicherungsdaten im erhalten haben "[Bereiten Sie die Knoten für ein Upgrade vor](#)" Abschnitt dieses Verfahrens.

e. Nachdem das System wieder zum speziellen Startmenü gestartet wurde, führen Sie die Option **(1) Normal Boot** aus



In dieser Phase ist möglicherweise ein Fehler aufgetreten. Wenn ein Fehler auftritt, wiederholen Sie die Teilschritte in [Schritt 8](#) , bis das System ordnungsgemäß gebootet wird.

9. Wenn Sie von einem System mit externen Festplatten auf ein System aktualisieren, das interne und externe Festplatten unterstützt (z. B. AFF A800 -Systeme), legen Sie das Knoten1-Aggregat als Root-Aggregat fest, um zu bestätigen, dass Knoten3 vom Root-Aggregat von Knoten1 bootet. Um das Root-Aggregat festzulegen, gehen Sie zum Boot-Menü auf Knoten3 und wählen Sie die Option 5 um in den Wartungsmodus zu wechseln.



Die folgenden Teilschritte müssen in der angegebenen Reihenfolge ausgeführt werden; andernfalls kann es zu einem Ausfall oder sogar zu Datenverlust kommen.

Im folgenden Verfahren wird node3 vom Root-Aggregat von node1 gestartet:

a. Wechseln in den Wartungsmodus:

```
boot_ontap maint
```

b. Überprüfen Sie die RAID-, Plex- und Prüfsummeninformationen für das node1 Aggregat:

```
aggr status -r
```

c. Überprüfen Sie den Status des node1-Aggregats:

```
aggr status
```

d. Bei Bedarf das node1 Aggregat online bringen:

```
aggr_online root_aggr_from_node1
```

e. Verhindern Sie, dass das node3 vom ursprünglichen Root-Aggregat gebootet wird:

```
aggr offline root_aggr_on_node3
```

f. Legen Sie das node1-Root-Aggregat als das neue Root-Aggregat für node3 fest:

```
aggr options aggr_from_node1 root
```

g. Überprüfen Sie, ob das Root-Aggregat von node3 offline ist und das Root-Aggregat für die von node1 hergebrachten Festplatten online ist und in den Root-Status eingestellt ist:

```
aggr status
```



Wenn der vorherige Unterschritt nicht ausgeführt wird, kann node3 vom internen Root-Aggregat booten, oder es kann dazu führen, dass das System eine neue Cluster-Konfiguration übernimmt oder Sie aufgefordert werden, eine zu identifizieren.

Im Folgenden wird ein Beispiel für die Befehlsausgabe angezeigt:

```
-----  
Aggr              State    Status              Options  
  
aggr0_nst_fas8080_15 online  raid_dp, aggr      root, nosnap=on  
                                fast zeroed  
                                64-bit  
  
aggr0              offline  raid_dp, aggr      diskroot  
                                fast zeroed  
                                64-bit  
-----
```

Ports von node1 nach node3 zuordnen

Sie müssen überprüfen, ob die physischen Ports auf node1 den physischen Ports auf node3 korrekt zugeordnet sind. Dadurch kann node3 nach dem Upgrade mit anderen Knoten im Cluster und mit dem Netzwerk kommunizieren.

Über diese Aufgabe

Siehe "[Quellen](#)" Verknüpfen mit *Hardware Universe*, um Informationen über die Ports auf den neuen Nodes zu erfassen. Die Informationen werden später in diesem Abschnitt verwendet.

Die Port-Einstellungen können je nach Modell der Nodes variieren. Sie müssen die Port- und LIF-Konfiguration auf dem ursprünglichen Node mit der geplanten Verwendung und Konfiguration des neuen Node kompatibel machen. Dies liegt daran, dass der neue Node beim Booten der gleichen Konfiguration wiedergibt. Dies bedeutet, dass ONTAP beim Booten von node3 versuchen wird, LIFs auf den gleichen Ports zu hosten, die in node1 verwendet wurden.

Wenn also die physischen Ports auf node1 nicht direkt den physischen Ports auf node3 zugeordnet werden, sind daher Änderungen der Software-Konfiguration erforderlich, um nach dem Booten die Cluster-, Management- und Netzwerkkonnektivität wiederherzustellen. Wenn die Cluster-Ports auf node1 nicht direkt den Cluster-Ports auf node3 zugeordnet werden, wird node3 möglicherweise nicht automatisch dem Quorum beitreten, wenn er neu gestartet wird, bis Sie die Software-Konfiguration ändern, um die Cluster-LIFs auf den richtigen physischen Ports zu hosten.

Schritte

1. Notieren Sie in der Tabelle alle Kabelinformationen für node1, die Ports, Broadcast-Domänen und IPspaces:

LIF	Anzahl an Knoten1-Ports	Node1-IPspaces	Broadcast-Domänen der Nr. 1	Node3-Ports	Node3-IPspaces	Node3 Broadcast-Domänen
Cluster 1						
Cluster 2						
Cluster 3						
Cluster 4						
Node-Management						
Cluster-Management						
Daten 1						
Daten 2						
Daten 3						
Daten 4						
San						
Intercluster-Port						

2. Zeichnen Sie alle Kabelinformationen für node3, die Ports, Broadcast-Domänen und IPspaces in der Tabelle auf.
3. Führen Sie die folgenden Schritte aus, um zu überprüfen, ob es sich bei dem Setup um ein 2-Node-Cluster ohne Switches handelt:
 - a. Legen Sie die Berechtigungsebene auf erweitert fest:

```
cluster::> set -privilege advanced
```

- b. Überprüfen Sie, ob es sich um ein 2-Node-Cluster ohne Switches handelt:

```
cluster::> network options switchless-cluster show
```

```
cluster::*> network options switchless-cluster show
```

```
Enable Switchless Cluster: false/true
```

+

Der Wert dieser Befehlsausgabe muss dem physischen Status des Systems entsprechen.

- a. Zurück zur Administratorberechtigungsebene:

```
cluster::*> set -privilege admin
```

```
cluster::>
```

4. Gehen Sie folgendermaßen vor, um Node3 in Quorum zu platzieren:

- a. Boot-Knoten 3. Siehe ["Installieren und booten Sie node3"](#) Um den Node zu booten, wenn Sie dies noch nicht getan haben.
- b. Vergewissern Sie sich, dass sich die neuen Cluster-Ports in der Cluster Broadcast-Domäne befinden:

```
network port show -node node -port port -fields broadcast-domain
```

Das folgende Beispiel zeigt, dass Port „e0a“ sich in der Cluster-Domäne auf node3 befindet:

```
cluster::> network port show -node _node3_ -port e0a -fields  
broadcast-domain
```

node	port	broadcast-domain
node3	e0a	Cluster

- c. Wenn sich die Cluster-Ports nicht in der Cluster Broadcast-Domäne befinden, fügen Sie sie mit dem folgenden Befehl hinzu:

```
broadcast-domain add-ports -ip-space Cluster -broadcast-domain Cluster -ports  
node:port
```

Dieses Beispiel fügt Cluster-Port „e1b“ auf Knoten3 hinzu:

```
network port modify -node node3 -port e1b -ip-space Cluster -mtu 9000
```


d. Fügen Sie die korrekten Ports zur Cluster Broadcast-Domäne hinzu:

```
network port modify -node -port -ipSPACE Cluster -mtu 9000
```

Dieses Beispiel fügt Cluster-Port „e1b“ auf node4 hinzu:

```
network port modify -node node4 -port e1b -ipSPACE Cluster -mtu 9000
```

e. Migrieren Sie die Cluster-LIFs zu den neuen Ports, einmal für jede LIF:

```
network interface migrate -vserver Cluster -lif lif_name -source-node node3  
-destination-node node3 -destination-port port_name
```

f. Ändern Sie den Startport der Cluster-LIFs:

```
network interface modify -vserver Cluster -lif lif_name -home-port port_name
```

g. Entfernen Sie die alten Ports aus der Cluster Broadcast-Domäne:

```
network port broadcast-domain remove-ports
```

Mit dem folgenden Befehl wird der Port „e0d“ auf node3 entfernt:

```
network port broadcast-domain remove-ports -ipSPACE Cluster -broadcast-domain  
Cluster -ports node3:e0d
```

a. Vergewissern Sie sich, dass node3 erneut dem Quorum beigetreten ist:

```
cluster show -node node3 -fields health
```

5. Anpassen der Broadcast-Domänen, die Ihre Cluster-LIFs hosten, sowie Node-Management/clustermanagement-LIFs. Vergewissern Sie sich, dass jede Broadcast-Domäne die richtigen Ports enthält. Ein Port kann nicht zwischen Broadcast-Domänen verschoben werden, wenn er als Host oder Home für eine LIF ist, sodass Sie die LIFs möglicherweise wie folgt migrieren und ändern müssen:

a. Zeigen Sie den Startport einer logischen Schnittstelle an:

```
network interface show -fields home-node,home-port
```

b. Zeigen Sie die Broadcast-Domäne an, die diesen Port enthält:

```
network port broadcast-domain show -ports node_name:port_name
```

c. Ports aus Broadcast-Domänen hinzufügen oder entfernen:

```
network port broadcast-domain add-ports
```

```
network port broadcast-domain remove-ports
```

a. Ändern Sie den Home-Port eines LIF:

```
network interface modify -vserver vserver -lif lif_name -home-port port_name
```

6. Passen Sie die Broadcast-Domänenmitgliedschaft der Netzwerkports an, die für Intercluster LIFs verwendet werden, mit denselben Befehlen an, wie in dargestellt [Schritt 5](#).
7. Passen Sie alle anderen Broadcast-Domänen an und migrieren Sie die Daten-LIFs, falls erforderlich, mit denselben Befehlen in [Schritt 5](#).
8. Wenn auf node1 keine Ports mehr vorhanden waren, löschen Sie sie mit den folgenden Schritten:
 - a. Zugriff auf die erweiterte Berechtigungsebene auf beiden Nodes:

```
set -privilege advanced
```

- b. So löschen Sie die Ports:

```
network port delete -node node_name -port port_name
```

- c. Zurück zur Administratorebene:

```
set -privilege admin
```

9. Passen Sie alle LIF Failover-Gruppen an:

```
network interface modify -failover-group failover_group -failover-policy failover_policy
```

Mit dem folgenden Befehl wird die Failover-Richtlinie auf festgelegt `broadcast-domain-wide` Und verwendet die Ports in der Failover-Gruppe „fg1“ als Failover-Ziele für LIF „data1“ auf node3:

```
network interface modify -vserver node3 -lif data1 failover-policy broadcast-domainwide -failover-group fg1
```

Siehe "[Quellen](#)" Link zu *Netzwerkverwaltung* oder den Befehlen *ONTAP 9: Manual Page Reference* für weitere Informationen.

10. Überprüfen Sie die Änderungen auf node3:

```
network port show -node node3
```

11. Jedes Cluster-LIF muss an Port 7700 zuhören. Vergewissern Sie sich, dass die Cluster-LIFs an Port 7700 zuhören:

```
::> network connections listening show -vserver Cluster
```

Port 7700, der auf Cluster-Ports hört, ist das erwartete Ergebnis, wie im folgenden Beispiel für ein Cluster mit zwei Nodes dargestellt:

```
Cluster::> network connections listening show -vserver Cluster
```

Vserver Name	Interface Name:Local Port	Protocol/Service

Node: NodeA		
Cluster	NodeA_clus1:7700	TCP/ctlopcp
Cluster	NodeA_clus2:7700	TCP/ctlopcp
Node: NodeB		
Cluster	NodeB_clus1:7700	TCP/ctlopcp
Cluster	NodeB_clus2:7700	TCP/ctlopcp
4 entries were displayed.		

12. Legen Sie für jede Cluster-LIF, die nicht an Port 7700 angehört, den Administrationsstatus der LIF auf fest down Und dann up:

```
::> net int modify -vserver Cluster -lif cluster-lif -status-admin down; net
int modify -vserver Cluster -lif cluster-lif -status-admin up
```

Wiederholen Sie Schritt 11, um zu überprüfen, ob die Cluster-LIF jetzt auf Port 7700 angehört.

Fügen Sie dem Quorum bei, wenn ein Node über einen anderen Satz an Netzwerkports verfügt

Der Node mit dem neuen Controller bootet und versucht zuerst, dem Cluster automatisch beizutreten. Wenn der neue Node jedoch einen anderen Satz an Netzwerkports aufweist, müssen Sie die folgenden Schritte durchführen, um zu bestätigen, dass der Node dem Quorum erfolgreich hinzugefügt wurde.

Über diese Aufgabe

Sie können diese Anweisungen für alle relevanten Knoten verwenden. Node3 wird in der folgenden Probe verwendet.

Schritte

- Überprüfen Sie, ob sich die neuen Cluster-Ports in der Cluster Broadcast-Domäne befinden, indem Sie den folgenden Befehl eingeben und die Ausgabe überprüfen:

```
network port show -node node -port port -fields broadcast-domain
```

Das folgende Beispiel zeigt, dass sich der Port „e1a“ in der Cluster-Domäne auf node3 befindet:

```
cluster::> network port show -node node3 -port e1a -fields broadcast-
domain
node   port broadcast-domain
-----
node3  e1a  Cluster
```

- Fügen Sie die korrekten Ports der Cluster Broadcast-Domäne hinzu, indem Sie den folgenden Befehl

eingeben und die Ausgabe überprüfen:

```
network port modify -node -port -ipSPACE Cluster -mtu 9000
```

Dieses Beispiel fügt Cluster-Port „e1b“ auf Knoten3 hinzu:

```
network port modify -node node3 -port e1b -ipSPACE Cluster -mtu 9000
```

3. Migrieren Sie die Cluster-LIFs zu den neuen Ports, einmal für jede LIF, und verwenden Sie den folgenden Befehl:

```
network interface migrate -vserver Cluster -lif lif_name -source-node node3 -  
destination-node node3 -destination-port port_name
```

4. Ändern Sie den Startport der Cluster-LIFs:

```
network interface modify -vserver Cluster -lif lif_name -home-port port_name
```

5. Wenn sich die Cluster-Ports nicht in der Cluster Broadcast-Domäne befinden, fügen Sie sie mit folgendem Befehl hinzu:

```
network port broadcast-domain add-ports -ipSPACE Cluster -broadcast-domain  
Cluster - ports node:port
```

6. Entfernen Sie die alten Ports aus der Cluster Broadcast-Domäne. Sie können für jeden relevanten Knoten verwenden. Mit dem folgenden Befehl wird der Port „e0d“ auf node3 entfernt:

```
network port broadcast-domain remove-ports network port broadcast-domain  
remove-ports ipSPACE Cluster -broadcast-domain Cluster -ports node3:e0d
```

7. Vergewissern Sie sich, dass der Node erneut dem Quorum beigetreten ist:

```
cluster show -node node3 -fields health
```

8. Passen Sie die Broadcast-Domänen an, die Ihre Cluster-LIFs und LIFs für das Node-Management/Cluster-Management hosten. Vergewissern Sie sich, dass jede Broadcast-Domäne die richtigen Ports enthält. Ein Port kann nicht zwischen Broadcast-Domänen verschoben werden, wenn er als Host oder Home für eine LIF ist, sodass Sie die LIFs möglicherweise wie folgt migrieren und ändern müssen:

- a. Zeigen Sie den Startport einer logischen Schnittstelle an:

```
network interface show -fields home-node,home-port
```

- b. Zeigen Sie die Broadcast-Domäne an, die diesen Port enthält:

```
network port broadcast-domain show -ports node_name:port_name
```

- c. Ports aus Broadcast-Domänen hinzufügen oder entfernen:

```
network port broadcast-domain add-ports network port broadcast-domain  
remove-port
```

- d. Ändern eines Startports einer LIF:

```
network interface modify -vserver vservice -lif lif_name -home-port port_name
```

Passen Sie die Intercluster-Broadcast-Domänen an und migrieren Sie gegebenenfalls die Intercluster LIFs. Die Daten-LIFs bleiben unverändert.

Überprüfen Sie die Installation von node3

Nach der Installation und dem Booten von node3 müssen Sie überprüfen, ob die Installation korrekt ist. Sie müssen warten, bis Knoten 3 dem Quorum beitreten und dann den Umzugsvorgang fortsetzen.

Über diese Aufgabe

An diesem Punkt des Verfahrens wird der Vorgang angehalten, da node3 dem Quorum beitrifft.

Schritte

1. Vergewissern Sie sich, dass node3 dem Quorum beigetreten ist:

```
cluster show -node node3 -fields health
```

2. Vergewissern Sie sich, dass node3 Teil desselben Clusters wie node2 ist und dass er sich in einem ordnungsgemäßen Zustand befindet:

```
cluster show
```

3. Überprüfen Sie den Status des Vorgangs, und überprüfen Sie, ob die Konfigurationsinformationen für node3 identisch sind mit node1:

```
system controller replace show-details
```

Wenn sich die Konfiguration für node3 unterscheidet, kann zu einem späteren Zeitpunkt eine Systemunterbrechung auftreten.

4. Überprüfen Sie, ob der ersetzte Controller für die MetroCluster-Konfiguration ordnungsgemäß konfiguriert ist, die MetroCluster-Konfiguration sollte sich im ordnungsgemäßen Zustand befinden und nicht im Switchover-Modus. Siehe ["Überprüfen Sie den Systemzustand der MetroCluster-Konfiguration"](#).

Erneutes Erstellen von VLANs, Schnittstellengruppen und Broadcast-Domänen auf Knoten3

Nachdem Sie bestätigt haben, dass node3 sich im Quorum befindet und mit node2 kommunizieren kann, müssen Sie die VLANs, Schnittstellengruppen und Broadcast-Domänen von node1 auf node3 neu erstellen. Sie müssen auch die node3-Ports zu den neu erstellten Broadcast-Domänen hinzufügen.

Über diese Aufgabe

Weitere Informationen zum Erstellen und Neuerstellen von VLANs, Schnittstellengruppen und Broadcast-Domänen finden Sie unter ["Quellen"](#) Und Link zu *Network Management*.

Schritte

1. Erstellen Sie die VLANs auf Node3 anhand der Node1-Informationen, die im aufgezeichnet wurden, erneut ["Verschieben von Aggregaten ohne Root-Wurzeln und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, auf Knoten 2"](#) Abschnitt:

```
network port vlan create -node node_name -vlan vlan-names
```

2. Erstellen Sie die Schnittstellengruppen auf node3 mit den node1-Informationen, die im aufgezeichnet wurden, erneut ["Verschieben von Aggregaten ohne Root-Wurzeln und NAS-Daten-LIFs, die sich im Besitz"](#)

von node1 befinden, auf Knoten 2" Abschnitt:

```
network port ifgrp create -node node_name -ifgrp port_ifgrp_names-distr-func
```

- Erstellen Sie die Broadcast-Domänen auf node3 mithilfe der node1-Informationen, die im aufgezeichnet wurden, erneut "Verschieben von Aggregaten ohne Root-Wurzeln und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, auf Knoten 2" Abschnitt:

```
network port broadcast-domain create -ipspace Default -broadcast-domain  
broadcast_domain_names -mtu mtu_size -ports  
node_name:port_name,node_name:port_name
```

- Fügen Sie die node3-Ports zu den neu erstellten Broadcast-Domänen hinzu:

```
network port broadcast-domain add-ports -broadcast-domain  
broadcast_domain_names -ports node_name:port_name,node_name:port_name
```

Wiederherstellung der Key-Manager-Konfiguration auf Knoten 3

Wenn Sie NetApp Aggregate Encryption (NAE) oder NetApp Volume Encryption (NVE) zur Verschlüsselung von Volumes auf dem System verwenden, das Sie aktualisieren, muss die Verschlüsselungskonfiguration mit den neuen Nodes synchronisiert werden. Wenn Sie den Schlüsselmanager nicht wiederherstellen, werden beim Verschieben der Node1-Aggregate mit ARL von node2 auf Knoten 3 verschlüsselte Volumes offline geschaltet.

Schritte

- Führen Sie zum Synchronisieren der Verschlüsselungskonfiguration für Onboard Key Manager den folgenden Befehl an der Cluster-Eingabeaufforderung aus:

Für diese ONTAP-Version...	Befehl
ONTAP 9.6 oder 9.7	<code>security key-manager onboard sync</code>
ONTAP 9.5	<code>security key-manager setup -node node_name</code>

- Geben Sie die Cluster-weite Passphrase für das Onboard Key Manager ein.

Verschieben Sie Aggregate ohne Root-Root-Fehler und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, von node2 auf node3

Nachdem Sie die Installation node3 überprüft haben und bevor Sie Aggregate von node2 auf node3 verschieben, müssen Sie die NAS-Daten-LIFs von node1 verschieben, die sich derzeit in node2 von node2 auf node3 befinden. Sie müssen außerdem überprüfen, ob die SAN-LIFs auf node3 vorhanden sind.

Über diese Aufgabe

Remote-LIFs verarbeiten den Datenverkehr zu SAN-LUNs während des Upgrades. Das Verschieben von SAN-LIFs ist für den Zustand des Clusters oder des Service während des Upgrades nicht erforderlich. SAN LIFs werden nicht verschoben, es sei denn, sie müssen neuen Ports zugeordnet werden. Sie überprüfen, ob die LIFs sich in einem ordnungsgemäßen Zustand befinden und sich auf den entsprechenden Ports befinden, nachdem Sie node3 in den Online-Modus versetzt haben.

Schritte

1. Wiederaufnahme des Betriebs der Versetzung:

```
system controller replace resume
```

Das System führt die folgenden Aufgaben aus:

- Cluster-Quorum-Prüfung
- System-ID-Prüfung
- Prüfung der Bildversion
- Überprüfung der Zielplattform
- Prüfung der Netzwerkanachbarkeit

Der Vorgang unterbricht in dieser Phase in der Überprüfung der Netzwerknachprüfbarkeit.

2. Überprüfen Sie manuell, ob das Netzwerk und alle VLANs, Schnittstellengruppen und Broadcast-Domänen korrekt konfiguriert wurden.

3. Wiederaufnahme des Betriebs der Versetzung:

```
system controller replace resume
```

To complete the "Network Reachability" phase, ONTAP network configuration must be manually adjusted to match the new physical network configuration of the hardware. This includes assigning network ports to the correct broadcast domains, creating any required ifgrps and VLANs, and modifying the home-port parameter of network interfaces to the appropriate ports. Refer to the "Using aggregate relocation to upgrade controller hardware on a pair of nodes running ONTAP 9.x" documentation, Stages 3 and 5. Have all of these steps been manually completed? [y/n]

4. Eingabe `y` Um fortzufahren.

5. Das System führt folgende Prüfungen durch:

- Cluster-Zustandsprüfung
- LIF-Statusüberprüfung für Cluster

Nach Durchführung dieser Prüfungen verschiebt das System die nicht-Root-Aggregate und NAS-Daten-LIFs, die sich im Besitz von node1 befinden, auf den neuen Controller, node3. Das System hält an, sobald die Ressourcenverlagerung abgeschlossen ist.

6. Überprüfen Sie den Status der Aggregatverschiebung und der LIF-Verschiebung von NAS-Daten:

```
system controller replace show-details
```

7. Überprüfen Sie, ob die nicht-Root-Aggregate und NAS-Daten-LIFs erfolgreich in node3 verschoben wurden.

Falls Aggregate nicht verschoben oder ein Veto eingesetzt werden kann, müssen sie die Aggregate manuell verschieben oder – falls erforderlich – entweder die Vetos oder die Zielprüfungen außer Kraft setzen. Siehe "[Verschiebung ausgefallener oder Vetos von Aggregaten](#)" Finden Sie weitere Informationen.

8. Überprüfen Sie, ob sich die SAN-LIFs auf den richtigen Ports auf node3 befinden, indem Sie die folgenden Teilschritte ausführen:

- a. Geben Sie den folgenden Befehl ein und überprüfen Sie die Ausgabe:

```
network interface show -data-protocol iscsi|fc -home-node node3
```

Das System gibt die Ausgabe wie im folgenden Beispiel zurück:

```
cluster::> net int show -data-protocol iscsi|fc -home-node node3
```

	Logical	Status	Network	Current	Current	Is
Vserver	Interface	Admin/Oper	Address/Mask	Node	Port	Home
-----	-----	-----	-----	-----	-----	----
vs0						
	a0a	up/down	10.63.0.53/24	node3	a0a	true
	data1	up/up	10.63.0.50/18	node3	e0c	true
	rads1	up/up	10.63.0.51/18	node3	e1a	true
	rads2	up/down	10.63.0.52/24	node3	e1b	true
vs1						
	lif1	up/up	172.17.176.120/24	node3	e0c	true
	lif2	up/up	172.17.176.121/24	node3	e1a	true

- b. Wenn node3 irgendwelche SAN-LIFs oder Gruppen von SAN-LIFs hat, die sich auf einem Port befinden, der nicht in node1 vorhanden war oder einem anderen Port zugeordnet werden muss, verschieben Sie sie zu einem geeigneten Port auf node3, indem Sie die folgenden Teilschritte ausführen:

- i. Setzen Sie den LIF-Status auf „down“:

```
network interface modify -vserver Vserver_name -lif LIF_name -status  
-admin down
```

- ii. Entfernen Sie das LIF aus dem Portsatz:

```
portset remove -vserver Vserver_name -portset portset_name -port-name  
port_name
```

- iii. Geben Sie einen der folgenden Befehle ein:

- Verschieben eines einzelnen LIF:

```
network interface modify -vserver Vserver_name -lif LIF_name -home
```



```
-port new_home_port
```

- Verschieben Sie alle LIFs auf einem einzelnen nicht vorhandenen oder falschen Port in einen neuen Port:

```
network interface modify {-home-port port_on_node1 -home-node node1  
-role data} -home-port new_home_port_on_node3
```

- Fügen Sie die LIFs wieder dem Portsatz hinzu:

```
portset add -vserver Vserver_name -portset portset_name -port-name  
port_name
```



Sie müssen bestätigen, dass Sie SAN-LIFs zu einem Port mit der gleichen Verbindungsgeschwindigkeit wie der ursprüngliche Port verschoben haben.

- a. Ändern Sie den Status aller LIFs auf „up“, damit die LIFs den Datenverkehr auf dem Node akzeptieren und senden können:

```
network interface modify -home-port port_name -home-node node3 -lif data  
-status admin up
```

- b. Geben Sie an jedem Node den folgenden Befehl ein, und überprüfen Sie seine Ausgabe, um zu überprüfen, ob LIFs zu den richtigen Ports verschoben wurden und ob die LIFs den Status von aufweisen up:

```
network interface show -home-node node3 -role data
```

- c. Wenn irgendwelche LIFs ausgefallen sind, setzen Sie den Administratorstatus der LIFs auf up Geben Sie den folgenden Befehl ein, einmal für jede LIF:

```
network interface modify -vserver vserver_name -lif lif_name -status-admin  
up
```

9. Setzen Sie den Vorgang fort, um das System zur Durchführung der erforderlichen Nachprüfungen zu auffordern:

```
system controller replace resume
```

Das System führt die folgenden Nachprüfungen durch:

- Cluster-Quorum-Prüfung
- Cluster-Zustandsprüfung
- Aggregatrekonstruktion
- Aggregatstatus-Prüfung
- Überprüfung des Festplattenstatus
- LIF-Statusüberprüfung für Cluster

Phase 4: Knoten2 verschieben und ausmustern

Verschieben von Aggregaten und NAS-Daten-LIFs ohne Root-Wurzeln von Knoten 2 auf Knoten 3

Bevor Sie node2 durch node4 ersetzen, verschieben Sie die nicht-Root-Aggregate und NAS-Daten-LIFs, die im Besitz von node2 sind, auf node3.

Bevor Sie beginnen

Nach den Nachprüfungen aus der vorherigen Phase wird automatisch die Ressourcenfreigabe für node2 gestartet. Die Aggregate außerhalb des Root-Bereichs und LIFs für nicht-SAN-Daten werden von node2 auf node3 migriert.

Über diese Aufgabe

Remote-LIFs verarbeiten den Datenverkehr zu SAN-LUNs während des Upgrades. Das Verschieben von SAN-LIFs ist für den Zustand des Clusters oder des Service während des Upgrades nicht erforderlich.

Nach der Migration der Aggregate und LIFs wird der Vorgang zu Verifizierungszwecken angehalten. In dieser Phase müssen Sie überprüfen, ob alle Aggregate ohne Root-Root-Daten und LIFs außerhalb des SAN in node3 migriert werden.



Der Home-Inhaber für die Aggregate und LIFs werden nicht geändert, nur der aktuelle Besitzer wird geändert.

Schritte

1. Vergewissern Sie sich, dass alle nicht-Root-Aggregate online sind und ihren Status auf node3:

```
storage aggregate show -node <node3> -state online -root false
```

Das folgende Beispiel zeigt, dass die nicht-Root-Aggregate auf node2 online sind:

```
cluster::> storage aggregate show -node node3 state online -root false
```

Aggregate	Size	Available	Used%	State	#Vols	Nodes
RAID	Status					
-----	-----	-----	-----	-----	-----	-----
aggr_1	744.9GB	744.8GB	0%	online	5	node2
raid_dp	normal					
aggr_2	825.0GB	825.0GB	0%	online	1	node2
raid_dp	normal					

2 entries were displayed.

Wenn die Aggregate offline sind oder in node3 offline sind, bringen Sie sie mit dem folgenden Befehl auf node3 online, einmal für jedes Aggregat:

```
storage aggregate online -aggregate <aggregate_name>
```

2. Überprüfen Sie, ob alle Volumes auf node3 online sind, indem Sie den folgenden Befehl auf node3 verwenden und die Ausgabe überprüfen:

```
volume show -node <node3> -state offline
```

Wenn ein Volume auf node3 offline ist, schalten Sie sie online. Verwenden Sie dazu den folgenden Befehl auf node3, einmal für jedes Volume:

```
volume online -vserver <vserver_name> -volume <volume_name>
```

Der `vserver_name` Die Verwendung mit diesem Befehl finden Sie in der Ausgabe des vorherigen `volume show` Befehl.

3. Überprüfen Sie, ob die LIFs zu den richtigen Ports verschoben wurden und über den Status von verfügen up. Wenn irgendwelche LIFs ausgefallen sind, setzen Sie den Administratorstatus der LIFs auf up. Geben Sie den folgenden Befehl ein, einmal für jede LIF:

```
network interface modify -vserver <vserver_name> -lif <LIF_name> -home-node <node_name> -status-admin up
```

4. Wenn die Ports, die derzeit Daten-LIFs hosten, nicht auf der neuen Hardware vorhanden sind, entfernen Sie diese aus der Broadcast-Domäne:

```
network port broadcast-domain remove-ports
```

5. Überprüfen Sie, ob auf node2 keine Daten-LIFs bleiben, indem Sie den folgenden Befehl eingeben und die Ausgabe überprüfen:

```
network interface show -curr-node node2 -role data
```

6. Wenn Schnittstellengruppen oder VLANs konfiguriert sind, führen Sie die folgenden Teilschritte aus:

- a. Notieren Sie VLAN- und Schnittstellengruppeninformationen, damit Sie die VLANs und Schnittstellengruppen auf node3 neu erstellen können, nachdem node3 gestartet wurde.
- b. Entfernen Sie die VLANs aus den Schnittstellengruppen:

```
network port vlan delete -node nodename -port ifgrp -vlan-id VLAN_ID
```

- c. Überprüfen Sie, ob auf dem Node Schnittstellengruppen konfiguriert sind, indem Sie den folgenden Befehl eingeben und seine Ausgabe überprüfen:

```
network port ifgrp show -node node2 -ifgrp ifgrp_name -instance
```

Das System zeigt Schnittstellengruppeninformationen für den Node an, wie im folgenden Beispiel gezeigt:

```
cluster::> network port ifgrp show -node node2 -ifgrp a0a -instance
Node: node3
Interface Group Name: a0a
Distribution Function: ip
Create Policy: multimode_lacp
MAC Address: 02:a0:98:17:dc:d4
Port Participation: partial
Network Ports: e2c, e2d
Up Ports: e2c
Down Ports: e2d
```

- a. Wenn Schnittstellengruppen auf dem Node konfiguriert sind, notieren Sie die Namen dieser Gruppen und der ihnen zugewiesenen Ports. Löschen Sie dann die Ports, indem Sie den folgenden Befehl eingeben, und zwar einmal für jeden Port:

```
network port ifgrp remove-port -node nodename -ifgrp ifgrp_name -port
netport
```

Node2 ausmustern

Um Knoten2 außer Betrieb zu nehmen, fahren Sie Knoten2 ordnungsgemäß herunter und entfernen ihn dann aus dem Rack oder Gehäuse.

Schritte

1. Vorgang fortsetzen:

```
system controller replace resume
```

Der Knoten wird automatisch angehalten.

Nachdem Sie fertig sind

Sie können nach Abschluss des Upgrades die Decommission node2 deaktivieren. Siehe "[Ausmustern des alten Systems](#)".

Phase 5: installieren und booten sie node4

installieren und booten sie node4

Sie müssen node4 im Rack installieren, die node2-Verbindungen nach node4 übertragen, node4 booten und ONTAP installieren. Sie müssen dann freie Festplatten in node2 neu zuweisen, sämtliche Festplatten, die zum Root-Volume gehören, und alle nicht-Root-Aggregate, die zuvor nicht in node3 verschoben wurden, wie in diesem Abschnitt beschrieben.

Über diese Aufgabe

Der Umzugsvorgang wird zu Beginn dieser Phase angehalten. Dieser Vorgang wird größtenteils automatisch

durchgeführt. Der Vorgang hält an, damit Sie seinen Status überprüfen können. Sie müssen den Vorgang manuell fortsetzen. Außerdem müssen Sie überprüfen, ob die NAS-Daten-LIFs erfolgreich in node4 verschoben wurden.

Sie müssen node4 neu starten, wenn die ONTAP Version auf node4 von der ONTAP Version auf node2 abweicht. Nach der Installation von node4 starten Sie es vom ONTAP 9-Image, das auf dem Webserver gespeichert ist. Anschließend können Sie die richtigen Dateien für nachfolgende Systemstarts auf das Startmedium herunterladen, indem Sie den Anweisungen in ["Vorbereitungen für den Netzboot"](#)Die



- Bei einem AFF A800 oder AFF C800 -Controller-Upgrade müssen Sie sicherstellen, dass alle Laufwerke im Gehäuse fest an der Mittelplatine sitzen, bevor Sie Knoten 2 entfernen. Weitere Informationen finden Sie unter ["Ersetzen Sie die AFF A800- oder AFF C800-Controller-Module"](#).
- Wenn Sie ein System mit Speicherplatten aktualisieren, müssen Sie diesen Abschnitt vollständig ausfüllen und dann mit ["Legen Sie die FC- oder UTA/UTA2-Konfiguration auf node4 fest"](#) , indem Sie Befehle an der Cluster-Eingabeaufforderung eingeben.

Schritte

1. stellen Sie sicher, dass node4 über ausreichend Rack-Platz verfügt.

Wenn node4 sich in einem separaten Chassis von node2 befindet, können sie node4 an der gleichen Stelle wie node3 platzieren. Wenn sich Node2 und node4 im selben Chassis befinden, befindet sich node4 bereits in der entsprechenden Rack-Position.

2. installieren sie node4 im Rack gemäß den Anweisungen in der Anleitung *Installation and Setup Instructions* für das Node-Modell.
3. Kabel node4, ziehen Sie die Verbindungen von node2 nach node4.

Verkabeln Sie die folgenden Verbindungen mithilfe der *Installations- und Einrichtungsanweisungen* für die Node4-Plattform, des entsprechenden Disk-Shelf-Dokuments und der *HA-Paarverwaltung*-Dokumentation.

Siehe ["Quellen"](#) zur Verknüpfung mit *HA-Paarverwaltung*.

- Konsole (Remote-Management-Port)
- Cluster-Ports
- Datenports
- Cluster- und Node-Management-Ports
- Storage
- SAN-Konfigurationen: iSCSI Ethernet und FC Switch Ports



Möglicherweise müssen Sie die Interconnect-Karte/FC-VI-Karte oder die Interconnect/FC-VI-Kabelverbindung von node2 auf node4 nicht verschieben, da die meisten Plattform-Modelle über einzigartige Interconnect-Kartenmodelle verfügen. Bei der MetroCluster Konfiguration müssen Sie die FC-VI-Kabelverbindungen von node2 nach node4 verschieben. Wenn der neue Host keine FC-VI-Karte besitzt, müssen Sie möglicherweise die FC-VI-Karte verschieben.

4. Schalten Sie node4 ein, und unterbrechen Sie den Bootvorgang, indem Sie am Konsolenterminal Strg-C drücken, um auf die Eingabeaufforderung der Boot-Umgebung zuzugreifen.



Wenn Sie node4 booten, wird möglicherweise die folgende Warnmeldung angezeigt:

```
WARNING: The battery is unfit to retain data during a power outage. This
is likely
        because the battery is discharged but could be due to other
temporary
        conditions.
        When the battery is ready, the boot process will complete
        and services will be engaged. To override this delay, press 'c'
followed
        by 'Enter'
```

5. Wenn die Warnmeldung in Schritt 4 angezeigt wird, führen Sie die folgenden Schritte aus:

- a. Überprüfen Sie auf Meldungen der Konsole, die auf ein anderes Problem als eine schwache NVRAM-Batterie hinweisen und ergreifen Sie gegebenenfalls erforderliche Korrekturmaßnahmen.
- b. Warten Sie, bis der Akku geladen ist und der Bootvorgang abgeschlossen ist.



Überschreiben Sie die Verzögerung nicht. Wenn Sie den Akku nicht aufladen, kann dies zu Datenverlust führen.



Siehe "[Vorbereitungen für den Netzboot](#)".

6. Konfigurieren Sie die Netzboot-Verbindung, indem Sie eine der folgenden Aktionen auswählen.



Sie müssen den Management-Port und die IP als Netzboot-Verbindung verwenden. Verwenden Sie keine Daten-LIF-IP, oder es kann während des Upgrades ein Datenausfall auftreten.

Wenn DHCP (Dynamic Host Configuration Protocol) lautet...	Dann...
Wird Ausgeführt	Konfigurieren Sie die Verbindung automatisch, indem Sie an der Eingabeaufforderung der Boot-Umgebung den folgenden Befehl eingeben: <code>ifconfig e0M -auto</code>

Wenn DHCP (Dynamic Host Configuration Protocol) lautet...	Dann...
Nicht ausgeführt	Konfigurieren Sie die Verbindung manuell, indem Sie an der Eingabeaufforderung der Boot-Umgebung den folgenden Befehl eingeben: <pre>ifconfig e0M -addr=<i>filer_addr</i> -mask=<i>netmask</i> -gw=<i>gateway</i> -dns=<i>dns_addr</i> -domain=<i>dns_domain</i></pre> <i>filer_addr</i> Ist die IP-Adresse des Speichersystems (obligatorisch). <i>netmask</i> Ist die Netzwerkmaske des Storage-Systems (erforderlich). <i>gateway</i> Ist das Gateway für das Speichersystem (erforderlich). <i>dns_addr</i> Ist die IP-Adresse eines Namensservers in Ihrem Netzwerk (optional). <i>dns_domain</i> Der Domain Name (DNS) ist der Domain-Name. Wenn Sie diesen optionalen Parameter verwenden, benötigen Sie in der Netzboot-Server-URL keinen vollqualifizierten Domännennamen. Sie benötigen nur den Host-Namen des Servers. HINWEIS: Andere Parameter können für Ihre Schnittstelle erforderlich sein. Eingabe <code>help ifconfig</code> Details finden Sie in der Firmware-Eingabeaufforderung.

7. Ausführen eines Netzboots auf node4:

Für...	Dann...
Systeme der FAS/AFF8000 Serie	<code>netboot http://<web_server_ip/path_to_web-accessible_directory>/netboot/kernel</code>
Alle anderen Systeme	<code>netboot http://<web_server_ip/path_to_web-accessible_directory>/<ontap_version>_image.tgz</code>

Der `<path_to_the_web-accessible_directory>` Sollten Sie dazu führen, wo Sie das heruntergeladen haben `<ontap_version>_image.tgz` In Schritt 1 im Abschnitt "[Vorbereitungen für den Netzboot](#)".



Unterbrechen Sie den Startvorgang nicht.

8. Wählen Sie im Startmenü Option (7) Install new software first.

Mit dieser Menüoption wird das neue ONTAP-Image auf das Startgerät heruntergeladen und installiert.

Ignorieren Sie die folgende Meldung:

`This procedure is not supported for Non-Disruptive Upgrade on an HA pair`

Der Hinweis gilt für unterbrechungsfreie Upgrades der ONTAP und keine Upgrades von Controllern.



Aktualisieren Sie den neuen Node immer als Netzboot auf das gewünschte Image. Wenn Sie eine andere Methode zur Installation des Images auf dem neuen Controller verwenden, wird möglicherweise das falsche Image installiert. Dieses Problem gilt für alle ONTAP Versionen. Das Netzboot wird mit der Option kombiniert (7) `Install new software` Entfernt das Boot-Medium und platziert dieselbe ONTAP-Version auf beiden Image-Partitionen.

9. Wenn Sie aufgefordert werden, den Vorgang fortzusetzen, geben Sie ein `y`. Und wenn Sie zur Eingabe des Pakets aufgefordert werden, geben Sie die URL ein:

```
http://<web_server_ip/path_to_web-  
accessible_directory>/<ontap_version>_image.tgz
```

10. Führen Sie die folgenden Teilschritte durch, um das Controller-Modul neu zu booten:

- a. Eingabe `n` So überspringen Sie die Backup-Recovery, wenn folgende Eingabeaufforderung angezeigt wird:

```
Do you want to restore the backup configuration now? {y|n}
```

- b. Starten Sie den Neustart durch Eingabe `y` Wenn die folgende Eingabeaufforderung angezeigt wird:

```
The node must be rebooted to start using the newly installed  
software. Do you want to reboot now? {y|n}
```

Das Controller-Modul wird neu gestartet, stoppt aber im Startmenü, da das Boot-Gerät neu formatiert wurde und die Konfigurationsdaten wiederhergestellt werden müssen.

11. Wählen Sie Wartungsmodus 5 Öffnen Sie das Startmenü, und geben Sie ein `y` Wenn Sie aufgefordert werden, den Startvorgang fortzusetzen.
12. Vergewissern Sie sich, dass Controller und Chassis als HA konfiguriert sind:

```
ha-config show
```

Das folgende Beispiel zeigt die Ausgabe von `ha-config show` Befehl:

```
Chassis HA configuration: ha  
Controller HA configuration: ha
```



Das System zeichnet in einem PROM auf, ob es sich um ein HA-Paar oder eine eigenständige Konfiguration handelt. Der Status muss auf allen Komponenten im Standalone-System oder im HA-Paar der gleiche sein.

13. Wenn der Controller und das Chassis nicht als HA konfiguriert wurden, verwenden Sie zum Korrigieren der Konfiguration die folgenden Befehle:

```
ha-config modify controller ha
```



```
ha-config modify chassis ha
```

Wenn Sie eine MetroCluster-Konfiguration haben, verwenden Sie die folgenden Befehle, um den Controller und das Chassis zu ändern:

```
ha-config modify controller mcc
```

```
ha-config modify chassis mcc
```

14. Beenden des Wartungsmodus:

```
halt
```

Unterbrechen Sie DAS AUTOBOOT, indem Sie an der Eingabeaufforderung der Boot-Umgebung Strg-C drücken.

15. auf node3 überprüfen Sie Datum, Uhrzeit und Zeitzone des Systems:

```
date
```

16. Überprüfen Sie am node4 das Datum mithilfe des folgenden Befehls an der Eingabeaufforderung der Boot-Umgebung:

```
show date
```

17. Legen Sie bei Bedarf das Datum auf node4 fest:

```
set date mm/dd/yyyy
```

18. Überprüfen Sie auf node4 die Zeit mit dem folgenden Befehl an der Eingabeaufforderung der Boot-Umgebung:

```
show time
```

19. Stellen Sie bei Bedarf die Uhrzeit auf node4 ein:

```
set time hh:mm:ss
```

20. Legen Sie im Boot-Loader die Partner-System-ID auf node4 fest:

```
setenv partner-sysid node3_sysid
```

Für node4, `partner-sysid` Muss das der Node3 sein.

Einstellungen speichern:

```
saveenv
```

21. `[[Auto_install4_step21]` Verify the `partner-sysid` für node4:

```
printenv partner-sysid
```

22. Wenn Sie NetApp Storage Encryption (NSE)-Laufwerke installiert haben, führen Sie die folgenden Schritte aus:



Falls Sie dies noch nicht bereits in der Prozedur getan haben, lesen Sie den Artikel in der Knowledge Base "[Wie erkennen Sie, ob ein Laufwerk FIPS-zertifiziert ist](#)" Ermitteln der Art der verwendeten Self-Encrypting Drives.

a. Einstellen `bootarg.storageencryption.support` Bis `true` Oder `false`:

Wenn die folgenden Laufwerke verwendet werden...	Dann...
NSE-Laufwerke, die den Self-Encryption-Anforderungen von FIPS 140-2 Level 2 entsprechen	<code>setenv bootarg.storageencryption.support true</code>
NetApp ohne FIPS SEDs	<code>setenv bootarg.storageencryption.support false</code>



FIPS-Laufwerke können nicht mit anderen Laufwerkstypen auf demselben Node oder HA-Paar kombiniert werden. SEDs können mit Laufwerken ohne Verschlüsselung auf demselben Node oder HA-Paar kombiniert werden.

b. Wenden Sie sich an den NetApp Support, um Hilfe beim Wiederherstellen der integrierten Schlüsselmanagementinformationen zu erhalten.

23. Starten Sie den Node im Boot-Menü:

```
boot_ontap menu
```

Was kommt als Nächstes?

- Wenn Sie ein System mit einer FC- oder UTA/UTA2-Konfiguration haben, "[Legen Sie die FC- oder UTA/UTA2-Konfiguration auf Knoten4 fest](#)" .
- Wenn Sie keine FC- oder UTA/UTA2-Konfiguration haben, "[Neuzuweisung der Node2-Festplatten zu Node4, Schritt 1](#)" damit Knoten4 die Festplatten von Knoten2 erkennen kann.
- Wenn Sie eine MetroCluster Konfiguration haben, "[Legen Sie die FC- oder UTA/UTA2-Konfiguration auf Knoten4 fest](#)" um die an den Knoten angeschlossenen Festplatten zu erkennen.

Legen Sie die FC- oder UTA/UTA2-Konfiguration auf node4 fest

Wenn node4 über integrierte FC-Ports, integrierte Unified Target Adapter (UTA/UTA2)-Ports oder eine UTA/UTA2-Karte verfügt, müssen Sie die Einstellungen konfigurieren, bevor Sie den Rest des Verfahrens abschließen.

Über diese Aufgabe

Möglicherweise müssen Sie Folgendes ausfüllen: [Konfigurieren Sie FC-Ports auf node4](#) oder [UTA/UTA2-Ports auf node4 prüfen und konfigurieren](#) oder beide Abschnitte.



Wenn node4 nicht über integrierte FC-Ports, Onboard UTA/UTA2-Ports oder eine UTA/UTA2-Karte verfügt und Sie ein System mit Storage-Festplatten aktualisieren, können Sie zu überspringen "[Weisen Sie Node2-Festplatten node4 neu zu](#)".

stellen Sie sicher, dass node4 über ausreichend Rack-Platz verfügt. Wenn node4 sich in einem separaten Chassis von node2 befindet, können sie node4 an der gleichen Stelle wie node3 platzieren. Wenn sich Node2 und node4 im selben Chassis befinden, befindet sich node4 bereits in der entsprechenden Rack-Position.

Konfigurieren Sie FC-Ports auf node4

Wenn Knoten 4 über FC-Ports verfügt (entweder integriert oder auf einem zusätzlichen FC-Adapter), müssen Sie die Portkonfigurationen auf dem Knoten festlegen, bevor Sie ihn in Betrieb nehmen, da die Ports bei der Auslieferung der Systeme nicht vorkonfiguriert sind. Wenn Sie die Ports nicht wie erforderlich konfigurieren, kann es zu einer Dienstunterbrechung kommen.

Bevor Sie beginnen

Sie müssen die Werte der FC-Port-Einstellungen von node2 haben, die Sie im Abschnitt gespeichert haben "[Bereiten Sie die Knoten für ein Upgrade vor](#)".

Über diese Aufgabe

Sie können diesen Abschnitt überspringen, wenn Ihr System über keine FC-Konfigurationen verfügt. Wenn Ihr System über integrierte UTA/UTA2-Ports oder einen UTA/UTA2-Adapter verfügt, konfigurieren Sie sie in [UTA/UTA2-Ports auf node4 prüfen und konfigurieren](#).



Geben Sie die Befehle in diesem Abschnitt an der Shell-Eingabeaufforderung im Wartungsmodus ein.

Schritte

1. Zeigen Sie Informationen zu allen FC- und konvergenten Netzwerkadaptern auf dem System an:

```
system node hardware unified-connect show
```

2. Vergleichen Sie die FC-Einstellungen auf node4 mit den Einstellungen, die Sie zuvor aus node1 erfasst haben.
3. Ändern Sie die FC-Ports auf node4 nach Bedarf:

- So programmieren Sie als Zielports:

```
ucadmin modify -m fc -t target adapter
```

Zum Beispiel: `ucadmin modify -m fc -t target 2a`

- So programmieren Sie Initiator-Ports:

```
ucadmin modify -m fc -t initiator adapter
```

-t Ist der FC4-Typ: Target oder Initiator.

Zum Beispiel: `ucadmin modify -m fc -t initiator 2b`

4. Stoppen Sie den Knoten:

halt

5. Booten Sie das System über die LOADER-Eingabeaufforderung:

```
boot_ontap menu
```

6. Nachdem Sie den Befehl eingegeben haben, warten Sie, bis das System an der Eingabeaufforderung der Boot-Umgebung angehalten wird.

7. Wählen Sie die Option 5 Wählen Sie im Bootmenü für den Wartungsmodus aus.

8. Nehmen Sie eine der folgenden Aktionen:

- Gehen Sie zu [UTA/UTA2-Ports auf node4 prüfen und konfigurieren](#) Bei node4 mit einer UTA/UTA2-Karte oder Onboard-Ports UTA/UTA2:
- Wenn Knoten4 keine UTA/UTA2-Karte oder UTA/UTA2-Onboard-Ports hat, überspringen Sie *Überprüfen und Konfigurieren der UTA/UTA2-Ports auf Knoten4* und gehen Sie zu "[Weisen Sie Node2-Festplatten node4 neu zu](#)".

UTA/UTA2-Ports auf node4 prüfen und konfigurieren

Wenn node4 Onboard UTA/UTA2-Ports oder eine UTA/UTA2-Karte hat, müssen Sie die Konfiguration der Ports überprüfen und sie je nach Nutzung des aktualisierten Systems konfigurieren.

Bevor Sie beginnen

Sie müssen die richtigen SFP+ Module für die UTA/UTA2-Ports besitzen.

Über diese Aufgabe

DIE UTA2-Ports können im nativen FC-Modus oder im UTA/UTA2-Modus konfiguriert werden. Der FC-Modus unterstützt FC Initiator und FC Target. Der UTA-/UTA2-Modus ermöglicht es, gleichzeitig NIC- und FCoE-Datenverkehr die gleiche 10-GbE-SFP+-Schnittstelle zu nutzen und das FC-Ziel zu unterstützen.



Bei NetApp Marketingmaterialien wird möglicherweise der Begriff UTA2 verwendet, um sich auf CNA-Adapter und Ports zu beziehen. Allerdings verwendet die CLI den Begriff CNA.

UTA2-Ports können an einem Adapter oder auf dem Controller mit den folgenden Konfigurationen verwendet werden:

- UTA-/UTA2-Karten, die gleichzeitig mit dem Controller bestellt wurden, werden vor dem Versand konfiguriert, um die von Ihnen angeforderte Persönlichkeit zu erhalten.
- DIE UTA2-Karten, die separat vom Controller bestellt werden, werden mit der standardmäßigen FC-Zielgruppe ausgeliefert.
- Onboard UTA/UTA2-Ports auf neuen Controllern werden konfiguriert (vor dem Versand), um die von Ihnen angeforderte Persönlichkeit zu besitzen.

Sie sollten jedoch die Konfiguration der UTA/UTA2-Ports auf node4 überprüfen und sie gegebenenfalls ändern.



Geben Sie die Befehle in diesem Abschnitt an der Shell-Eingabeaufforderung im Wartungsmodus ein.

Schritte

1. Überprüfen Sie, wie die Ports derzeit auf Knoten4 konfiguriert sind:

```
system node hardware unified-connect show
```

Das System zeigt eine Ausgabe wie im folgenden Beispiel an:

```
*> ucadmin show
```

Node	Adapter	Current Mode	Current Type	Pending Mode	Pending Type	Admin Status
f-a	0e	fc	initiator	-	-	online
f-a	0f	fc	initiator	-	-	online
f-a	0g	cna	target	-	-	online
f-a	0h	cna	target	-	-	online
f-a	0e	fc	initiator	-	-	online
f-a	0f	fc	initiator	-	-	online
f-a	0g	cna	target	-	-	online
f-a	0h	cna	target	-	-	online

```
*>
```

2. Wenn das aktuelle SFP+-Modul nicht mit der gewünschten Verwendung übereinstimmt, ersetzen Sie es durch das richtige SFP+-Modul.

Wenden Sie sich an Ihren NetApp Ansprechpartner, um das richtige SFP+ Modul zu erhalten.

3. Überprüfen Sie die Einstellungen:

```
ucadmin show
```

Überprüfen Sie die Ausgabe des `ucadmin show` Führen Sie einen Befehl aus, und bestimmen Sie, ob die UTA/UTA2-Ports die gewünschte Persönlichkeit haben.

Die Ausgabe in den folgenden Beispielen zeigt, dass sich der Adaptertyp „1b“ in ändert `initiator` Und dass sich der Modus der Adapter „2a“ und „2b“ in ändert `cna`:

```
*> ucaadmin show
Node   Adapter  Current Mode  Current Type  Pending Mode  Pending Type
Admin Status
-----
-----
f-a    1a        fc           initiator     -             -
online
f-a    1b        fc           target        -             initiator
online
f-a    2a        fc           target        cna           -
online
f-a    2b        fc           target        cna           -
online
4 entries were displayed.
*>
```

4. Führen Sie eine der folgenden Aktionen durch:

Wenn die CNA-Ports...	Dann...
Haben Sie nicht die Persönlichkeit, die Sie wollen	Gehen Sie zu Schritt 5 .
Haben Sie die Persönlichkeit, die Sie wollen	Überspringen Sie Schritt 5 bis Schritt 9 und gehen Sie zu Schritt 10 .

5. Nehmen Sie eine der folgenden Aktionen:

Wenn Sie konfigurieren...	Dann...
Ports auf einer UTA/UTA2-Karte	Gehe zu Schritt 6
Onboard UTA/UTA2-Ports	Überspringen Sie Schritt 6 und gehen Sie zu Schritt 7 .

6. Wenn sich der Adapter im Initiatormodus befindet und der UTA/UTA2-Port online ist, schalten Sie den UTA/UTA2-Port offline:

```
storage disable adapter adapter_name
```

Adapter im Zielmodus sind im Wartungsmodus automatisch offline.

7. Wenn die aktuelle Konfiguration nicht mit der gewünschten Verwendung übereinstimmt, ändern Sie die Konfiguration nach Bedarf:

```
ucaadmin modify -m fc|cna -t initiator|target <adapter_name>
```

- -m Ist der Personality-Modus, FC oder 10GbE UTA.
- -t Ist der Typ FC4, target Oder initiator.



Sie müssen den FC-Initiator für Bandlaufwerke und MetroCluster -Konfigurationen verwenden. Sie müssen das FC-Ziel für SAN-Clients verwenden.

8. Schalten Sie alle Zielports online, indem Sie den folgenden Befehl einmal für jeden Port eingeben:

```
storage enable adapter <adapter_name>
```

9. Verkabeln Sie den Port.

10. Wartungsmodus beenden:

```
halt
```

11. Starten Sie den Node im Boot-Menü:

```
boot_ontap menu
```

Was kommt als Nächstes?

- Wenn Sie ein Upgrade auf ein AFF A800-System durchführen, gehen Sie zu ["Weisen Sie Node2-Festplatten node4, Schritt 9, neu zu"](#).
- Für alle anderen System-Upgrades gehen Sie zu ["Weisen Sie Node2-Festplatten node4, Schritt 1, neu zu"](#).

Weisen Sie Node2-Festplatten node4 neu zu

Sie müssen die Festplatten, die zu node 2 gehörten, vor der Überprüfung der Installation von node 4 neu zuweisen.

Schritte

1. Überprüfen Sie, ob Knoten2 im Bootmenü angehalten wurde, und weisen Sie die Festplatten von Knoten2 Knoten4 neu zu:

```
boot_after_controller_replacement
```

Nach einer kurzen Verzögerung werden Sie aufgefordert, den Namen des Node einzugeben, der ersetzt wird. Wenn gemeinsam genutzte Festplatten vorhanden sind (auch Advanced Disk Partitioning (ADP) oder partitionierte Festplatten), werden Sie aufgefordert, den Node-Namen des HA-Partners einzugeben.

Diese Eingabeaufforderungen sind möglicherweise in den Konsolenmeldungen verborgen. Wenn Sie keinen Node-Namen eingeben oder einen falschen Namen eingeben, werden Sie aufgefordert, den Namen erneut einzugeben.

Erweitern Sie das Ausgabebeispiel der Konsole

```
LOADER-A> boot_ontap menu ...
*****
*
* Press Ctrl-C for Boot Menu. *
*
*****

.
.
Please choose one of the following:

(1) Normal Boot.
(2) Boot without /etc/rc.
(3) Change password.
(4) Clean configuration and initialize all disks.
(5) Maintenance mode boot.
(6) Update flash from backup config.
(7) Install new software first.
(8) Reboot node.
(9) Configure Advanced Drive Partitioning.
Selection (1-9)? 22/7
.
.
(boot_after_controller_replacement) Boot after controller upgrade
(9a) Unpartition all disks and remove
their ownership information.
(9b) Clean configuration and
initialize node with partitioned disks.
(9c) Clean configuration and
initialize node with whole disks.
(9d) Reboot the node.
(9e) Return to main boot menu.

Please choose one of the following:

(1) Normal Boot.
(2) Boot without /etc/rc.
(3) Change password.
(4) Clean configuration and initialize all disks.
(5) Maintenance mode boot.
(6) Update flash from backup config.
(7) Install new software first.
(8) Reboot node.
(9) Configure Advanced Drive Partitioning.
Selection (1-9)? boot_after_controller_replacement
```



```

.
This will replace all flash-based configuration with the last backup
to disks. Are you sure you want to continue?: yes
.
.
Controller Replacement: Provide name of the node you would like to
replace: <name of the node being replaced>
Controller Replacement: Provide High Availability partner of node1:
<nodename of the partner of the node being replaced>
Changing sysid of node <node being replaced> disks.
Fetched sanown old_owner_sysid = 536953334 and calculated old sys id
= 536953334
Partner sysid = 4294967295, owner sysid = 536953334
.
.
.
Terminated
<node reboots>
.
.
System rebooting...
.
Restoring env file from boot media...
copy_env_file:scenario = head upgrade
Successfully restored env file from boot media...
.
.
System rebooting...
.
.
.
WARNING: System ID mismatch. This usually occurs when replacing a
boot device or NVRAM cards!
Override system ID? {y|n} y
Login: ...

```

2. Wenn das System in eine Neustartschleife gerät und die Meldung `no disks found`, liegt dies daran, dass die Ports wieder auf den Zielmodus zurückgesetzt wurden und daher keine Festplatten erkannt werden können. Ausführen [Schritt 3](#) durch [Schritt 8](#) auf Knoten4, um dieses Problem zu beheben.
3. Drücken Sie während des AUTOBOOTS Strg-C, um den Knoten an der Eingabeaufforderung `Loader>` anzuhalten.
4. Wechseln Sie an der LOADER-Eingabeaufforderung in den Wartungsmodus:

```
boot_ontap maint
```

5. Zeigen Sie im Wartungsmodus alle zuvor festgelegten Initiator-Ports an, die sich jetzt im Ziel-Modus befinden:

```
ucadmin show
```

Ändern Sie die Ports zurück in den Initiatormodus:

```
ucadmin modify -m fc -t initiator -f adapter name
```

6. Vergewissern Sie sich, dass die Ports in den Initiatormodus geändert wurden:

```
ucadmin show
```

7. Beenden des Wartungsmodus:

```
halt
```



Wenn Sie ein Upgrade von einem System durchführen, das externe Festplatten unterstützt, auf ein System, das auch externe Festplatten unterstützt, gehen Sie zu [Schritt 8](#).

Wenn Sie ein Upgrade von einem System durchführen, das externe Festplatten auf ein System verwendet, das sowohl interne als auch externe Festplatten unterstützt, [Schritt 9z](#).
B. ein AFF A800-System, gehen Sie zu .

8. Starten Sie an der Loader-Eingabeaufforderung:

```
boot_ontap menu
```

Beim Booten erkennt der Node jetzt alle Festplatten, die zuvor ihm zugewiesen waren, und kann wie erwartet gebootet werden.

Wenn die Clusterknoten, die Sie ersetzen, die Root-Volume-Verschlüsselung verwenden, kann ONTAP die Volume-Informationen von den Festplatten nicht lesen. Stellen Sie die Schlüssel für das Root-Volume wieder her:

- a. Zurück zum speziellen Startmenü:

```
LOADER> boot_ontap menu
```

Please choose one of the following:

- (1) Normal Boot.
- (2) Boot without /etc/rc.
- (3) Change password.
- (4) Clean configuration and initialize all disks.
- (5) Maintenance mode boot.
- (6) Update flash from backup config.
- (7) Install new software first.
- (8) Reboot node.
- (9) Configure Advanced Drive Partitioning.
- (10) Set Onboard Key Manager recovery secrets.
- (11) Configure node for external key management.

Selection (1-11)? 10

a. Wählen Sie **(10) Set Onboard Key Manager Recovery Secrets**

b. Eingabe `y` An der folgenden Eingabeaufforderung:

```
This option must be used only in disaster recovery procedures. Are you sure?  
(y or n): y
```

c. Geben Sie an der Eingabeaufforderung die Passphrase für das Schlüsselmanagement ein.

d. Geben Sie bei Aufforderung die Backup-Daten ein.



Sie müssen die Passphrase und Sicherungsdaten im erhalten haben "[Bereiten Sie die Knoten für ein Upgrade vor](#)" Abschnitt dieses Verfahrens.

e. Nachdem das System wieder zum speziellen Startmenü gestartet wurde, führen Sie die Option **(1) Normal Boot** aus



Es könnte in dieser Phase ein Fehler auftreten. Falls ein Fehler auftritt, wiederholen Sie die Teilschritte in [Schritt 8](#) bis das System normal hochfährt.

9. Wenn Sie von einem System mit externen Festplatten auf ein System aktualisieren, das interne und externe Festplatten unterstützt (z. B. AFF A800 -Systeme), legen Sie das Knoten2-Aggregat als Root-Aggregat fest, um zu bestätigen, dass Knoten4 vom Root-Aggregat von Knoten2 bootet. Um das Root-Aggregat festzulegen, gehen Sie zum Boot-Menü auf Knoten4 und wählen Sie die Option 5 um in den Wartungsmodus zu wechseln.



Die folgenden Teilschritte müssen in der angegebenen Reihenfolge ausgeführt werden; andernfalls kann es zu einem Ausfall oder sogar zu Datenverlust kommen.

Mit dem folgenden Verfahren wird node4 vom Root-Aggregat von node2 gestartet:

a. Wechseln in den Wartungsmodus:

```
boot_ontap maint
```

- b. Überprüfen Sie die RAID-, Plex- und Prüfsummeninformationen für das node2 Aggregat:

```
aggr status -r
```

- c. Überprüfen Sie den Status des node2-Aggregats:

```
aggr status
```

- d. Bei Bedarf das node2 Aggregat online bringen:

```
aggr_online root_aggr_from_node2
```

- e. Verhindern Sie, dass das node4 aus dem ursprünglichen Root-Aggregat gebootet wird:

```
aggr offline root_aggr_on_node4
```

- f. Legen Sie das node2-Root-Aggregat als das neue Root-Aggregat für node4 fest:

```
aggr options aggr_from_node2 root
```

Weisen Sie Ports von node2 nach node4 zu

Sie müssen überprüfen, ob die physischen Ports auf node2 den physischen Ports auf node4 korrekt zugeordnet sind. Dadurch kann node4 nach dem Upgrade mit anderen Knoten im Cluster und mit dem Netzwerk kommunizieren.

Über diese Aufgabe

Siehe "[Quellen](#)" Verknüpfen mit *Hardware Universe*, um Informationen über die Ports auf den neuen Nodes zu erfassen. Die Informationen werden später in diesem Abschnitt verwendet.

die Softwarekonfiguration von node4 muss mit der physischen Konnektivität von node4 übereinstimmen und die IP-Konnektivität muss wiederhergestellt werden, bevor Sie das Upgrade fortsetzen.

Die Port-Einstellungen können je nach Modell der Nodes variieren. Sie müssen den Port des ursprünglichen Node und die LIF-Konfiguration mit dem kompatibel machen, was Sie planen, die Konfiguration des neuen Node zu verwenden. Dies liegt daran, dass der neue Node beim Booten die gleiche Konfiguration wiedergibt. Dies bedeutet, wenn Sie node4 booten, wird Data ONTAP versuchen, LIFs auf den gleichen Ports zu hosten, die in node2 verwendet wurden.

Wenn die physischen Ports auf node2 also nicht direkt den physischen Ports auf node4 zugeordnet werden, sind daher Änderungen der Software-Konfiguration erforderlich, um nach dem Booten die Cluster-, Management- und Netzwerkkonnektivität wiederherzustellen. Wenn die Cluster-Ports auf node2 zudem nicht direkt den Cluster-Ports auf node4 zugeordnet werden, wird node4 beim Neustart möglicherweise nicht automatisch dem Quorum beitreten, bis die Software-Konfiguration geändert wird, um die Cluster LIFs auf den korrekten physischen Ports zu hosten.

Schritte

1. Notieren Sie alle node2-Verkabelungsinformationen für node2, die Ports, Broadcast-Domänen und IPspaces in der Tabelle:

LIF	Anzahl an Knoten2-Ports	Node2-IPspaces	Node2 Broadcast-Domänen	Node4-Ports	Node4 IPspaces	Node4 Broadcast-Domänen
Cluster 1						
Cluster 2						
Cluster 3						
Cluster 4						
Node-Management						
Cluster-Management						
Daten 1						
Daten 2						
Daten 3						
Daten 4						
San						
Intercluster-Port						

2. Zeichnen Sie alle Kabelinformationen für node4, die Ports, Broadcast-Domänen und IPspaces in der Tabelle auf.
3. Führen Sie die folgenden Schritte aus, um zu überprüfen, ob es sich bei dem Setup um ein 2-Node-Cluster ohne Switches handelt:

- a. Legen Sie die Berechtigungsebene auf erweitert fest:

```
cluster::> set -privilege advanced
```

- b. Überprüfen Sie, ob es sich um ein 2-Node-Cluster ohne Switches handelt:

```
cluster::> network options switchless-cluster show
```

```
cluster::*> network options switchless-cluster show
Enable Switchless Cluster: false/true
```

+

Der Wert dieses Befehls muss mit dem physischen Status des Systems übereinstimmen.

- a. Zurück zur Administratorberechtigungsebene:

```
cluster::*> set -privilege admin
cluster::>
```

4. Gehen Sie folgendermaßen vor, um node4 in Quorum zu platzieren:

- a. Boot node4. Siehe ["installieren und booten sie node4"](#) Um den Node zu booten, wenn Sie dies noch nicht getan haben.
- b. Vergewissern Sie sich, dass sich die neuen Cluster-Ports in der Cluster Broadcast-Domäne befinden:

```
network port show -node node -port port -fields broadcast-domain
```

Das folgende Beispiel zeigt, dass Port „e0a“ in der Cluster-Domäne auf node4 ist:

```
cluster::> network port show -node node4 -port e0a -fields broadcast-  
domain  
node      port broadcast-domain  
-----  
node4     e0a  Cluster
```

- c. Wenn sich die Cluster-Ports nicht in der Cluster Broadcast-Domäne befinden, fügen Sie sie mit dem folgenden Befehl hinzu:

```
broadcast-domain add-ports -ip-space Cluster -broadcast-domain Cluster -ports  
node:port
```

- d. Fügen Sie die korrekten Ports zur Cluster Broadcast-Domäne hinzu:

```
network port modify -node -port -ip-space Cluster -mtu 9000
```

Dieses Beispiel fügt Cluster-Port „e1b“ auf node4 hinzu:

```
network port modify -node node4 -port e1b -ip-space Cluster -mtu 9000
```

- e. Migrieren Sie die Cluster-LIFs zu den neuen Ports, einmal für jede LIF:

```
network interface migrate -vserver Cluster -lif lif_name -source-node node4  
destination-node node4 -destination-port port_name
```

- f. Ändern Sie den Startport der Cluster-LIFs:

```
network interface modify -vserver Cluster -lif lif_name -home-port port_name
```

- g. Entfernen Sie die alten Ports aus der Cluster Broadcast-Domäne:

```
network port broadcast-domain remove-ports
```

Dieser Befehl entfernt Port „e0d“ auf node4:

```
network port broadcast-domain remove-ports -ip-space Cluster -broadcast-domain  
Cluster -ports node4:e0d
```

- a. Vergewissern Sie sich, dass node4 Quorum erneut verbunden hat:

```
cluster show -node node4 -fields health
```

5. Anpassen der Broadcast-Domänen, die Ihre Cluster-LIFs hosten, sowie Node-Management/clustermanagement-LIFs. Vergewissern Sie sich, dass jede Broadcast-Domäne die richtigen Ports enthält. Ein Port kann nicht zwischen Broadcast-Domänen verschoben werden, wenn er als Host oder Home für eine LIF ist, sodass Sie möglicherweise die LIFs migrieren und ändern müssen, wie in den folgenden Schritten dargestellt:

- a. Zeigen Sie den Startport einer logischen Schnittstelle an:

```
network interface show -fields home-node,home-port
```

- b. Zeigen Sie die Broadcast-Domäne an, die diesen Port enthält:

```
network port broadcast-domain show -ports node_name:port_name
```

- c. Ports aus Broadcast-Domänen hinzufügen oder entfernen:

```
network port broadcast-domain add-ports
network port broadcast-domain remove-ports
```

- d. Ändern Sie den Home-Port eines LIF:

```
network interface modify -vserver vservers -lif lif_name -home-port port_name
```

6. Passen Sie die Intercluster-Broadcast-Domänen an und migrieren Sie gegebenenfalls die Intercluster LIFs mithilfe derselben Befehle, die in dargestellt sind [Schritt 5](#).
7. Passen Sie alle anderen Broadcast-Domänen an und migrieren Sie die Daten-LIFs, falls erforderlich, mit denselben Befehlen in [Schritt 5](#).
8. Wenn in node4 keine Ports mehr vorhanden sind, löschen Sie diese wie folgt:

- a. Zugriff auf die erweiterte Berechtigungsebene auf beiden Nodes:

```
set -privilege advanced
```

- b. So löschen Sie die Ports:

```
network port delete -node node_name -port port_name
```

- c. Zurück zur Administratorebene:

```
set -privilege admin
```

9. Passen Sie alle LIF Failover-Gruppen an:

```
network interface modify -failover-group failover_group -failover-policy
failover_policy
```

Mit dem folgenden Befehl wird die Failover-Richtlinie auf festgelegt `broadcast-domain-wide` Und verwendet die Ports in Failover-Gruppe `fg1` Als Failover-Ziele für LIF `data1` Ein node4:

```
network interface modify -vserver node4 -lif data1 failover-policy broadcast-
domainwide -failover-group fg1
```

Siehe "[Quellen](#)" Weitere Informationen finden Sie unter *Netzwerkverwaltung* oder den Befehlen *ONTAP 9: Manual Page Reference* unter *Failover-Einstellungen auf LIF konfigurieren*.

10. Überprüfen Sie die Änderungen auf node4:

```
network port show -node node4
```

11. Jedes Cluster-LIF muss an Port 7700 zuhören. Vergewissern Sie sich, dass die Cluster-LIFs an Port 7700 zuhören:

```
::> network connections listening show -vserver Cluster
```

Port 7700, der auf Cluster-Ports hört, ist das erwartete Ergebnis, wie im folgenden Beispiel für ein Cluster mit zwei Nodes dargestellt:

```
Cluster::> network connections listening show -vserver Cluster
Vserver Name      Interface Name:Local Port      Protocol/Service
-----
Node: NodeA
Cluster           NodeA_clus1:7700              TCP/ctlopcp
Cluster           NodeA_clus2:7700              TCP/ctlopcp
Node: NodeB
Cluster           NodeB_clus1:7700              TCP/ctlopcp
Cluster           NodeB_clus2:7700              TCP/ctlopcp
4 entries were displayed.
```

12. Legen Sie für jede Cluster-LIF, die nicht an Port 7700 angehört, den Administrationsstatus der LIF auf fest down Und dann up:

```
::> net int modify -vserver Cluster -lif cluster-lif -status-admin down; net
int modify -vserver Cluster -lif cluster-lif -status-admin up
```

Wiederholen Sie Schritt 11, um zu überprüfen, ob die Cluster-LIF jetzt auf Port 7700 angehört.

Fügen Sie dem Quorum bei, wenn ein Node über einen anderen Satz an Netzwerkports verfügt

Der Node mit dem neuen Controller bootet und versucht zuerst, dem Cluster automatisch beizutreten. Wenn der neue Node jedoch einen anderen Satz an Netzwerkports aufweist, müssen Sie die folgenden Schritte durchführen, um zu bestätigen, dass der Node dem Quorum erfolgreich hinzugefügt wurde.

Über diese Aufgabe

Sie können diese Anweisungen für alle relevanten Knoten verwenden. Node3 wird in der folgenden Probe verwendet.

Schritte

1. Überprüfen Sie, ob sich die neuen Cluster-Ports in der Cluster Broadcast-Domäne befinden, indem Sie den folgenden Befehl eingeben und die Ausgabe überprüfen:

```
network port show -node node -port port -fields broadcast-domain
```


Das folgende Beispiel zeigt, dass sich der Port „e1a“ in der Cluster-Domäne auf node3 befindet:

```
cluster::> network port show -node node3 -port e1a -fields broadcast-  
domain  
node    port    broadcast-domain  
-----  
node3   e1a     Cluster
```

2. Fügen Sie die korrekten Ports zur Cluster Broadcast-Domäne hinzu, indem Sie den folgenden Befehl eingeben und die Ausgabe überprüfen:

```
network port modify -node -port -ipSPACE Cluster -mtu 9000
```

Dieses Beispiel fügt Cluster-Port „e1b“ auf Knoten3 hinzu:

```
network port modify -node node3 -port e1b -ipSPACE Cluster -mtu 9000
```

3. Migrieren Sie die Cluster-LIFs zu den neuen Ports, einmal für jede LIF, und verwenden Sie den folgenden Befehl:

```
network interface migrate -vserver Cluster -lif lif_name -source-node node3  
destination-node node3 -destination-port port_name
```

4. Ändern Sie den Home-Port der Cluster-LIFs wie folgt:

```
network interface modify -vserver Cluster -lif lif_name -home-port port_name
```

5. Wenn sich die Cluster-Ports nicht in der Cluster Broadcast-Domäne befinden, fügen Sie sie mit dem folgenden Befehl hinzu:

```
network port broadcast-domain add-ports -ipSPACE Cluster -broadcastdomain  
Cluster ports node:port
```

6. Entfernen Sie die alten Ports aus der Cluster Broadcast-Domäne. Sie können für jeden relevanten Knoten verwenden. Mit dem folgenden Befehl wird der Port „e0d“ auf node3 entfernt:

```
network port broadcast-domain remove-ports network port broadcast-domain  
remove-ports ipSPACE Cluster -broadcast-domain Cluster -ports node3:e0d
```

7. Vergewissern Sie sich, dass sich der Node erneut dem Quorum angeschlossen hat:

```
cluster show -node node3 -fields health
```

8. Passen Sie die Broadcast-Domänen an, die Ihre Cluster-LIFs und LIFs für das Node-Management/Cluster-Management hosten. Vergewissern Sie sich, dass jede Broadcast-Domäne die richtigen Ports enthält. Ein Port kann nicht zwischen Broadcast-Domänen verschoben werden, wenn er als Host oder Home für eine LIF ist, sodass Sie die LIFs möglicherweise wie folgt migrieren und ändern müssen:

- a. Zeigen Sie den Startport einer logischen Schnittstelle an:

```
network interface show -fields home-node,home-port
```

- b. Zeigen Sie die Broadcast-Domäne an, die diesen Port enthält:

```
network port broadcast-domain show -ports node_name:port_name
```

- c. Ports aus Broadcast-Domänen hinzufügen oder entfernen:

```
network port broadcast-domain add-ports network port broadcast-domain  
remove-port
```

- d. Ändern eines Startports einer LIF:

```
network interface modify -vserver vservice-name -lif lif_name -home-port  
port_name
```

Passen Sie die Intercluster-Broadcast-Domänen an und migrieren Sie gegebenenfalls die Intercluster LIFs. Die Daten-LIFs bleiben unverändert.

Überprüfen Sie die installation von node4

Nach der Installation und dem Booten von node4 müssen Sie überprüfen, ob node4 korrekt installiert ist, dass er Teil des Clusters ist und mit node3 kommunizieren kann.

Über diese Aufgabe

An diesem Punkt des Verfahrens wird der Vorgang angehalten, da node4 dem Quorum beitrifft.

Schritte

1. Vergewissern Sie sich, dass node4 dem Quorum beigetreten ist:

```
cluster show -node node4 -fields health
```

2. Vergewissern Sie sich, dass node4 Teil desselben Clusters wie node3 und des entsprechenden Clusters ist, indem Sie den folgenden Befehl eingeben:

```
cluster show
```

3. Überprüfen Sie den Status des Vorgangs, und überprüfen Sie, ob die Konfigurationsinformationen für node4 identisch sind mit node2:

```
system controller replace show-details
```

Wenn sich die Konfiguration für node4 unterscheidet, kann es zu einem späteren Zeitpunkt zu einer Systemunterbrechung kommen.

4. Überprüfen Sie, ob der ersetzte Controller für MetroCluster-Konfiguration und nicht im Switch-Over-Modus korrekt konfiguriert ist.



Zu diesem Zeitpunkt befindet sich die MetroCluster Konfiguration nicht im Normalzustand und es müssen möglicherweise Fehler behoben werden. Sehen ["Überprüfen Sie den Systemzustand der MetroCluster-Konfiguration"](#) .

VLANs, Schnittstellengruppen und Broadcast-Domänen auf node4 neu erstellen

Nachdem Sie bestätigt haben, dass node4 sich im Quorum befindet und mit node3 kommunizieren kann, müssen Sie die VLANs, Schnittstellengruppen und Broadcast-Domänen von node2 auf node4 neu erstellen. Sie müssen auch die node3-Ports zu den neu erstellten Broadcast-Domänen hinzufügen.

Über diese Aufgabe

Weitere Informationen zum Erstellen und Neuerstellen von VLANs, Schnittstellengruppen und Broadcast-Domänen finden Sie unter ["Quellen"](#) Und Link zu *Network Management*.

Schritte

1. Erstellen Sie die VLANs auf node4 mithilfe der node2-Informationen, die im aufgezeichnet wurden, neu ["Verschieben von Aggregaten und NAS-Daten-LIFs ohne Root-Wurzeln von Knoten 2 auf Knoten 3"](#) Abschnitt:

```
network port vlan create -node node4 -vlan vlan-names
```

2. Erstellen Sie die Schnittstellengruppen auf node4 mithilfe der node2-Informationen, die im aufgezeichnet wurden, neu ["Verschieben von Aggregaten und NAS-Daten-LIFs ohne Root-Wurzeln von Knoten 2 auf Knoten 3"](#) Abschnitt:

```
network port ifgrp create -node node4 -ifgrp port_ifgrp_names-distr-func
```

3. Erstellen Sie die Broadcast-Domänen auf node4 mithilfe der node2-Informationen, die im aufgezeichnet wurden, erneut ["Verschieben von Aggregaten und NAS-Daten-LIFs ohne Root-Wurzeln von Knoten 2 auf Knoten 3"](#) Abschnitt:

```
network port broadcast-domain create -ipspace Default -broadcast-domain  
broadcast_domain_names -mtu mtu_size -ports  
node_name:port_name,node_name:port_name
```

4. Fügen Sie die node4-Ports zu den neu erstellten Broadcast-Domänen hinzu:

```
network port broadcast-domain add-ports -broadcast-domain  
broadcast_domain_names -ports node_name:port_name,node_name:port_name
```

Wiederherstellen der Key-Manager-Konfiguration auf node4

Wenn Sie NetApp Aggregate Encryption (NAE) oder NetApp Volume Encryption (NVE) zur Verschlüsselung von Volumes auf dem System verwenden, das Sie aktualisieren, muss die Verschlüsselungskonfiguration mit den neuen Nodes synchronisiert werden. Wenn Sie den Schlüsselmanager nicht wiederherstellen, werden beim Verschieben der Node2-Aggregate mit ARL von node3 auf node4 verschlüsselte Volumes offline geschaltet.

Schritte

1. Führen Sie zum Synchronisieren der Verschlüsselungskonfiguration für Onboard Key Manager den folgenden Befehl an der Cluster-Eingabeaufforderung aus:

Für diese ONTAP-Version...	Befehl
ONTAP 9.6 oder 9.7	<code>security key-manager onboard sync</code>

Für diese ONTAP-Version...	Befehl
ONTAP 9.5	<code>security key-manager setup -node node_name</code>

2. Geben Sie die Cluster-weite Passphrase für das Onboard Key Manager ein.

Verschieben Sie Aggregate ohne Root-Root-Fehler und NAS-Daten-LIFs, die sich im Besitz von node2 befinden, von node3 auf node4

Nachdem Sie die node4-Installation überprüft haben und bevor Sie Aggregate von node3 auf node4 verschieben, müssen Sie die NAS-Daten-LIFs von node2 verschieben, die sich derzeit in node3 von node3 nach node4 befinden. Sie müssen außerdem überprüfen, ob die SAN-LIFs auf node4 vorhanden sind.

Über diese Aufgabe

Remote-LIFs verarbeiten den Datenverkehr zu SAN-LUNs während des Upgrades. Das Verschieben von SAN-LIFs ist für den Zustand des Clusters oder des Service während des Upgrades nicht erforderlich. SAN LIFs werden nicht verschoben, es sei denn, sie müssen neuen Ports zugeordnet werden. Sie überprüfen, ob die LIFs ordnungsgemäß sind und sich in den entsprechenden Ports befinden, nachdem Sie node4 in den Online-Modus versetzt haben.

Schritte

1. Wiederaufnahme des Betriebs der Versetzung:

```
system controller replace resume
```

Das System führt die folgenden Aufgaben aus:

- Cluster-Quorum-Prüfung
- System-ID-Prüfung
- Prüfung der Bildversion
- Überprüfung der Zielplattform
- Prüfung der Netzwerkanachhabilität

Der Vorgang unterbricht in dieser Phase in der Überprüfung der Netzwerknachprüfbarkeit.

2. Überprüfen Sie manuell, ob das Netzwerk und alle VLANs, Schnittstellengruppen und Broadcast-Domänen korrekt konfiguriert wurden.

3. Wiederaufnahme des Betriebs der Versetzung:

```
system controller replace resume
```

To complete the "Network Reachability" phase, ONTAP network configuration must be manually adjusted to match the new physical network configuration of the hardware. This includes assigning network ports to the correct broadcast domains, creating any required ifgrps and VLANs, and modifying the home-port parameter of network interfaces to the appropriate ports. Refer to the "Using aggregate relocation to upgrade controller hardware on a pair of nodes running ONTAP 9.x" documentation, Stages 3 and 5. Have all of these steps been manually completed? [y/n]

4. Eingabe `y` Um fortzufahren.

5. Das System führt folgende Prüfungen durch:

- Cluster-Zustandsprüfung
- LIF-Statusüberprüfung für Cluster

Nach Durchführung dieser Prüfungen werden die nicht-Root-Aggregate und NAS-Daten-LIFs, die sich im Besitz von node2 befinden, an den neuen Controller node4 verschoben. Das System hält an, sobald die Ressourcenverlagerung abgeschlossen ist.

6. Überprüfen Sie den Status der Aggregatverschiebung und der LIF-Verschiebung von NAS-Daten:

```
system controller replace show-details
```

7. Überprüfen Sie manuell, ob die nicht-Root-Aggregate und NAS-Daten-LIFs erfolgreich in node4 verschoben wurden.

Falls Aggregate nicht verschoben oder ein Veto eingesetzt werden kann, müssen sie die Aggregate manuell verschieben oder – falls erforderlich – entweder die Vetos oder Zielprüfungen überschreiben. Siehe Abschnitt "[Verschiebung ausgefallener oder Vetos von Aggregaten](#)" Finden Sie weitere Informationen.

8. Vergewissern Sie sich, dass sich die SAN-LIFs auf den richtigen Ports auf node4 befinden, indem Sie die folgenden Teilschritte ausführen:

a. Geben Sie den folgenden Befehl ein und überprüfen Sie die Ausgabe:

```
network interface show -data-protocol iscsi|fc -home-node node4
```

Das System gibt die Ausgabe wie im folgenden Beispiel zurück:

```
cluster::> net int show -data-protocol iscsi|fc -home-node node3
```

	Logical	Status	Network	Current	Current	Is
Vserver	Interface	Admin/Oper	Address/Mask	Node	Port	Home

vs0						
	a0a	up/down	10.63.0.53/24	node3	a0a	true
	data1	up/up	10.63.0.50/18	node3	e0c	true
	rads1	up/up	10.63.0.51/18	node3	e1a	true
	rads2	up/down	10.63.0.52/24	node3	e1b	true
vs1						
	lif1	up/up	172.17.176.120/24	node3	e0c	true
	lif2	up/up	172.17.176.121/24	node3	e1a	true

- b. Wenn node4 eine SAN-LIFs oder Gruppen von SAN-LIFs hat, die sich auf einem Port befinden, der nicht in node2 vorhanden war oder einem anderen Port zugeordnet werden muss, verschieben Sie sie in einen entsprechenden Port auf node4, indem Sie die folgenden Teilschritte ausführen:

- i. Stellen Sie den LIF-Status auf „down“ ein, indem Sie den folgenden Befehl eingeben:

```
network interface modify -vserver vs0 -lif lif1 -status
-admin down
```

- ii. Entfernen Sie das LIF aus dem Portset:

```
portset remove -vserver vs0 -portset portset1 -port-name
lif1
```

- iii. Geben Sie einen der folgenden Befehle ein:

- Verschieben Sie ein einzelnes LIF durch Eingabe des folgenden Befehls:

```
network interface modify -vserver vs0 -lif lif1 -home
-port new_home_port
```

- Verschieben Sie alle LIFs auf einem einzelnen nicht vorhandenen oder falschen Port in einen neuen Port, indem Sie den folgenden Befehl eingeben:

```
network interface modify {-home-port port_on_node1 -home-node node1
-role data} -home-port new_home_port_on_node3
```

- Fügen Sie die LIFs wieder dem Portset hinzu:

```
portset add -vserver vs0 -portset portset1 -port-name
lif1
```



Sie müssen bestätigen, dass Sie SAN LIFs zu einem Port mit der gleichen Verbindungsgeschwindigkeit wie der ursprüngliche Port verschieben.

- a. Ändern Sie den Status aller LIFs in `up`. Damit die LIFs Datenverkehr auf dem Node akzeptieren und senden können, indem Sie den folgenden Befehl eingeben:

```
network interface modify -home-port port_name -home-node node4 -lif data  
-statusadmin up
```

- b. Geben Sie den folgenden Befehl ein und überprüfen Sie seine Ausgabe, um zu überprüfen, ob LIFs zu den richtigen Ports verschoben wurden und ob die LIFs den Status von aufweisen `up`. Wenn Sie auf einem der beiden Nodes den folgenden Befehl eingeben und die Ausgabe überprüfen:

```
network interface show -home-node <node4> -role data
```

- c. Wenn irgendwelche LIFs ausgefallen sind, setzen Sie den Administratorstatus der LIFs auf `up`. Geben Sie den folgenden Befehl ein, einmal für jede LIF:

```
network interface modify -vserver vserver_name -lif lif_name -status-admin  
up
```

9. Setzen Sie den Vorgang fort, um das System zur Durchführung der erforderlichen Nachprüfungen zu auffordern:

```
system controller replace resume
```

Das System führt die folgenden Nachprüfungen durch:

- Cluster-Quorum-Prüfung
- Cluster-Zustandsprüfung
- Aggregatrekonstruktion
- Aggregatstatus-Prüfung
- Überprüfung des Festplattenstatus
- LIF-Statusüberprüfung für Cluster

Phase 6: Schließen Sie das Upgrade ab

Authentifizierungsmanagement mit KMIP-Servern

Von ONTAP 9.5 bis 9.7 können KMIP-Server (Key Management Interoperability Protocol) Authentifizierungsschlüssel managen.

Schritte

1. Hinzufügen eines neuen Controllers:

```
security key-manager setup -node new_controller_name
```

2. Fügen Sie den Schlüsselmanager hinzu:

```
security key-manager -add key_management_server_ip_address
```

3. Vergewissern Sie sich, dass die Verschlüsselungsmanagement-Server konfiguriert und für alle Nodes im Cluster verfügbar sind:

```
security key-manager show -status
```

4. Stellen Sie die Authentifizierungsschlüssel von allen verknüpften Verschlüsselungsmanagementservern auf den neuen Knoten wieder her:

```
security key-manager restore -node new_controller_name
```

Vergewissern Sie sich, dass die neuen Controller ordnungsgemäß eingerichtet sind

Um das korrekte Setup zu bestätigen, müssen Sie das HA-Paar aktivieren. Sie müssen außerdem überprüfen, dass Node3 und node4 auf den Storage der jeweils anderen Person zugreifen können und dass keine der logischen Datenschnittstellen zu anderen Nodes im Cluster vorhanden sind. Darüber hinaus müssen Sie bestätigen, dass Node3 zu Aggregaten node1 gehört und dass node4 die Aggregate von node2 besitzt und dass die Volumes für beide Nodes online sind.

Schritte

1. Nach den Tests nach der Prüfung von „node2“ werden das Storage Failover und das Cluster HA-Paar für den node2 Cluster aktiviert. Nach Abschluss des Vorgangs werden beide Nodes als abgeschlossen angezeigt, und das System führt einige Bereinigungsvorgänge aus.
2. Vergewissern Sie sich, dass Storage-Failover aktiviert ist:

```
storage failover show
```

Im folgenden Beispiel wird die Ausgabe des Befehls angezeigt, wenn ein Storage Failover aktiviert ist:

```
cluster::> storage failover show
```

Takeover			
Node	Partner	Possible	State Description
node3	node4	true	Connected to node4
node4	node3	true	Connected to node3

3. Überprüfen Sie, ob node3 und node4 zum selben Cluster gehören, indem Sie den folgenden Befehl verwenden und die Ausgabe überprüfen:

```
cluster show
```

4. Stellen Sie sicher, dass node3 und node4 über den folgenden Befehl und eine Analyse der Ausgabe auf den Storage des jeweils anderen zugreifen können:

```
storage failover show -fields local-missing-disks, partner-missing-disks
```

5. Vergewissern Sie sich, dass weder node3 noch node4 Eigentümer von Daten-LIFs sind, die im Besitz anderer Nodes im Cluster sind. Verwenden Sie dazu den folgenden Befehl und prüfen Sie die Ausgabe:

```
network interface show
```

Wenn keine der Knoten „Node3“ oder „node4“ Daten-LIFs besitzt, die sich im Besitz anderer Nodes im Cluster befinden, setzen Sie die Daten-LIFs auf ihren Home-Eigentümer zurück:


```
network interface revert
```

6. Überprüfen Sie, ob node3 die Aggregate von node1 besitzt und dass node4 die Aggregate von node2 besitzt:

```
storage aggregate show -owner-name <node3>
```

```
storage aggregate show -owner-name <node4>
```

7. Legen Sie fest, ob Volumes offline sind:

```
volume show -node <node3> -state offline
```

```
volume show -node <node4> -state offline
```

8. Wenn Volumes offline sind, vergleichen Sie sie mit der Liste der Offline-Volumes, die Sie im Abschnitt erfasst haben "[Bereiten Sie die Knoten für ein Upgrade vor](#)", Und Online-Laufwerk eines der Offline-Volumes, nach Bedarf, durch Verwendung des folgenden Befehls, einmal für jedes Volume:

```
volume online -vserver <vserver_name> -volume <volume_name>
```

9. Installieren Sie neue Lizenzen für die neuen Nodes mithilfe des folgenden Befehls für jeden Node:

```
system license add -license-code <license_code,license_code,license_code...>
```

Der Lizenzcode-Parameter akzeptiert eine Liste von 28 alphabetischen Zeichenschlüsseln für Großbuchstaben. Sie können jeweils eine Lizenz hinzufügen, oder Sie können mehrere Lizenzen gleichzeitig hinzufügen, indem Sie jeden Lizenzschlüssel durch ein Komma trennen.

10. Entfernen Sie alle alten Lizenzen mithilfe eines der folgenden Befehle von den ursprünglichen Nodes:

```
system license clean-up -unused -expired
```

```
system license delete -serial-number <node_serial_number> -package  
<licensable_package>
```

- Alle abgelaufenen Lizenzen löschen:

```
system license clean-up -expired
```

- Alle nicht verwendeten Lizenzen löschen:

```
system license clean-up -unused
```

- Löschen Sie eine bestimmte Lizenz von einem Cluster mit den folgenden Befehlen auf den Nodes:

```
system license delete -serial-number <node1_serial_number> -package *
```

```
system license delete -serial-number <node2_serial_number> -package *
```

Die folgende Ausgabe wird angezeigt:

```
Warning: The following licenses will be removed:
<list of each installed package>
Do you want to continue? {y|n}: y
```

Eingabe `y` Um alle Pakete zu entfernen.

11. Überprüfen Sie, ob die Lizenzen korrekt installiert sind, indem Sie den folgenden Befehl verwenden und die Ausgabe überprüfen:

```
system license show
```

Sie können die Ausgabe mit der im Abschnitt erfassten Ausgabe vergleichen "[Bereiten Sie die Knoten für ein Upgrade vor](#)".

12. Wenn in der Konfiguration selbstverschlüsselnde Laufwerke verwendet werden und Sie die folgende Einstellung vorgenommen haben: `kmip.init.maxwait` Variable zu `off` (zum Beispiel in "[Installieren und starten Sie Node4, Schritt 22](#)"), müssen Sie die Variable löschen:

```
set diag; systemshell -node node_name -command sudo kenv -u -p
kmip.init.maxwait
```

13. Konfigurieren Sie die SPs mit dem folgenden Befehl auf beiden Knoten:

```
system service-processor network modify -node node_name
```

Siehe "[Quellen](#)" Link zur *Systemverwaltungsreferenz* für Informationen zu den SPs und den Befehlen *ONTAP 9: Manual Page Reference* für detaillierte Informationen zum `system service-processor network modify` Befehl.

14. Informationen zum Einrichten eines Clusters ohne Switches auf den neuen Nodes finden Sie unter "[Quellen](#)" Um eine Verbindung zur NetApp Support Site_ zu erhalten, befolgen Sie die Anweisungen unter „Wechsel zu einem 2-Node-Cluster ohne Switch_“.

Nachdem Sie fertig sind

Wenn die Speicherverschlüsselung auf node3 und node4 aktiviert ist, füllen Sie den Abschnitt aus "[Richten Sie Storage Encryption auf dem neuen Controller-Modul ein](#)". Andernfalls füllen Sie den Abschnitt aus "[Ausmustern des alten Systems](#)".

Richten Sie Storage Encryption auf dem neuen Controller-Modul ein

Wenn der ersetzte Controller oder der HA-Partner des neuen Controllers Storage Encryption verwendet, müssen Sie das neue Controller-Modul für Storage Encryption konfigurieren, einschließlich der Installation von SSL-Zertifikaten und der Einrichtung von Key Management-Servern.

Über diese Aufgabe

Dieses Verfahren umfasst Schritte, die auf dem neuen Controller-Modul ausgeführt werden. Sie müssen den Befehl auf dem richtigen Node eingeben.

Schritte

1. Vergewissern Sie sich, dass die Verschlüsselungsmanagement-Server weiterhin verfügbar sind, deren Status und ihre Authentifizierungsdaten folgendermaßen sind:

```
security key-manager show -status
```

```
security key-manager query
```

2. Fügen Sie die im vorherigen Schritt aufgeführten Verschlüsselungsmanagement-Server der Liste des zentralen Management-Servers im neuen Controller hinzu.

- a. Fügen Sie den Schlüsselverwaltungsserver hinzu:

```
security key-manager -add key_management_server_ip_address
```

- b. Wiederholen Sie den vorherigen Schritt für jeden aufgeführten Key Management Server. Sie können bis zu vier Verschlüsselungsmanagement-Server verknüpfen.
- c. Überprüfen Sie, ob die Verschlüsselungsmanagementserver erfolgreich hinzugefügt wurden:

```
security key-manager show
```

3. Führen Sie auf dem neuen Controller-Modul den Setup-Assistenten für das Verschlüsselungsmanagement aus, um die wichtigsten Management-Server einzurichten und zu installieren.

Sie müssen dieselben Key Management-Server installieren, die auf dem vorhandenen Controller-Modul installiert sind.

- a. Starten Sie den Setup-Assistenten für den Schlüsselmanagementserver auf dem neuen Knoten:

```
security key-manager setup -node new_controller_name
```

- b. Führen Sie die Schritte im Assistenten zum Konfigurieren von Verschlüsselungsmanagementservern durch.

4. Stellen Sie Authentifizierungsschlüssel von allen verknüpften Verschlüsselungsmanagementservern mit dem neuen Knoten wieder her:

```
security key-manager restore -node new_controller_name
```

Einrichtung von NetApp Volume oder Aggregate Encryption auf dem neuen Controller-Modul

Wenn der ersetzte Controller oder der HA-Partner (High Availability, Hochverfügbarkeit) des neuen Controllers NetApp Volume Encryption (NVE) oder NetApp Aggregate Encryption (NAE) verwendet, muss das neue Controller-Modul für NVE oder NAE konfiguriert werden.

Über diese Aufgabe

Dieses Verfahren umfasst Schritte, die auf dem neuen Controller-Modul ausgeführt werden. Sie müssen den Befehl auf dem richtigen Node eingeben.

ONTAP 9.6 und 9.7

Konfigurieren Sie NVE oder NAE auf Controllern mit ONTAP 9.6 oder 9.7

Schritte

1. Vergewissern Sie sich, dass die Verschlüsselungsmanagement-Server weiterhin verfügbar sind, deren Status und ihre Authentifizierungsdaten folgendermaßen sind:

```
security key-manager key query -node node
```

2. Fügen Sie die im vorherigen Schritt aufgeführten Verschlüsselungsmanagement-Server der Liste des zentralen Management-Servers des neuen Controllers hinzu:

- a. Fügen Sie den Schlüsselverwaltungsserver hinzu:

```
security key-manager -add key_management_server_ip_address
```

- b. Wiederholen Sie den vorherigen Schritt für jeden aufgeführten Key Management Server.

Sie können bis zu vier Verschlüsselungsmanagement-Server verknüpfen.

- c. Überprüfen Sie, ob die Verschlüsselungsmanagementserver erfolgreich hinzugefügt wurden:

```
security key-manager show
```

3. Führen Sie auf dem neuen Controller-Modul den Setup-Assistenten für das Verschlüsselungsmanagement aus, um die wichtigsten Management-Server einzurichten und zu installieren.

Sie müssen dieselben Key Management-Server installieren, die auf dem vorhandenen Controller-Modul installiert sind.

- a. Starten Sie den Setup-Assistenten für den Schlüsselmanagementserver auf dem neuen Knoten:

```
security key-manager setup -node new_controller_name
```

- b. Führen Sie die Schritte im Assistenten zum Konfigurieren von Verschlüsselungsmanagementservern durch.

4. Stellen Sie Authentifizierungsschlüssel von allen verknüpften Verschlüsselungsmanagementservern mit dem neuen Knoten wieder her.

- Authentifizierung für externen Schlüsselmanager wiederherstellen:

```
security key-manager external restore
```

Für diesen Befehl ist die OKM-Passphrase (Onboard Key Manager) erforderlich.

Weitere Informationen finden Sie im Knowledge Base-Artikel ["So stellen Sie die Konfiguration des externen Schlüsselmanager-Servers aus dem ONTAP-Startmenü wieder her"](#).

- Authentifizierung für den OKM wiederherstellen:

```
security key-manager onboard sync
```

ONTAP 9.5

NVE oder NAE auf Controllern mit ONTAP 9.5 konfigurieren

Schritte

1. Vergewissern Sie sich, dass die Verschlüsselungsmanagement-Server weiterhin verfügbar sind, deren Status und ihre Authentifizierungsdaten folgendermaßen sind:

```
security key-manager key show
```

2. Fügen Sie die im vorherigen Schritt aufgeführten Verschlüsselungsmanagement-Server der Liste des zentralen Management-Servers des neuen Controllers hinzu:

- a. Fügen Sie den Schlüsselverwaltungsserver hinzu:

```
security key-manager -add key_management_server_ip_address
```

- b. Wiederholen Sie den vorherigen Schritt für jeden aufgeführten Key Management Server.

Sie können bis zu vier Verschlüsselungsmanagement-Server verknüpfen.

- c. Überprüfen Sie, ob die Verschlüsselungsmanagementserver erfolgreich hinzugefügt wurden:

```
security key-manager show
```

3. Führen Sie auf dem neuen Controller-Modul den Setup-Assistenten für das Verschlüsselungsmanagement aus, um die wichtigsten Management-Server einzurichten und zu installieren.

Sie müssen dieselben Key Management-Server installieren, die auf dem vorhandenen Controller-Modul installiert sind.

- a. Starten Sie den Setup-Assistenten für den Schlüsselmanagementserver auf dem neuen Knoten:

```
security key-manager setup -node new_controller_name
```

- b. Führen Sie die Schritte im Assistenten zum Konfigurieren von Verschlüsselungsmanagementservern durch.

4. Stellen Sie Authentifizierungsschlüssel von allen verknüpften Verschlüsselungsmanagementservern mit dem neuen Knoten wieder her.

- Authentifizierung für externen Schlüsselmanager wiederherstellen:

```
security key-manager external restore
```

Für diesen Befehl ist die OKM-Passphrase (Onboard Key Manager) erforderlich.

Weitere Informationen finden Sie im Knowledge Base-Artikel ["So stellen Sie die Konfiguration des externen Schlüsselmanager-Servers aus dem ONTAP-Startmenü wieder her"](#).

- Restore-Authentifizierung für OKM:

```
security key-manager setup -node node_name
```

Nachdem Sie fertig sind

Überprüfen Sie, ob Volumes offline geschaltet wurden, da Authentifizierungsschlüssel nicht verfügbar waren oder externe Schlüsselverwaltungsserver nicht erreicht werden konnten. Bringen Sie diese Volumes wieder online, indem Sie die `volume online` Befehl.

Ausmustern des alten Systems

Nach dem Upgrade kann das alte System über die NetApp Support Site außer Betrieb gesetzt werden. Die Deaktivierung des Systems sagt NetApp, dass das System nicht mehr in Betrieb ist und dass es aus Support-Datenbanken entfernt wird.

Schritte

1. Siehe "[Quellen](#)" Um auf die *NetApp Support Site* zu verlinken und sich anzumelden.
2. Wählen Sie im Menü die Option **Produkte > Meine Produkte**.
3. Wählen Sie auf der Seite **installierte Systeme anzeigen** die **Auswahlkriterien** aus, mit denen Sie Informationen über Ihr System anzeigen möchten.

Sie können eine der folgenden Optionen wählen, um Ihr System zu finden:

- Seriennummer (auf der Rückseite des Geräts)
- Seriennummern für „My Location“

4. Wählen Sie **Los!**

In einer Tabelle werden Cluster-Informationen, einschließlich der Seriennummern, angezeigt.

5. Suchen Sie den Cluster in der Tabelle und wählen Sie im Dropdown-Menü Product Tool Set die Option **Decommission this System** aus.

Setzen Sie den SnapMirror Betrieb fort

Sie können SnapMirror Transfers, die vor dem Upgrade stillgelegt wurden, fortsetzen und die SnapMirror Beziehungen fortsetzen. Die Updates sind nach Abschluss des Upgrades im Zeitplan.

Schritte

1. Überprüfen Sie den SnapMirror Status auf dem Ziel:

```
snapmirror show
```

2. Wiederaufnahme der SnapMirror Beziehung:

```
snapmirror resume -destination-vserver vserver_name
```

Fehlerbehebung

Fehler bei der Aggregatverschiebung

Bei der Aggregatverschiebung (ARL, Aggregate Relocation) fallen während des

Upgrades möglicherweise an verschiedenen Punkten aus.

Prüfen Sie, ob Aggregate Relocation Failure vorhanden sind

Während des Verfahrens kann ARL in Phase 2, Phase 3 oder Phase 5 fehlschlagen.

Schritte

1. Geben Sie den folgenden Befehl ein und überprüfen Sie die Ausgabe:

```
storage aggregate relocation show
```

Der `storage aggregate relocation show` Befehl zeigt Ihnen, welche Aggregate erfolgreich umgezogen wurden und welche nicht, zusammen mit den Ursachen des Ausfalls.

2. Überprüfen Sie die Konsole auf EMS-Nachrichten.
3. Führen Sie eine der folgenden Aktionen durch:
 - Führen Sie die entsprechenden Korrekturmaßnahmen durch, je nach der Ausgabe des `storage aggregate relocation show` Befehl und Ausgabe der EMS-Nachricht.
 - Erzwingen Sie das Verlagern des Aggregats oder der Aggregate mit dem `override-veto` oder die Option `override-destination-checks` Option des `storage aggregate relocation start` Befehl.

Ausführliche Informationen zum `storage aggregate relocation start`, `override-veto`, und `override-destination-checks` Optionen finden Sie unter "[Quellen](#)" Link zu den Befehlen *ONTAP 9: Manual Page Reference*.

Aggregate, die ursprünglich auf node1 waren, gehören node4 nach Abschluss des Upgrades

Beim Abschluss des Upgrade-Verfahrens sollte die Knoten3 der neue Home-Node von Aggregaten sein, die ursprünglich als Home-Node die Knoten1 hatten. Sie können sie nach dem Upgrade verschieben.

Über diese Aufgabe

Unter den folgenden Umständen kann es nicht gelingen, Aggregate ordnungsgemäß zu verschieben und Node 1 als Home Node anstelle von Knoten3 zu verwenden:

- In Phase 3, wenn Aggregate von node2 auf node3 verschoben werden. Einige der verlagerten Aggregate haben die Nr. 1 als Home-Node. Ein solches Aggregat könnte zum Beispiel „aggr_Node_1“ heißen. Wenn die Verlagerung von aggr_Node_1 während Phase 3 fehlschlägt und eine Verlagerung nicht erzwungen werden kann, dann wird das Aggregat auf node2 zurückgelassen.
- Nach Stufe 4, wenn node2 durch node4 ersetzt wird. Wenn node2 ersetzt wird, kommt aggr_Node_1 mit node4 als Home-Node statt node3 online.

Sie können das falsche Eigentümerproblem nach Phase 6 beheben, wenn ein Storage-Failover aktiviert wurde, indem Sie die folgenden Schritte durchführen:

Schritte

1. Geben Sie den folgenden Befehl ein, um eine Liste der Aggregate zu erhalten:

```
storage aggregate show -nodes node4 -is-home true
```

Informationen zur Identifizierung von Aggregaten, die nicht korrekt verschoben wurden, finden Sie in der

Liste der Aggregate mit dem Home-Inhaber von node1, die Sie im Abschnitt erhalten haben ["Bereiten Sie die Knoten für ein Upgrade vor"](#) Und vergleichen Sie ihn mit der Ausgabe des obigen Befehls.

2. Vergleichen Sie die Ausgabe von Schritt 1 mit der Ausgabe, die Sie für Knoten 1 im Abschnitt aufgenommen haben ["Bereiten Sie die Knoten für ein Upgrade vor"](#) Und beachten Sie alle Aggregate, die nicht korrekt verschoben wurden.
3. Verschiebung der Aggregate links auf node4:

```
storage aggregate relocation start -node node4 -aggr aggr_node_1 -destination node3
```

Verwenden Sie das nicht `-ndo-controller-upgrade` Parameter während dieser Verschiebung.

4. Vergewissern Sie sich, dass node3 jetzt der Haupteigentümer der Aggregate ist:

```
storage aggregate show -aggregate aggr1,aggr2,aggr3... -fields home-name
```

aggr1,aggr2,aggr3... Ist die Liste der Aggregate, die node1 als ursprünglichen Besitzer hatten.

Aggregate, die nicht über Node3 als Hausbesitzer verfügen, können mit dem gleichen Relocation-Befehl in auf node3 verschoben werden [Schritt 3](#).

Neustarts, Panikspiele oder Energiezyklen

Das System kann in verschiedenen Phasen des Upgrades abstürzt, z. B. neu gebootet, in Panik geraten oder aus- und wieder eingeschaltet werden.

Die Lösung dieser Probleme hängt davon ab, wann sie auftreten.

Neustarts, Panikzugänge oder Energiezyklen während der Vorprüfphase

Node1 oder node2 stürzt vor der Pre-Check-Phase ab, während das HA-Paar noch aktiviert ist

Wenn node1 oder node2 vor der Pre-Check-Phase abstürzt, wurden noch keine Aggregate verschoben und die HA-Paar-Konfiguration ist noch aktiviert.

Über diese Aufgabe

Takeover und Giveback können normal fortgesetzt werden.

Schritte

1. Überprüfen Sie die Konsole auf EMS-Meldungen, die das System möglicherweise ausgegeben hat, und ergreifen Sie die empfohlenen Korrekturmaßnahmen.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Neustarts, Panikzugänge oder Energiezyklen während der ersten Ressourcenfreigabephase

Node1 stürzt während der ersten Resource-Release-Phase ab, während das HA-Paar noch aktiviert ist

Einige oder alle Aggregate wurden von node1 in node2 verschoben und das HA-Paar ist noch aktiviert. Node2 übernimmt das Root-Volume von node1 und alle nicht-Root-Aggregate, die nicht verschoben wurden.

Über diese Aufgabe

Eigentum an Aggregaten, die verschoben wurden, sehen genauso aus wie das Eigentum von nicht-Root-Aggregaten, die übernommen wurden, da sich der Home-Eigentümer nicht geändert hat.

Wenn node1 in den eintritt `waiting for giveback` Status, node2 gibt alle node1 nicht-Root-Aggregate zurück.

Schritte

1. Nachdem node1 gestartet wurde, sind alle nicht-Root-Aggregate von node1 zurück in node1 verschoben. Sie müssen eine manuelle Aggregatverschiebung der Aggregate von node1 nach node2 durchführen:
`storage aggregate relocation start -node node1 -destination node2 -aggregate -list * -ndocontroller-upgrade true`
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node1 stürzt während der ersten Ressourcen-Release-Phase ab, während das HA-Paar deaktiviert ist

Node2 übernimmt nicht, aber es stellt immer noch Daten aus allen nicht-Root-Aggregaten bereit.

Schritte

1. Knoten 1 aufbring.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node2 schlägt während der ersten Phase der Ressourcenfreigabe fehl, während das HA-Paar noch aktiviert ist

Node1 hat einige oder alle seine Aggregate in node2 verschoben. Das HA-Paar ist aktiviert.

Über diese Aufgabe

Node1 übernimmt alle node2 Aggregate sowie jedes seiner eigenen Aggregate, die auf node2 verschoben wurden. Beim Booten von node2 wird die Aggregatverschiebung automatisch abgeschlossen.

Schritte

1. Knoten 2 aufbring.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node2 stürzt während der ersten Resource-Release-Phase ab und nachdem HA-Paar deaktiviert ist

Node1 übernimmt nicht.

Schritte

1. Knoten 2 aufbring.

Ein Client-Ausfall tritt für alle Aggregate auf, während node2 gestartet wird.
2. Fahren Sie mit dem verbleibenden Upgrade des Node-Paars fort.

Startet während der ersten Verifikationsphase neu, erzeugt eine Panik oder schaltet die Stromversorgung aus

Node2 stürzt in der ersten Überprüfungsphase ab, wobei das HA-Paar deaktiviert ist

Node3 übernimmt nach einem Absturz nach einem node2 nicht, da das HA-Paar bereits deaktiviert ist.

Schritte

1. Knoten 2 aufbring.

Ein Client-Ausfall tritt für alle Aggregate auf, während node2 gestartet wird.

2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node3 stürzt in der ersten Verifizierungsphase ab, wobei das HA-Paar deaktiviert ist

Node2 übernimmt nicht, aber es stellt immer noch Daten aus allen nicht-Root-Aggregaten bereit.

Schritte

1. Knoten 3 aufbring.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Neustarts, Panikzucken oder Energiezyklen während der ersten Ressourcen-Wiederholen-Phase

Knoten 2 stürzt während der ersten Ressourcen-Wiederholen Phase während der Aggregat-Verschiebung ab

Node2 hat einige oder alle seine Aggregate von node1 in node3 verschoben. Node3 stellt Daten von Aggregaten bereit, die verlagert wurden. Das HA-Paar ist deaktiviert und somit gibt es keine Übernahme.

Über diese Aufgabe

Es gibt einen Client-Ausfall für Aggregate, die nicht verschoben wurden. Beim Booten von node2 werden die Aggregate von node1 auf node3 verschoben.

Schritte

1. Knoten 2 aufbring.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node3 stürzt während der ersten Phase zur Ressourcenrückgewinnung während der Aggregatverschiebung ab

Falls node3 abstürzt, während node2 Aggregate zu node3 verschoben wird, wird die Aufgabe nach dem Booten von node3 fortgesetzt.

Über diese Aufgabe

Node2 dient weiterhin verbleibenden Aggregaten, doch Aggregate, die bereits in Knoten 3 verlagert wurden, begegnen ein Client-Ausfall, während node3 gebootet wird.

Schritte

1. Knoten 3 aufbring.
2. Führen Sie das Controller-Upgrade fort.

Neustarts, Panikspiele oder Energiezyklen während der Nachprüfphase

Node2 oder node3 stürzt während der Post-Check-Phase ab

Das HA-Paar ist deaktiviert, damit dies keine Übernahme ist. Es gibt einen Client-Ausfall für Aggregate, die zum neu gebooteten Node gehören.

Schritte

1. Bringen Sie den Node hoch.

2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Neustarts, Panikzucken oder Energiezyklen während der zweiten Ressourcenfreigabephase

Node3 stürzt während der zweiten Resource-Release-Phase ab

Wenn node3 abstürzt, während node2 Aggregate verschoben, wird die Aufgabe nach dem Booten von node3 fortgesetzt.

Über diese Aufgabe

Node2 dient weiterhin verbleibenden Aggregaten, doch Aggregate, die bereits in Node3 verlagert wurden, und Node3 eigene Aggregate stoßen auf Client-Ausfälle, während Node3 gebootet wird.

Schritte

1. Knoten 3 aufbring.
2. Fahren Sie mit dem Controller-Upgrade fort.

Node2 stürzt während der zweiten Resource-Release-Phase ab

Wenn node2 während der Aggregatverschiebung abstürzt, wird node2 nicht übernommen.

Über diese Aufgabe

Node3 dient weiterhin den Aggregaten, die verschoben wurden, doch die Aggregate von node2 stoßen auf Client-Ausfälle.

Schritte

1. Knoten 2 aufbring.
2. Fahren Sie mit dem Controller-Upgrade fort.

Startet während der zweiten Verifikationsphase neu, erzeugt eine Panik oder schaltet die Stromversorgung aus

Node3 stürzt während der zweiten Verifikationsphase ab

Wenn während dieser Phase node3 abstürzt, wird die Übernahme nicht durchgeführt, da HA bereits deaktiviert ist.

Über diese Aufgabe

Es gibt einen Ausfall für nicht-Root-Aggregate, die bereits verschoben wurden, bis nach einem Neustart von Knoten3.

Schritte

1. Knoten 3 aufbring.

Ein Client-Ausfall tritt für alle Aggregate auf, während Node3 gestartet wird.

2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Node4 stürzt während der zweiten Verifikationsphase ab

Wenn node4 während dieser Phase abstürzt, wird die Übernahme nicht durchgeführt. Node3 stellt Daten aus den Aggregaten bereit.

Über diese Aufgabe

Es gibt einen Ausfall für nicht-Root-Aggregate, die bereits verschoben wurden, bis node4 neu startet.

Schritte

1. bringen sie node4 auf.
2. Fahren Sie mit dem Upgrade des Node-Paars fort.

Probleme, die in mehreren Phasen des Verfahrens auftreten können

Einige Probleme können in verschiedenen Phasen des Verfahrens auftreten.

Unerwartete Ausgabe des „Storage Failover show“-Befehls

Wenn während der Prozedur der Node, der alle Daten hostet, „Panic und“ oder versehentlich neu gebootet wird, wird möglicherweise die unerwartete Ausgabe für den angezeigt `storage failover show` Befehl vor und nach dem Neubooten, Panic oder aus- und Wiedereinschalten.

Über diese Aufgabe

Möglicherweise wird eine unerwartete Ausgabe von der angezeigt `storage failover show` Befehl in Phase 2, Stufe 3, Stufe 4 oder Stufe 5.

Das folgende Beispiel zeigt die erwartete Ausgabe von `storage failover show` Befehl, wenn auf dem Node, der alle Datenaggregate hostet, kein Neubooten oder „Panic“ erfolgt:

```
cluster::> storage failover show
```

Node	Partner	Takeover	
		Possible	State Description
node1	node2	false	Unknown
node2	node1	false	Node owns partner aggregates as part of the non-disruptive head upgrade procedure. Takeover is not possible: Storage failover is disabled.

Das folgende Beispiel zeigt die Ausgabe von `storage failover show` Befehl nach einem Neubooten oder Panic:

```
cluster::> storage failover show
```

Node	Partner	Takeover	
		Possible	State Description
node1	node2	-	Unknown
node2	node1	false	Waiting for node1, Partial giveback, Takeover is not possible: Storage failover is disabled

Obwohl die Ausgabe sagt, dass sich ein Node im teilweise Giveback befindet und der Storage-Failover

deaktiviert ist, können Sie diese Meldung ignorieren.

Schritte

Es ist keine Aktion erforderlich. Fahren Sie mit dem Upgrade des Node-Paars fort.

Fehler bei der LIF-Migration

Nach der Migration der LIFs sind diese nach der Migration in Phase 2, Phase 3 oder Phase 5 möglicherweise nicht online.

Schritte

1. Vergewissern Sie sich, dass die MTU-Port-Größe mit der Größe des Quell-Nodes identisch ist.

Wenn beispielsweise die MTU-Größe des Cluster-Ports am Quell-Node 9000 ist, sollte sie auf dem Ziel-Node 9000 sein.

2. Überprüfen Sie die physische Konnektivität des Netzkabels, wenn der physische Status des Ports lautet down.

Quellen

Wenn Sie die Verfahren in diesem Inhalt ausführen, müssen Sie möglicherweise Referenzinhalt konsultieren oder zu Referenzwebsites gehen.

Referenzinhalt

Die für dieses Upgrade spezifischen Inhalte sind in der folgenden Tabelle aufgeführt.

Inhalt	Beschreibung
"Administrationsübersicht mit der CLI"	Beschreibt das Verwalten von ONTAP Systemen, zeigt die Verwendung der CLI-Schnittstelle, den Zugriff auf das Cluster, das Managen von Nodes und vieles mehr.
"Entscheiden Sie, ob Sie System Manager oder die ONTAP CLI für das Cluster-Setup verwenden möchten"	Beschreibt die Einrichtung und Konfiguration von ONTAP.
"Festplatten- und Aggregatmanagement mit CLI"	Beschreibt das Verwalten von physischem ONTAP Storage mit der CLI. Hier erfahren Sie, wie Sie Aggregate erstellen, erweitern und managen, wie Sie mit Flash Pool Aggregaten arbeiten, Festplatten managen und RAID-Richtlinien managen.
"HA-Paar-Management"	Beschreibt die Installation und das Management von hochverfügbaren geclusterten Konfigurationen, einschließlich Storage Failover und Takeover/Giveback.
"Logisches Storage-Management mit der CLI"	Beschreibt, wie Sie Ihre logischen Storage-Ressourcen mithilfe von Volumes, FlexClone Volumes, Dateien und LUNs effizient managen FlexCache Volumes, Deduplizierung, Komprimierung, qtrees und Quotas.

Inhalt	Beschreibung
"MetroCluster Management und Disaster Recovery"	Beschreibt die Durchführung von MetroCluster-Switchover- und Switchback-Vorgängen sowohl bei geplanten Wartungsvorgängen als auch bei einem Notfall.
"MetroCluster Upgrade und Erweiterung"	Bietet Verfahren zum Upgrade von Controller- und Storage-Modellen in der MetroCluster Konfiguration, zum Wechsel von einer MetroCluster FC- zu einer MetroCluster IP-Konfiguration und zum erweitern der MetroCluster-Konfiguration durch Hinzufügen weiterer Nodes
"Netzwerkmanagement"	Beschreibt die Konfiguration und das Management von physischen und virtuellen Netzwerk-Ports (VLANs und Schnittstellengruppen), LIFs, Routing- und Host-Resolution-Services in Clustern; Optimierung des Netzwerk-Traffic durch Lastenausgleich; und Überwachung des Clusters mit SNMP
"ONTAP 9.0-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.0-Befehle.
"ONTAP 9.1-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.1-Befehle.
"ONTAP 9.2-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.2-Befehle.
"ONTAP 9.3-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.3-Befehle.
"ONTAP 9.4-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.4-Befehle.
"ONTAP 9.5-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.5-Befehle.
"ONTAP 9.6-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.6-Befehle.
"ONTAP 9.7-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.7-Befehle.
"ONTAP 9.8-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.8-Befehle.
"ONTAP 9.9.1-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.9.1-Befehle.
"ONTAP 9.10.1-Befehle: Manuelle Seitenreferenz"	Beschreibt die Syntax und die Verwendung der unterstützten ONTAP 9.10.1-Befehle.
"SAN-Management mit CLI"	In wird beschrieben, wie LUNs, Initiatorgruppen und Ziele mithilfe der iSCSI- und FC-Protokolle sowie Namespaces und Subsysteme mit dem NVMe/FC-Protokoll konfiguriert und gemanagt werden.
"Referenz zur SAN-Konfiguration"	Hier finden Sie Informationen zu FC- und iSCSI-Topologien sowie Kabelschemata.

Inhalt	Beschreibung
"Upgrade durch Verschieben von Volumes oder Storage"	Beschreibt das schnelle Upgrade von Controller Hardware in einem Cluster durch Verschieben von Storage oder Volumes. Beschreibt zudem, wie ein unterstütztes Modell in ein Festplatten-Shelf konvertiert wird.
"Upgrade von ONTAP"	Die Anleitungen zum Herunterladen und Aktualisieren von ONTAP.
"Verwenden Sie Befehle zum Austauschen von System-Controllern, um die mit ONTAP 9.15.1 und höher eingeführte Controller-Hardware zu aktualisieren"	Beschreibt die Verfahren für die Aggregatverschiebung, die für die in ONTAP 9.15.1 und höher eingeführten unterbrechungsfreien Upgrades von Controllern mit Befehlen „System Controller Replace“ erforderlich sind.
"Aktualisieren Sie Controller-Modelle im selben Chassis mit Befehlen „System-Controller ersetzen“"	Beschreibt die Verfahren zur Aggregatverschiebung, die für ein unterbrechungsfreies Upgrade eines Systems erforderlich sind, wobei das alte System-Chassis und die alten Festplatten erhalten bleiben.
"Verwenden Sie „System Controller Replace“-Befehle, um das Upgrade der Controller Hardware mit ONTAP 9.8 oder höher durchzuführen"	Beschreibt die Verfahren für Aggregatverschiebung, die nötig sind, um Controller, die ONTAP 9.8 ausführen, durch den „System-Controller-Austausch“-Befehl unterbrechungsfrei zu aktualisieren.
"Nutzen Sie die Aggregatverschiebung, um manuell ein Upgrade der Controller-Hardware mit ONTAP 9.8 oder höher durchzuführen"	Beschreibt das Verfahren für die Aggregatverschiebung, die erforderlich sind, um manuelle, unterbrechungsfreie Controller-Upgrades mit ONTAP 9.8 oder höher durchzuführen.
"Verwenden Sie „System Controller Replace“-Befehle, um Controller Hardware mit ONTAP 9.5 auf ONTAP 9.7 zu aktualisieren"	Beschreibt die Verfahren für Aggregatverschiebung, die nötig sind, um ein unterbrechungsfreies Upgrade der Controller, die ONTAP 9.5 auf ONTAP 9.7 mithilfe von Befehlen zum Austausch des System-Controllers durchführen, durchzuführen.
"Nutzen Sie die Aggregatverschiebung, um manuell ein Upgrade der Controller-Hardware mit ONTAP 9.7 oder einer älteren Version durchzuführen"	Beschreibt die Verfahren für die Aggregatverschiebung, die erforderlich sind, um manuelle, unterbrechungsfreie Controller-Upgrades mit ONTAP 9.7 oder früher durchzuführen.

Referenzstandorte

Der ["NetApp Support Website"](#) Enthält auch Dokumentation zu Netzwerkschnittstellenkarten (NICs) und anderer Hardware, die Sie mit Ihrem System verwenden könnten. Es enthält auch die ["Hardware Universe"](#), Die Informationen über die Hardware liefert, die das neue System unterstützt.

Datenzugriff ["ONTAP 9-Dokumentation"](#).

Auf das zugreifen ["Active IQ Config Advisor"](#) Werkzeug.

Copyright-Informationen

Copyright © 2026 NetApp. Alle Rechte vorbehalten. Gedruckt in den USA. Dieses urheberrechtlich geschützte Dokument darf ohne die vorherige schriftliche Genehmigung des Urheberrechtsinhabers in keiner Form und durch keine Mittel – weder grafische noch elektronische oder mechanische, einschließlich Fotokopieren, Aufnehmen oder Speichern in einem elektronischen Abrufsystem – auch nicht in Teilen, vervielfältigt werden.

Software, die von urheberrechtlich geschütztem NetApp Material abgeleitet wird, unterliegt der folgenden Lizenz und dem folgenden Haftungsausschluss:

DIE VORLIEGENDE SOFTWARE WIRD IN DER VORLIEGENDEN FORM VON NETAPP ZUR VERFÜGUNG GESTELLT, D. H. OHNE JEGLICHE EXPLIZITE ODER IMPLIZITE GEWÄHRLEISTUNG, EINSCHLIESSLICH, JEDOCH NICHT BESCHRÄNKT AUF DIE STILLSCHWEIGENDE GEWÄHRLEISTUNG DER MARKTGÄNGIGKEIT UND EIGNUNG FÜR EINEN BESTIMMTEN ZWECK, DIE HIERMIT AUSGESCHLOSSEN WERDEN. NETAPP ÜBERNIMMT KEINERLEI HAFTUNG FÜR DIREKTE, INDIREKTE, ZUFÄLLIGE, BESONDERE, BEISPIELHAFT SCHÄDEN ODER FOLGESCHÄDEN (EINSCHLIESSLICH, JEDOCH NICHT BESCHRÄNKT AUF DIE BESCHAFFUNG VON ERSATZWAREN ODER -DIENSTLEISTUNGEN, NUTZUNGS-, DATEN- ODER GEWINNVERLUSTE ODER UNTERBRECHUNG DES GESCHÄFTSBETRIEBS), UNABHÄNGIG DAVON, WIE SIE VERURSACHT WURDEN UND AUF WELCHER HAFTUNGSTHEORIE SIE BERUHEN, OB AUS VERTRAGLICH FESTGELEGTER HAFTUNG, VERSCHULDENSUNABHÄNGIGER HAFTUNG ODER DELIKTSHAFTUNG (EINSCHLIESSLICH FAHRLÄSSIGKEIT ODER AUF ANDEREM WEGE), DIE IN IRGEND EINER WEISE AUS DER NUTZUNG DIESER SOFTWARE RESULTIEREN, SELBST WENN AUF DIE MÖGLICHKEIT DERARTIGER SCHÄDEN HINGEWIESEN WURDE.

NetApp behält sich das Recht vor, die hierin beschriebenen Produkte jederzeit und ohne Vorankündigung zu ändern. NetApp übernimmt keine Verantwortung oder Haftung, die sich aus der Verwendung der hier beschriebenen Produkte ergibt, es sei denn, NetApp hat dem ausdrücklich in schriftlicher Form zugestimmt. Die Verwendung oder der Erwerb dieses Produkts stellt keine Lizenzierung im Rahmen eines Patentrechts, Markenrechts oder eines anderen Rechts an geistigem Eigentum von NetApp dar.

Das in diesem Dokument beschriebene Produkt kann durch ein oder mehrere US-amerikanische Patente, ausländische Patente oder anhängige Patentanmeldungen geschützt sein.

ERLÄUTERUNG ZU „RESTRICTED RIGHTS“: Nutzung, Vervielfältigung oder Offenlegung durch die US-Regierung unterliegt den Einschränkungen gemäß Unterabschnitt (b)(3) der Klausel „Rights in Technical Data – Noncommercial Items“ in DFARS 252.227-7013 (Februar 2014) und FAR 52.227-19 (Dezember 2007).

Die hierin enthaltenen Daten beziehen sich auf ein kommerzielles Produkt und/oder einen kommerziellen Service (wie in FAR 2.101 definiert) und sind Eigentum von NetApp, Inc. Alle technischen Daten und die Computersoftware von NetApp, die unter diesem Vertrag bereitgestellt werden, sind gewerblicher Natur und wurden ausschließlich unter Verwendung privater Mittel entwickelt. Die US-Regierung besitzt eine nicht ausschließliche, nicht übertragbare, nicht unterlizenzierbare, weltweite, limitierte unwiderrufliche Lizenz zur Nutzung der Daten nur in Verbindung mit und zur Unterstützung des Vertrags der US-Regierung, unter dem die Daten bereitgestellt wurden. Sofern in den vorliegenden Bedingungen nicht anders angegeben, dürfen die Daten ohne vorherige schriftliche Genehmigung von NetApp, Inc. nicht verwendet, offengelegt, vervielfältigt, geändert, aufgeführt oder angezeigt werden. Die Lizenzrechte der US-Regierung für das US-Verteidigungsministerium sind auf die in DFARS-Klausel 252.227-7015(b) (Februar 2014) genannten Rechte beschränkt.

Markeninformationen

NETAPP, das NETAPP Logo und die unter <http://www.netapp.com/TM> aufgeführten Marken sind Marken von NetApp, Inc. Andere Firmen und Produktnamen können Marken der jeweiligen Eigentümer sein.