



ASA Cluster for the Firepower 4100/9300

Clustering lets you group multiple Firepower 4100/9300 chassis ASAs together as a single logical device. The Firepower 4100/9300 chassis series includes the Firepower 9300 and Firepower 4100 series. A cluster provides all the convenience of a single device (management, integration into a network) while achieving the increased throughput and redundancy of multiple devices.



Note Some features are not supported when using clustering. See [Unsupported Features with Clustering](#), on page 74.

- [About Clustering on the Firepower 4100/9300 Chassis](#), on page 1
- [Requirements and Prerequisites for Clustering on the Firepower 4100/9300 Chassis](#), on page 8
- [Licenses for Clustering on the Firepower 4100/9300 Chassis](#), on page 9
- [Clustering Guidelines and Limitations](#), on page 11
- [Configure Clustering on the Firepower 4100/9300 Chassis](#), on page 15
- [FXOS: Remove a Cluster Node](#), on page 49
- [ASA: Manage Cluster Members](#), on page 50
- [ASA: Monitoring the ASA Cluster on the Firepower 4100/9300 chassis](#), on page 55
- [Troubleshooting Distributed S2S VPN](#), on page 65
- [Examples for ASA Clustering](#), on page 67
- [Reference for Clustering](#), on page 74
- [History for ASA Clustering on the Firepower 4100/9300](#), on page 89

About Clustering on the Firepower 4100/9300 Chassis

When you deploy a cluster on the Firepower 4100/9300 chassis, it does the following:

- Creates a *cluster-control link* (by default, port-channel 48) for node-to-node communication.

For a cluster isolated to security modules within one Firepower 9300 chassis, this link utilizes the Firepower 9300 backplane for cluster communications.

For clustering with multiple chassis, you need to manually assign physical interface(s) to this EtherChannel for communications between chassis.

- Creates the cluster bootstrap configuration within the application.

When you deploy the cluster, the chassis supervisor pushes a minimal bootstrap configuration to each unit that includes the cluster name, cluster control link interface, and other cluster settings. Some parts of the bootstrap configuration may be user-configurable within the application if you want to customize your clustering environment.

- Assigns data interfaces to the cluster as *Spanned* interfaces.

For a cluster isolated to security modules within one Firepower 9300 chassis, spanned interfaces are not limited to EtherChannels, like it is for clustering with multiple chassis. The Firepower 9300 supervisor uses EtherChannel technology internally to load-balance traffic to multiple modules on a shared interface, so any data interface type works for Spanned mode. For clustering with multiple chassis, you must use Spanned EtherChannels for all data interfaces.



Note Individual interfaces are not supported, with the exception of a management interface.

- Assigns a management interface to all units in the cluster.

See the following sections for more information about clustering.

Bootstrap Configuration

When you deploy the cluster, the Firepower 4100/9300 chassis supervisor pushes a minimal bootstrap configuration to each unit that includes the cluster name, cluster control link interface, and other cluster settings. Some parts of the bootstrap configuration are user-configurable if you want to customize your clustering environment.

Cluster Members

Cluster members work together to accomplish the sharing of the security policy and traffic flows.

One member of the cluster is the **control** unit. The control unit is determined automatically. All other members are **data** units.

You must perform all configuration on the control unit only; the configuration is then replicated to the data units.

Some features do not scale in a cluster, and the control unit handles all traffic for those features. See [Centralized Features for Clustering, on page 75](#).

Cluster Control Link

The cluster-control link is an EtherChannel (port-channel 48) for unit-to-unit communication. For intra-chassis clustering, this link utilizes the Firepower 9300 backplane for cluster communications. For inter-chassis clustering, you need to manually assign physical interface(s) to this EtherChannel on the Firepower 4100/9300 chassis for communications between chassis.

For a 2-chassis inter-chassis cluster, do not directly-connect the cluster control link from one chassis to the other chassis. If you directly connect the interfaces, then when one unit fails, the cluster control link fails, and

thus the remaining healthy unit fails. If you connect the cluster control link through a switch, then the cluster control link remains up for the healthy unit.

Cluster control link traffic includes both control and data traffic.

Control traffic includes:

- Control node election.
- Configuration replication.
- Health monitoring.

Data traffic includes:

- State replication.
- Connection ownership queries and data packet forwarding.

See the following sections for more information about the cluster control link.

Size the Cluster Control Link

If possible, you should size the cluster control link to match the expected throughput of each chassis so the cluster control link can handle the worst-case scenarios.

Cluster control link traffic is comprised mainly of state update and forwarded packets. The amount of traffic at any given time on the cluster control link varies. The amount of forwarded traffic depends on the load-balancing efficacy or whether there is a lot of traffic for centralized features. For example:

- NAT results in poor load balancing of connections, and the need to rebalance all returning traffic to the correct units.
- AAA for network access is a centralized feature, so all traffic is forwarded to the control unit.
- When membership changes, the cluster needs to rebalance a large number of connections, thus temporarily using a large amount of cluster control link bandwidth.

A higher-bandwidth cluster control link helps the cluster to converge faster when there are membership changes and prevents throughput bottlenecks.

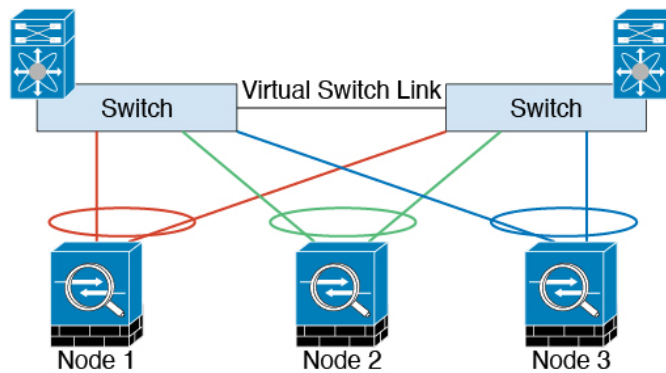


Note If your cluster has large amounts of asymmetric (rebalanced) traffic, then you should increase the cluster control link size.

Cluster Control Link Redundancy

We recommend using an EtherChannel for the cluster control link, so that you can pass traffic on multiple links in the EtherChannel while still achieving redundancy.

The following diagram shows how to use an EtherChannel as a cluster control link in a Virtual Switching System (VSS), Virtual Port Channel (vPC), StackWise, or StackWise Virtual environment. All links in the EtherChannel are active. When the switch is part of a redundant system, then you can connect firewall interfaces within the same EtherChannel to separate switches in the redundant system. The switch interfaces are members of the same EtherChannel port-channel interface, because the separate switches act like a single switch. Note that this EtherChannel is device-local, not a Spanned EtherChannel.



Cluster Control Link Reliability

To ensure cluster control link functionality, be sure the round-trip time (RTT) between units is less than 20 ms. This maximum latency enhances compatibility with cluster members installed at different geographical sites. To check your latency, perform a ping on the cluster control link between units.

The cluster control link must be reliable, with no out-of-order or dropped packets; for example, for inter-site deployment, you should use a dedicated link.

Cluster Control Link Network

The Firepower 4100/9300 chassis auto-generates the cluster control link interface IP address for each unit based on the chassis ID and slot ID: `127.2.chassis_id.slot_id`. You can customize this IP address when you deploy the cluster. The cluster control link network cannot include any routers between units; only Layer 2 switching is allowed. For inter-site traffic, Cisco recommends using Overlay Transport Virtualization (OTV).

Cluster Interfaces

For a cluster isolated to security modules within one Firepower 9300 chassis, you can assign both physical interfaces or EtherChannels (also known as port channels) to the cluster. Interfaces assigned to the cluster are Spanned interfaces that load-balance traffic across all members of the cluster.

For clustering with multiple chassis, you can only assign data EtherChannels to the cluster. These Spanned EtherChannels include the same member interfaces on each chassis; on the upstream switch, all of these interfaces are included in a single EtherChannel, so the switch does not know that it is connected to multiple devices.

Individual interfaces are not supported, with the exception of a management interface.

Connecting to a Redundant Switch System

We recommend connecting EtherChannels to a redundant switch system such as a VSS, vPC, StackWise, or StackWise Virtual system to provide redundancy for your interfaces.

Configuration Replication

All nodes in the cluster share a single configuration. You can only make configuration changes on the control node (with the exception of the bootstrap configuration), and changes are automatically synced to all other nodes in the cluster.

Secure Firewall ASA Cluster Management

One of the benefits of using ASA clustering is the ease of management. This section describes how to manage the cluster.

Management Network

We recommend connecting all units to a single management network. This network is separate from the cluster control link.

Management Interface

You must assign a Management type interface to the cluster. This interface is a special individual interface as opposed to a Spanned interface. The management interface lets you connect directly to each unit.

The Main cluster IP address is a fixed address for the cluster that always belongs to the current control unit. You also configure a range of addresses so that each unit, including the current control unit, can use a Local address from the range. The Main cluster IP address provides consistent management access to an address; when a control unit changes, the Main cluster IP address moves to the new control unit, so management of the cluster continues seamlessly.

For example, you can manage the cluster by connecting to the Main cluster IP address, which is always attached to the current control unit. To manage an individual member, you can connect to the Local IP address.



Note To-the-box traffic needs to be directed to the node's management IP address; to-the-box traffic is not forwarded over the cluster control link to any other node.

For outbound management traffic such as TFTP or syslog, each unit, including the control unit, uses the Local IP address to connect to the server.

Control Unit Management Vs. Data Unit Management

All management and monitoring can take place on the control node. From the control node, you can check runtime statistics, resource usage, or other monitoring information of all nodes. You can also issue a command to all nodes in the cluster, and replicate the console messages from data nodes to the control node.

You can monitor data nodes directly if desired. Although also available from the control node, you can perform file management on data nodes (including backing up the configuration and updating images). The following functions are not available from the control node:

- Monitoring per-node cluster-specific statistics.
- Syslog monitoring per node (except for syslogs sent to the console when console replication is enabled).
- SNMP

- NetFlow

Crypto Key Replication

When you create a crypto key on the control node, the key is replicated to all data nodes. If you have an SSH session to the Main cluster IP address, you will be disconnected if the control node fails. The new control node uses the same key for SSH connections, so that you do not need to update the cached SSH host key when you reconnect to the new control node.

ASDM Connection Certificate IP Address Mismatch

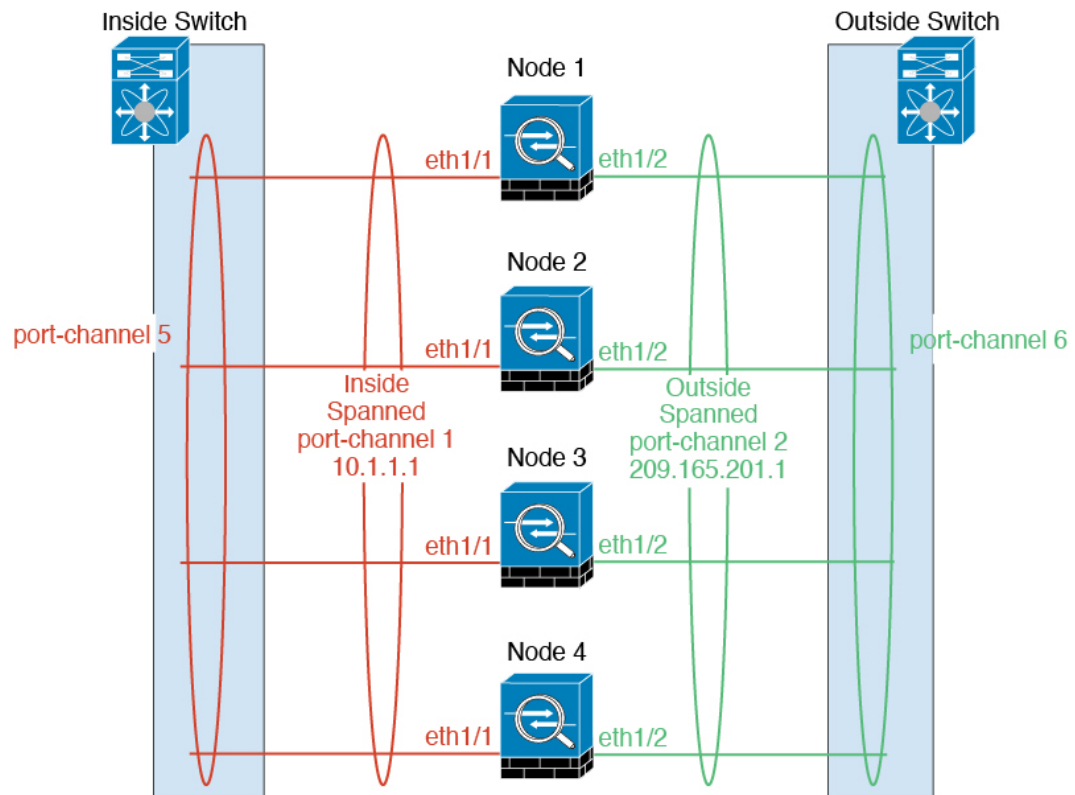
By default, a self-signed certificate is used for the ASDM connection based on the Local IP address. If you connect to the Main cluster IP address using ASDM, then a warning message about a mismatched IP address might appear because the certificate uses the Local IP address, and not the Main cluster IP address. You can ignore the message and establish the ASDM connection. However, to avoid this type of warning, you can enroll a certificate that contains the Main cluster IP address and all the Local IP addresses from the IP address pool. You can then use this certificate for each cluster member. See <https://www.cisco.com/c/en/us/td/docs/security/asdm/identity-cert/cert-install.html> for more information.

Spanned EtherChannels (Recommended)

You can group one or more interfaces per chassis into an EtherChannel that spans all chassis in the cluster. The EtherChannel aggregates the traffic across all the available active interfaces in the channel.

A Spanned EtherChannel can be configured in both routed and transparent firewall modes. In routed mode, the EtherChannel is configured as a routed interface with a single IP address. In transparent mode, the IP address is assigned to the BVI, not to the bridge group member interface.

The EtherChannel inherently provides load balancing as part of basic operation.



Inter-Site Clustering

For inter-site installations, you can take advantage of ASA clustering as long as you follow the recommended guidelines.

You can configure each cluster chassis to belong to a separate site ID.

Site IDs work with site-specific MAC addresses and IP addresses. Packets egressing the cluster use a site-specific MAC address and IP address, while packets received by the cluster use a global MAC address and IP address. This feature prevents the switches from learning the same global MAC address from both sites on two different ports, which causes MAC flapping; instead, they only learn the site MAC address. Site-specific MAC addresses and IP address are supported for routed mode using Spanned EtherChannels only.

Site IDs are also used to enable flow mobility using LISP inspection, director localization to improve performance and reduce round-trip time latency for inter-site clustering for data centers, and site redundancy for connections where a backup owner of a traffic flow is always at a different site from the owner.

See the following sections for more information about inter-site clustering:

- Sizing the Data Center Interconnect—[Requirements and Prerequisites for Clustering on the Firepower 4100/9300 Chassis, on page 8](#)
- Inter-Site Guidelines—[Clustering Guidelines and Limitations, on page 11](#)
- Configure Cluster Flow Mobility—[Configure Cluster Flow Mobility, on page 38](#)
- Enable Director Localization—[Enable Director Localization, on page 36](#)

- Enable Site Redundancy—[Enable Director Localization](#), on page 36

Requirements and Prerequisites for Clustering on the Firepower 4100/9300 Chassis

Maximum Clustering Units Per Model

- Firepower 4100—16 chassis
- Firepower 9300—16 modules. For example, you can use 1 module in 16 chassis, or 2 modules in 8 chassis, or any combination that provides a maximum of 16 modules.

Hardware and Software Requirements for Inter-Chassis Clustering

All chassis in a cluster:

- For the Firepower 4100: All chassis must be the same model. For the Firepower 9300: All security modules must be the same type. For example, if you use clustering, all modules in the Firepower 9300 must be SM-40s. You can have different quantities of installed security modules in each chassis, although all modules present in the chassis must belong to the cluster including any empty slots.
- Must run the identical FXOS and application software except at the time of an image upgrade. Mismatched software versions can lead to poor performance, so be sure to upgrade all nodes in the same maintenance window.
- Must include the same interface configuration for interfaces you assign to the cluster, such as the same Management interface, EtherChannels, active interfaces, speed and duplex, and so on. You can use different network module types on the chassis as long as the capacity matches for the same interface IDs and interfaces can successfully bundle in the same spanned EtherChannel. Note that all data interfaces must be EtherChannels in clusters with multiple chassis. If you change the interfaces in FXOS after you enable clustering (by adding or removing interface modules, or configuring EtherChannels, for example), then perform the same changes on each chassis, starting with the data nodes, and ending with the control node. Note that if you remove an interface in FXOS, the ASA configuration retains the related commands so that you can make any necessary adjustments; removing an interface from the configuration can have wide effects. You can manually remove the old interface configuration.
- Must use the same NTP server. Do not set the time manually.
- ASA: Each FXOS chassis must be registered with the License Authority or satellite server. There is no extra cost for data nodes. For permanent license reservation, you must purchase separate licenses for each chassis. For Firewall Threat Defense, all licensing is handled by the Firewall Management Center.

Switch Requirements

- Be sure to complete the switch configuration and successfully connect all the EtherChannels from the chassis to the switch(es) before you configure clustering on the Firepower 4100/9300 chassis.
- For supported switch characteristics, see [Cisco FXOS Compatibility](#).

Sizing the Data Center Interconnect for Inter-Site Clustering

You should reserve bandwidth on the data center interconnect (DCI) for cluster control link traffic equivalent to the following calculation:

$$\frac{\text{\# of cluster members per site}}{2} \times \text{cluster control link size per member}$$

If the number of members differs at each site, use the larger number for your calculation. The minimum bandwidth for the DCI should not be less than the size of the cluster control link for one member.

For example:

- For 4 members at 2 sites:
 - 4 cluster members total
 - 2 members at each site
 - 5 Gbps cluster control link per member

Reserved DCI bandwidth = 5 Gbps (2/2 x 5 Gbps).

- For 6 members at 3 sites, the size increases:
 - 6 cluster members total
 - 3 members at site 1, 2 members at site 2, and 1 member at site 3
 - 10 Gbps cluster control link per member

Reserved DCI bandwidth = 15 Gbps (3/2 x 10 Gbps).

- For 2 members at 2 sites:
 - 2 cluster members total
 - 1 member at each site
 - 10 Gbps cluster control link per member

Reserved DCI bandwidth = 10 Gbps (1/2 x 10 Gbps = 5 Gbps; but the minimum bandwidth should not be less than the size of the cluster control link (10 Gbps)).

Licenses for Clustering on the Firepower 4100/9300 Chassis

Smart Software Manager Regular and On-Prem

The clustering feature itself does not require any licenses. To use Strong Encryption and other optional licenses, each Firepower 4100/9300 chassis must be registered with the License Authority or Smart Software Manager Regular and On-Prem server. There is no extra cost for data units.

The Strong Encryption license is automatically enabled for qualified customers when you apply the registration token. When using the token, each chassis must have the same encryption license. For the optional Strong Encryption (3DES/AES) feature license enabled in the ASA configuration, see below.

In the ASA license configuration, you can only configure smart licensing on the control unit. The configuration is replicated to the data units, but for some licenses, they do not use the configuration; it remains in a cached state, and only the control unit requests the license. The licenses are aggregated into a single cluster license that is shared by the cluster units, and this aggregated license is also cached on the data units to be used if one of them becomes the control unit in the future. Each license type is managed as follows:

- **Essentials**—Only the control unit requests the Essentials license from the server, and both units can use it due to license aggregation.
- **Context**—Only the control unit requests the Context license from the server. The Essentials license includes 10 contexts by default and is present on all cluster members. The value from each unit's Essentials license plus the value of the Context license on the control unit are combined up to the platform limit in an aggregated cluster license. For example:
 - You have 6 Firepower 9300 modules in the cluster. The Essentials license includes 10 contexts; for 6 units, these licenses add up to 60 contexts. You configure an additional 20-Context license on the control unit. Therefore, the aggregated cluster license includes 80 contexts. Because the platform limit for one module is 250, the combined license allows a maximum of 250 contexts; the 80 contexts are within the limit. Therefore, you can configure up to 80 contexts on the control unit; each data unit will also have 80 contexts through configuration replication.
 - You have 3 Firepower 4112 units in the cluster. The Essentials license includes 10 contexts; for 3 units, these licenses add up to 30 contexts. You configure an additional 250-Context license on the control unit. Therefore, the aggregated cluster license includes 280 contexts. Because the platform limit for one unit is 250, the combined license allows a maximum of 250 contexts; the 280 contexts are over the limit. Therefore, you can only configure up to 250 contexts on the control unit; each data unit will also have 250 contexts through configuration replication. In this case, you should only configure the control unit Context license to be 220 contexts.
- **Carrier**—Required for Distributed S2S VPN. This license is a per-unit entitlement, and each unit requests its own license from the server.
- **Strong Encryption (3DES)**—For pre-2.3.0 Cisco Smart Software Manager On-Prem deployment; or if your Smart Account is not authorized for strong encryption, but Cisco has determined that you are allowed to use strong encryption, you can manually add a strong encryption license to your account. This license is a per-unit entitlement, and each unit requests its own license from the server.

If a new control unit is elected, the new control unit continues to use the aggregated license. It also uses the cached license configuration to re-request the control unit license. When the old control unit rejoins the cluster as a data unit, it releases the control unit license entitlement. Before the data unit releases the license, the control unit's license might be in a non-compliant state if there are no available licenses in the account. The retained license is valid for 30 days, but if it is still non-compliant after the grace period, you will not be able to make configuration changes to features requiring special licenses; operation is otherwise unaffected. The new active unit sends an entitlement authorization renewal request every 12 hours until the license is compliant. You should refrain from making configuration changes until the license requests are completely processed. If a unit leaves the cluster, the cached control configuration is removed, while the per-unit entitlements are retained. In particular, you would need to re-request the Context license on non-cluster units.

Permanent License Reservation

For permanent license reservation, you must purchase separate licenses for each chassis and enable the licenses *before* you configure clustering.

Clustering Guidelines and Limitations

Switches for Clustering

- Make sure connected switches match the MTU for both cluster data interfaces and the cluster control link interface. You should configure the cluster control link interface MTU to be at least 100 bytes higher than the data interface MTU, so make sure to configure the cluster control link connecting switch appropriately. Because the cluster control link traffic includes data packet forwarding, the cluster control link needs to accommodate the entire size of a data packet plus cluster traffic overhead. In addition, we do not recommend setting the cluster control link MTU between 2561 and 8362; due to block pool handling, this MTU size is not optimal for system operation. When a node joins the cluster, it checks MTU compatibility by sending a ping to the control node with a packet size matching the cluster control link MTU. If the ping fails, a notification is generated so you can fix the MTU mismatch on connecting switches and try again.
- For Cisco IOS XR systems, if you want to set a non-default MTU, set the IOS XR interface MTU to be 14 bytes higher than the cluster device MTU. Otherwise, OSPF adjacency peering attempts may fail unless the **mtu-ignore** option is used. Note that the cluster device MTU should match the IOS XR IPv4 MTU. This adjustment is not required for Cisco Catalyst and Cisco Nexus switches.
- On the switch(es) for the cluster control link interfaces, you can optionally enable Spanning Tree PortFast on the switch ports connected to the cluster unit to speed up the join process for new units.
- On the switch, we recommend that you use one of the following EtherChannel load-balancing algorithms: **source-dest-ip** or **src-dst-mixed-ip-port** (see the Cisco Nexus OS and Cisco IOS-XE **port-channel load-balance** command). Do not use a **vlan** keyword in the load-balance algorithm because it can cause unevenly distributed traffic to the devices in a cluster. *Do not* change the load-balancing algorithm from the default on the cluster device.
- If you change the load-balancing algorithm of the EtherChannel on the switch, the EtherChannel interface on the switch temporarily stops forwarding traffic, and the Spanning Tree Protocol restarts. There will be a delay before traffic starts flowing again.
- Switches on the cluster control link path should not verify the L4 checksum. Redirected traffic over the cluster control link does not have a correct L4 checksum. Switches that verify the L4 checksum could cause traffic to be dropped.
- Port-channel bundling downtime should not exceed the configured keepalive interval.
- On Supervisor 2T EtherChannels, the default hash distribution algorithm is adaptive. To avoid asymmetric traffic in a VSS design, change the hash algorithm on the port-channel connected to the cluster device to fixed:

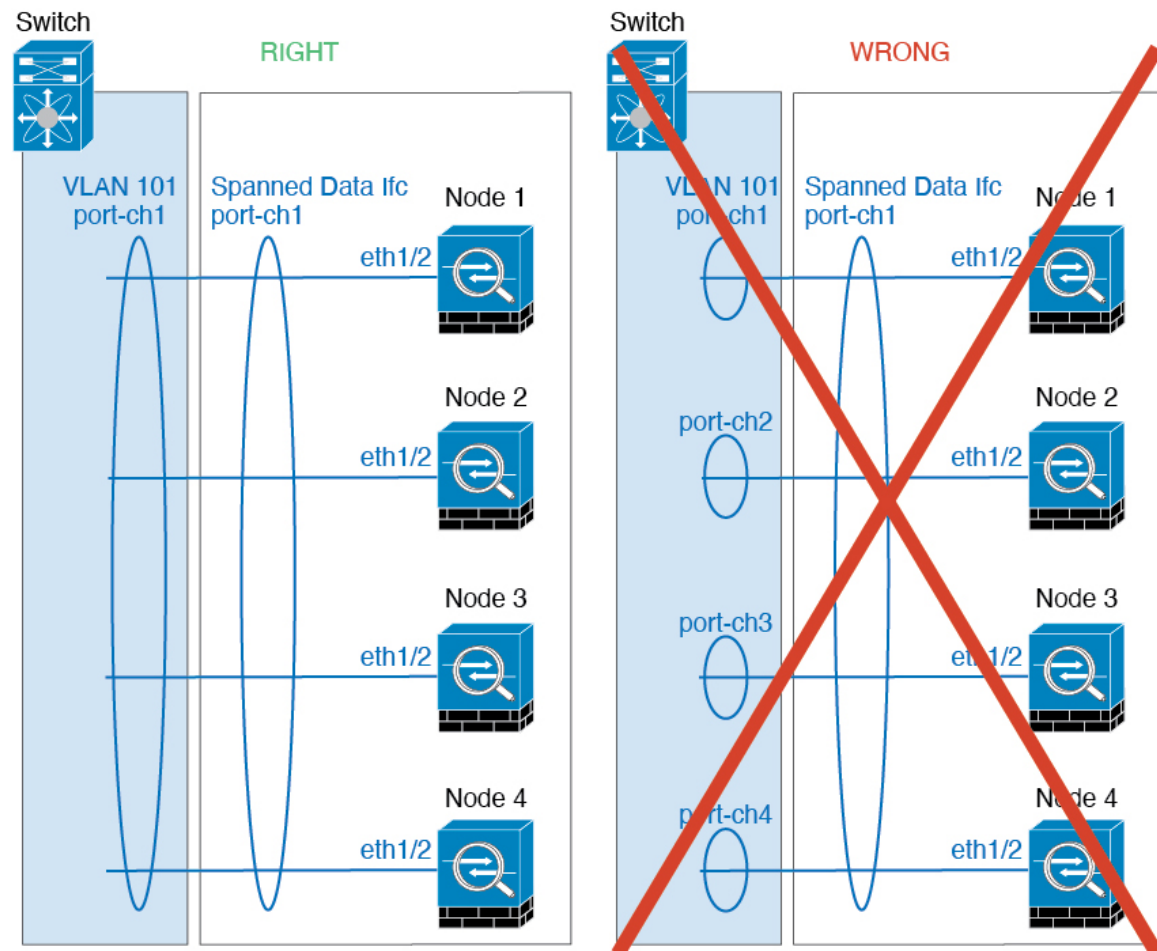
```
router(config)# port-channel id hash-distribution fixed
```

Do not change the algorithm globally; you may want to take advantage of the adaptive algorithm for the VSS peer link.
- Unlike ASA hardware clusters, Firepower 4100/9300 clusters support LACP graceful convergence. So for the platform, you can leave LACP graceful convergence enabled on connected Cisco Nexus switches.
- When you see slow bundling of a Spanned EtherChannel on the switch, you can enable LACP rate fast for an individual interface on the switch. FXOS EtherChannels have the LACP rate set to fast by default.

Note that some switches, such as the Nexus series, do not support LACP rate fast when performing in-service software upgrades (ISSUs), so we do not recommend using ISSUs with clustering.

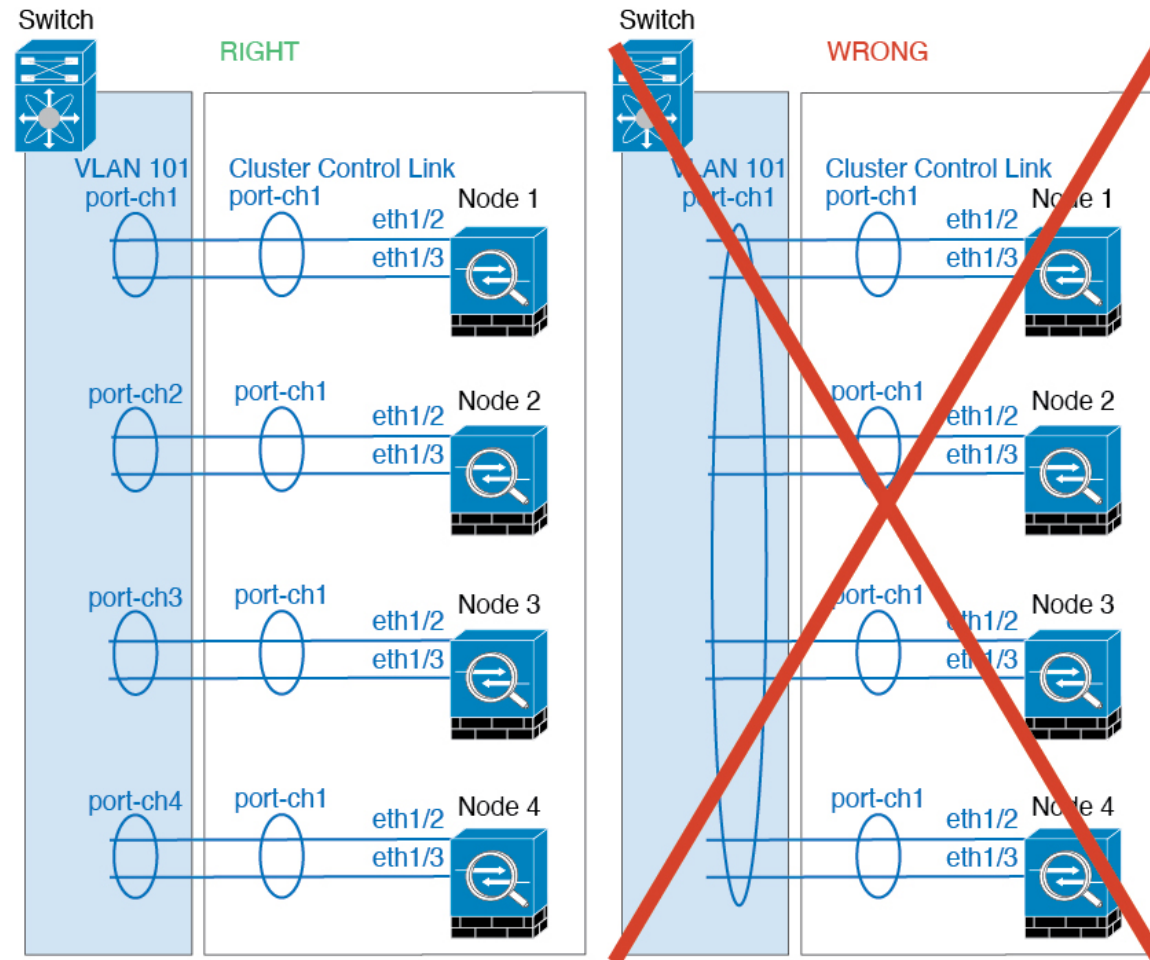
EtherChannels for Clustering

- In Catalyst 3750-X Cisco IOS software versions earlier than 15.1(1)S2, the cluster unit did not support connecting an EtherChannel to a switch stack. With default switch settings, if the cluster unit EtherChannel is connected cross stack, and if the control unit switch is powered down, then the EtherChannel connected to the remaining switch will not come up. To improve compatibility, set the **stack-mac persistent timer** command to a large enough value to account for reload time; for example, 8 minutes or 0 for indefinite. Or, you can upgrade to more a more stable switch software version, such as 15.1(1)S2.
- Spanned vs. Device-Local EtherChannel Configuration—Be sure to configure the switch appropriately for Spanned EtherChannels vs. Device-local EtherChannels.
 - Spanned EtherChannels—For cluster unit *Spanned* EtherChannels, which span across all members of the cluster, the interfaces are combined into a single EtherChannel on the switch. Make sure each interface is in the same channel group on the switch.



- Device-local EtherChannels—For cluster unit *Device-local* EtherChannels including any EtherChannels configured for the cluster control link, be sure to configure discrete EtherChannels

on the switch; do not combine multiple cluster unit EtherChannels into one EtherChannel on the switch.



Inter-Site Clustering

See the following guidelines for inter-site clustering:

- The cluster control link latency must be less than 20 ms round-trip time (RTT).
- The cluster control link must be reliable, with no out-of-order or dropped packets; for example, you should use a dedicated link.
- Do not configure connection rebalancing; you do not want connections rebalanced to cluster members at a different site.
- The ASA does not encrypt forwarded data traffic on the cluster control link because it is a dedicated link, even when used on a Data Center Interconnect (DCI). If you use Overlay Transport Virtualization (OTV), or are otherwise extending the cluster control link outside of the local administrative domain, you can configure encryption on your border routers such as 802.1AE MacSec over OTV.
- The cluster implementation does not differentiate between members at multiple sites for incoming connections; therefore, connection roles for a given connection may span across sites. This is expected.

behavior. However, if you enable director localization, the local director role is always chosen from the same site as the connection owner (according to site ID). Also, the local director chooses a new owner at the same site if the original owner fails (Note: if the traffic is asymmetric across sites, and there is continuous traffic from the remote site after the original owner fails, then a node from the remote site might become the new owner if it receives a data packet within the re-hosting window.).

- For director localization, the following traffic types do not support localization: NAT or PAT traffic; SCTP-inspected traffic; Fragmentation owner query.
- For UDP long-lived flows in a North-South deployment, routing loops can occur if nodes at the original flow owner site fail and then come back up, after which the flow is directed back to the original site. If the new owner at the other site doesn't have a route to the destination, it will route the flow back to the internet, causing a loop. In this case, use the **clear conn** command on the new owner to force the flow to be reestablished.
- For transparent mode, if the cluster is placed between a pair of inside and outside routers (AKA North-South insertion), you must ensure that both inside routers share a MAC address, and also that both outside routers share a MAC address. When a cluster member at site 1 forwards a connection to a member at site 2, the destination MAC address is preserved. The packet will only reach the router at site 2 if the MAC address is the same as the router at site 1.
- For transparent mode, if the cluster is placed between data networks and the gateway router at each site for firewalling between internal networks (AKA East-West insertion), then each gateway router should use a First Hop Redundancy Protocol (FHRP) such as HSRP to provide identical virtual IP and MAC address destinations at each site. The data VLANs are extended across the sites using Overlay Transport Virtualization (OTV), or something similar. You need to create filters to prevent traffic that is destined to the local gateway router from being sent over the DCI to the other site. If the gateway router becomes unreachable at one site, you need to remove any filters so traffic can successfully reach the other site's gateway.
- For transparent mode, if the cluster is connected to an HSRP router, you must add the router HSRP MAC address as a static MAC address table entry on the ASA (see [Add a Static MAC Address for Bridge Groups](#)). When adjacent routers use HSRP, traffic destined to the HSRP IP address will be sent to the HSRP MAC Address, but return traffic will be sourced from the MAC address of a particular router's interface in the HSRP pair. Therefore, the ASA MAC address table is typically only updated when the ASA ARP table entry for the HSRP IP address expires, and the ASA sends an ARP request and receives a reply. Because the ASA's ARP table entries expire after 14400 seconds by default, but the MAC address table entry expires after 300 seconds by default, a static MAC address entry is required to avoid MAC address table expiration traffic drops.
- For routed mode using Spanned EtherChannel, configure site-specific MAC addresses. Extend the data VLANs across the sites using OTV, or something similar. You need to create filters to prevent traffic that is destined to the global MAC address from being sent over the DCI to the other site. If the cluster becomes unreachable at one site, you need to remove any filters so traffic can successfully reach the other site's cluster nodes. Dynamic routing is not supported when an inter-site cluster acts as the first hop router for an extended segment.

Additional Guidelines

- When significant topology changes occur (such as adding or removing an EtherChannel interface, enabling or disabling an interface on the Firepower 4100/9300 chassis or the switch, adding an additional switch to form a VSS, vPC, StackWise, or StackWise Virtual) you should disable the health check feature, and

also disable interface monitoring for the disabled interfaces . When the topology change is complete, and the configuration change is synced to all units, you can re-enable the health check feature.

- When adding a unit to an existing cluster, or when reloading a unit, there will be a temporary, limited packet/connection drop; this is expected behavior. In some cases, the dropped packets can hang connections; for example, dropping a FIN/ACK packet for an FTP connection will make the FTP client hang. In this case, you need to reestablish the FTP connection.
- If you use a Windows 2003 server connected to a Spanned EtherChannel interface, when the syslog server port is down, and the server does not throttle ICMP error messages, then large numbers of ICMP messages are sent back to the cluster. These messages can result in some units of the cluster experiencing high CPU, which can affect performance. We recommend that you throttle ICMP error messages.
- We recommend connecting EtherChannels to a VSS, vPC, StackWise, or StackWise Virtual for redundancy.
- Within a chassis, you cannot cluster some security modules and run other security modules in standalone mode; you must include all security modules in the cluster.

Defaults

- The cluster health check feature is enabled by default with the holdtime of 3 seconds. Interface health monitoring is enabled on all interfaces by default.
- Connection rebalancing is disabled by default. If you enable connection rebalancing, the default time between load information exchanges is 5 seconds.
- The cluster auto-rejoin feature for a failed cluster control link is set to unlimited attempts every 5 minutes.
- The cluster auto-rejoin feature for a failed data interface is set to 3 attempts every 5 minutes, with the increasing interval set to 2.
- Connection replication delay of 5 seconds is enabled by default for HTTP traffic.

Configure Clustering on the Firepower 4100/9300 Chassis

You can easily deploy the cluster from the Firepower 4100/9300 chassis supervisor. All initial configuration is automatically generated for each unit. This section describes the default bootstrap configuration and optional customization you can perform on the ASA. This section also describes how to manage cluster members from within the ASA. You can also manage cluster membership from the Firepower 4100/9300 chassis. See the Firepower 4100/9300 chassis documentation for more information.

Procedure

-
- | | |
|---------------|--|
| Step 1 | FXOS: Add an ASA Cluster, on page 16 |
| Step 2 | ASA: Change the Firewall Mode and Context Mode, on page 26 |
| Step 3 | ASA: Configure Data Interfaces, on page 26 |
| Step 4 | ASA: Customize the Cluster Configuration, on page 29 |

Step 5 [ASA: Manage Cluster Members, on page 50](#)

FXOS: Add an ASA Cluster

You can add a single Firepower 9300 chassis as an intra-chassis cluster, or add multiple chassis for inter-chassis clustering. For inter-chassis clustering, you must configure each chassis separately. Add the cluster on one chassis; you can then enter most of the same settings on the next chassis.

Create an ASA Cluster

Set the scope to the image version.

You can easily deploy the cluster from the Firepower 4100/9300 chassis supervisor. All initial configuration is automatically generated for each unit.

For clustering on multiple chassis, you must configure each chassis separately. Deploy the cluster on one chassis; you can then copy the bootstrap configuration from the first chassis to the next chassis for ease of deployment.

In a Firepower 9300 chassis, you must enable clustering for all 3 module slots, even if you do not have a module installed. If you do not configure all 3 modules, the cluster will not come up.

For multiple context mode, you must first deploy the logical device, and then enable multiple context mode in the ASA application.

When you deploy a cluster, the Firepower 4100/9300 chassis supervisor configures each ASA application with the following bootstrap configuration. You can later modify parts of the bootstrap configuration from the ASA, if desired (shown in **Bold** text).

```
interface Port-channel48
  description Clustering Interface
  cluster group <service_type_name>
  key <secret>
  local-unit unit-<chassis#-module#>
  site-id <number>
  cluster-interface port-channel48 ip 127.2.<chassis#>.<module#> 255.255.255.0
  priority <auto>
  health-check holdtime 3
  health-check data-interface auto-rejoin 3 5 2
  health-check cluster-interface auto-rejoin unlimited 5 1
  enable

ip local pool cluster_ipv4_pool <ip_address>-<ip_address> mask <mask>

interface <management_ifc>
  management-only individual
  nameif management
  security-level 0
  ip address <ip_address> <mask> cluster-pool cluster_ipv4_pool
  no shutdown

http server enable
http 0.0.0.0 0.0.0.0 management
route management <management_host_ip> <mask> <gateway_ip> 1
```



Note The **local-unit** name can only be changed if you disable clustering.

Before you begin

- Download the application image you want to use for the logical device from Cisco.com, and then upload that image to the Firepower 4100/9300 chassis.
- Gather the following information:
 - Management interface ID, IP address, and network mask
 - Gateway IP address

Procedure

Step 1 Configure interfaces.

- a) Add at least one Data type interface or EtherChannel (also known as a port-channel) before you deploy the cluster. See [Add an EtherChannel \(Port Channel\)](#) or [Configure a Physical Interface](#).

For clustering on multiple chassis, all data interfaces must be Spanned EtherChannels with at least one member interface. Add the same EtherChannels on each chassis. Combine the member interfaces from all cluster units into a single EtherChannel on the switch. See [Clustering Guidelines and Limitations, on page 11](#) for more information about EtherChannels.

- b) Add a Management type interface or EtherChannel. See [Add an EtherChannel \(Port Channel\)](#) or [Configure a Physical Interface](#).

The management interface is required. Note that this management interface is not the same as the chassis management interface that is used only for chassis management (in FXOS, you might see the chassis management interface displayed as MGMT, management0, or other similar names).

For clustering on multiple chassis, add the same Management interface on each chassis.

- c) For clustering on multiple chassis, add a member interface to the cluster control link EtherChannel (by default, port-channel 48). See [Add an EtherChannel \(Port Channel\)](#).

Do not add a member interface for a cluster isolated to security modules within one Firepower 9300 chassis. If you add a member, the chassis assumes this cluster will be using multiple chassis, and will only allow you to use Spanned EtherChannels, for example.

Add the same member interfaces on each chassis. The cluster control link is a device-local EtherChannel on each chassis. Use separate EtherChannels on the switch per device. See [Clustering Guidelines and Limitations, on page 11](#) for more information about EtherChannels.

Step 2 Enter security services mode.

scope ssa

Example:

```
Firepower# scope ssa
```

```
Firepower /ssa #
```

Step 3 Set the application instance parameters, including the image version.

- a) View available images. Note the Version number that you want to use.

show app

Example:

```
Firepower /ssa # show app
```

Name	Version	Author	Supported Deploy Types	CSP Type	Is Default
App					
-----	-----	-----	-----	-----	-----
asa	9.9.1	cisco	Native	Application	No
asa	9.10.1	cisco	Native	Application	Yes
ftd	6.2.3	cisco	Native	Application	Yes
ftd	6.3.0	cisco	Native,Container	Application	Yes

- b) Set the scope to the image version.

scope app asa application_version

Example:

```
Firepower /ssa # scope app asa 9.10.1
Firepower /ssa/app #
```

- c) Set this version as the default.

set-default

Example:

```
Firepower /ssa/app # set-default
Firepower /ssa/app* #
```

- d) Exit to ssa mode.

exit

Example:

```
Firepower /ssa/app* # exit
Firepower /ssa* #
```

Example:

```
Firepower /ssa # scope app asa 9.12.1
Firepower /ssa/app # set-default
Firepower /ssa/app* # exit
Firepower /ssa* #
```

Step 4 Create the cluster.

enter logical-device device_name asa slots clustered

- *device_name*—Used by the Firepower 4100/9300 chassis supervisor to configure clustering settings and assign interfaces; it is not the cluster name used in the security module configuration. You must specify all three security modules, even if you have not yet installed the hardware.
- *slots*—Assigns the chassis modules to the cluster. For the Firepower 4100, specify **1**. For the Firepower 9300, specify **1,2,3**. You must enable clustering for all 3 module slots in a Firepower 9300 chassis, even if you do not have a module installed. If you do not configure all 3 modules, the cluster will not come up.

Example:

```
Firepower /ssa # enter logical-device ASA1 asa 1,2,3 clustered
Firepower /ssa/logical-device* #
```

Step 5 Configure the cluster bootstrap parameters.

These settings are meant for initial deployment only, or for disaster recovery. For normal operation, you can later change most values in the application CLI configuration.

- a) Create the cluster bootstrap object.

enter cluster-bootstrap

Example:

```
Firepower /ssa/logical-device* # enter cluster-bootstrap
Firepower /ssa/logical-device/cluster-bootstrap* #
```

- b) Set the chassis ID.

set chassis-id id

Each chassis in the cluster needs a unique ID.

- c) For inter-site clustering, set the site ID between 1 and 8.

set site-id number.

To remove the site ID, set the value to **0**.

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set site-id 1
Firepower /ssa/logical-device/cluster-bootstrap* #
```

- d) Configure an authentication key for control traffic on the cluster control link.

set key

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set key
Key: diamonddogs
```

You are prompted to enter the shared secret.

The shared secret is an ASCII string from 1 to 63 characters. The shared secret is used to generate the key. This option does not affect datapath traffic, including connection state update and forwarded packets, which are always sent in the clear.

- e) Set the cluster interface mode.

set mode spanned-etherchannel

Spanned EtherChannel mode is the only supported mode.

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set mode spanned-etherchannel
Firepower /ssa/logical-device/cluster-bootstrap* #
```

- f) Set the cluster group name in the security module configuration.

set service-type cluster_name

The name must be an ASCII string from 1 to 38 characters.

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set service-type cluster1
Firepower /ssa/logical-device/cluster-bootstrap* #
```

- g) (Optional) Set the cluster control link IP network.

set cluster-control-link network a.b.0.0

By default, the cluster control link uses the 127.2.0.0/16 network. However, some networking deployments do not allow 127.2.0.0/16 traffic to pass. In this case, you can specify a /16 address on a unique network for the cluster.

- **a.b.0.0**—Specify any /16 network address, except for loopback (127.0.0.0/8) and multicast (224.0.0.0/4) addresses. If you set the value to 0.0.0.0, then the default network is used: 127.2.0.0.

The chassis auto-generates the cluster control link interface IP address for each unit based on the chassis ID and slot ID: *a.b.chassis_id.slot_id*.

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set cluster-control-link network
10.10.0.0
```

- h) Configure the management IP address information.

This information is used to configure a management interface in the security module configuration.

1. Configure a pool of Local IP addresses, one of which will be assigned to each cluster unit for the interface.

set ipv4 pool start_ip end_ip

set ipv6 pool start_ip end_ip

Include at least as many addresses as there are units in the cluster. Note that for the Firepower 9300, you must include 3 addresses per chassis, even if you do not have all module slots filled. If you plan to expand the cluster, include additional addresses. The Virtual IP address (known as the Main cluster

IP address) that belongs to the current control unit is *not* a part of this pool; be sure to reserve an IP address on the same network for the Main cluster IP address. You can use IPv4 and/or IPv6 addresses.

2. Configure the Main cluster IP address for the management interface.

```
set virtual ipv4 ip_address mask mask
```

```
set virtual ipv6 ip_address prefix-length prefix
```

This IP address must be on the same network as the cluster pool addresses, but not be part of the pool.

3. Enter the network gateway address.

```
set ipv4 gateway ip_address
```

```
set ipv6 gateway ip_address
```

Example:

```
Firepower /ssa/logical-device/cluster-bootstrap* # set ipv4 gateway 10.1.1.254
Firepower /ssa/logical-device/cluster-bootstrap* # set ipv4 pool 10.1.1.11 10.1.1.27
Firepower /ssa/logical-device/cluster-bootstrap* # set ipv6 gateway 2001:DB8::AA
Firepower /ssa/logical-device/cluster-bootstrap* # set ipv6 pool 2001:DB8::11 2001:DB8::27
Firepower /ssa/logical-device/cluster-bootstrap* # set virtual ipv4 10.1.1.1 mask
255.255.255.0
Firepower /ssa/logical-device/cluster-bootstrap* # set virtual ipv6 2001:DB8::1
prefix-length 64
```

- i) Exit the cluster bootstrap mode.

```
exit
```

Example:

```
Firepower /ssa/logical-device* # enter cluster-bootstrap
Firepower /ssa/logical-device/cluster-bootstrap* # set chassis-id 1
Firepower /ssa/logical-device/cluster-bootstrap* # set key
Key: f@arscape
Firepower /ssa/logical-device/cluster-bootstrap* # set mode spanned-etherchannel
Firepower /ssa/logical-device/cluster-bootstrap* # set service-type cluster1
Firepower /ssa/logical-device/cluster-bootstrap* # exit
Firepower /ssa/logical-device/* #
```

Step 6 Configure the management bootstrap parameters.

These settings are meant for initial deployment only, or for disaster recovery. For normal operation, you can later change most values in the application CLI configuration.

- a) Create the management bootstrap object.

```
enter mgmt-bootstrap asa
```

Example:

```
Firepower /ssa/logical-device* # enter mgmt-bootstrap asa
Firepower /ssa/logical-device/mgmt-bootstrap* #
```

- b) Specify the admin and enable password.

create bootstrap-key-secret PASSWORD**set value**

Enter a value: *password*

Confirm the value: *password*

exit**Example:**

The pre-configured ASA admin user and enable password is useful for password recovery; if you have FXOS access, you can reset the admin user password if you forget it.

Example:

```
Firepower /ssa/logical-device/mgmt-bootstrap* # create bootstrap-key-secret PASSWORD
Firepower /ssa/logical-device/mgmt-bootstrap/bootstrap-key-secret* # set value
Enter a value: floppylampshade
Confirm the value: floppylampshade
Firepower /ssa/logical-device/mgmt-bootstrap/bootstrap-key-secret* # exit
Firepower /ssa/logical-device/mgmt-bootstrap* #
```

- c) Specify the firewall mode, routed or transparent.

create bootstrap-key FIREWALL_MODE**set value {routed | transparent}****exit**

In routed mode, the device is considered to be a router hop in the network. Each interface that you want to route between is on a different subnet. A transparent firewall, on the other hand, is a Layer 2 firewall that acts like a “bump in the wire,” or a “stealth firewall,” and is not seen as a router hop to connected devices.

The firewall mode is only set at initial deployment. If you re-apply the bootstrap settings, this setting is not used.

Example:

```
Firepower /ssa/logical-device/mgmt-bootstrap* # create bootstrap-key FIREWALL_MODE
Firepower /ssa/logical-device/mgmt-bootstrap/bootstrap-key* # set value routed
Firepower /ssa/logical-device/mgmt-bootstrap/bootstrap-key* # exit
Firepower /ssa/logical-device/mgmt-bootstrap* #
```

- d) Exit the management bootstrap mode.

exit**Example:**

```
Firepower /ssa/logical-device/mgmt-bootstrap* # exit
Firepower /ssa/logical-device* #
```

Step 7 Save the configuration.**commit-buffer**

The chassis deploys the logical device by downloading the specified software version and pushing the bootstrap configuration and management interface settings to the application instance. Check the status of the deployment using the **show app-instance** command. The application instance is running and ready to use when the **Admin State** is **Enabled** and the **Oper State** is **Online**.

Example:

```
Firepower /ssa/logical-device* # commit-buffer
Firepower /ssa/logical-device # exit
Firepower /ssa # show app-instance
```

App Name	Identifier	Slot ID	Admin State	Oper State	Running Version	Startup Version
Deploy Type	Profile Name	Cluster	State	Cluster Role		
ftd	cluster1	1	Enabled	Online	7.3.0.49	7.3.0.49
Native			In Cluster	Data Node		
ftd	cluster1	2	Enabled	Online	7.3.0.49	7.3.0.49
Native			In Cluster	Control Node		
ftd	cluster1	3	Disabled	Not Available		7.3.0.49
Native			Not Applicable	None		

Step 8 To add another chassis to the cluster, repeat this procedure except you must configure a unique **chassis-id** and the correct **site-id**; otherwise, use the same configuration for both chassis.

Make sure the interface configuration is the same on the new chassis. You can export and import FXOS chassis configuration to make this process easier.

Step 9 Connect to the control unit ASA to customize your clustering configuration.

Example

For chassis 1:

```
scope eth-uplink
  scope fabric a
    enter port-channel 1
    set port-type data
    enable
    enter member-port Ethernet1/1
    exit
    enter member-port Ethernet1/2
    exit
    exit
  enter port-channel 2
  set port-type data
  enable
  enter member-port Ethernet1/3
  exit
  enter member-port Ethernet1/4
  exit
  exit
enter port-channel 3
set port-type data
enable
enter member-port Ethernet1/5
exit
enter member-port Ethernet1/6
```

```

        exit
    exit
enter port-channel 4
    set port-type mgmt
    enable
    enter member-port Ethernet2/1
        exit
    enter member-port Ethernet2/2
        exit
    exit
enter port-channel 48
    set port-type cluster
    enable
    enter member-port Ethernet2/3
        exit
    exit
exit
exit
commit-buffer

scope ssa
    enter logical-device ASA1 asa "1,2,3" clustered
    enter cluster-bootstrap
        set chassis-id 1
        set ipv4 gateway 10.1.1.254
        set ipv4 pool 10.1.1.11 10.1.1.27
        set ipv6 gateway 2001:DB8::AA
        set ipv6 pool 2001:DB8::11 2001:DB8::27
        set key
        Key: f@arscape
        set mode spanned-etherchannel
        set service-type cluster1
        set virtual ipv4 10.1.1.1 mask 255.255.255.0
        set virtual ipv6 2001:DB8::1 prefix-length 64
    exit
exit
scope app asa 9.5.2.1
    set-default
    exit
commit-buffer

```

For chassis 2:

```

scope eth-uplink
    scope fabric a
        create port-channel 1
            set port-type data
            enable
            create member-port Ethernet1/1
                exit
            create member-port Ethernet1/2
                exit
        exit
        create port-channel 2
            set port-type data
            enable
            create member-port Ethernet1/3
                exit
            create member-port Ethernet1/4
                exit
        exit
        create port-channel 3
            set port-type data

```

```

    enable
    create member-port Ethernet1/5
    exit
    create member-port Ethernet1/6
    exit
    exit
create port-channel 4
    set port-type mgmt
    enable
    create member-port Ethernet2/1
    exit
    create member-port Ethernet2/2
    exit
    exit
create port-channel 48
    set port-type cluster
    enable
    create member-port Ethernet2/3
    exit
    exit
exit
exit
commit-buffer

scope ssa
    enter logical-device ASA1 asa "1,2,3" clustered
    enter cluster-bootstrap
        set chassis-id 2
        set ipv4 gateway 10.1.1.254
        set ipv4 pool 10.1.1.11 10.1.1.15
        set ipv6 gateway 2001:DB8::AA
        set ipv6 pool 2001:DB8::11 2001:DB8::19
        set key
        Key: f@rscape
        set mode spanned-etherchannel
        set service-type cluster1
        set virtual ipv4 10.1.1.1 mask 255.255.255.0
        set virtual ipv6 2001:DB8::1 prefix-length 64
    exit
scope app asa 9.5.2.1
    set-default
    exit
commit-buffer

```

Add More Cluster Members

Add or replace the ASA cluster member.



Note This procedure only applies to adding or replacing a *chassis*; if you are adding or replacing a module to a Firepower 9300 where clustering is already enabled, the module will be added automatically.

Before you begin

- Make sure your existing cluster has enough IP addresses in the management IP address pool for this new member. If not, you need to edit the existing cluster bootstrap configuration on each chassis before you add this new member. This change causes a restart of the logical device.

- The interface configuration must be the same on the new chassis. You can export and import FXOS chassis configuration to make this process easier.
- For multiple context mode, enable multiple context mode in the ASA application on the first cluster member; additional cluster members will inherit the multiple context mode configuration automatically.

Procedure

Step 1 Click **OK**.

Step 2 To add another chassis to the cluster, repeat the procedure in [Create an ASA Cluster, on page 16](#) except you must configure a unique **chassis-id** and the correct **site-id**; otherwise, use the same configuration for the new chassis.

ASA: Change the Firewall Mode and Context Mode

By default, the FXOS chassis deploys a cluster in routed firewall mode, and single context mode.

- Change the firewall mode— To change the mode after you deploy, change the mode on the control unit; the mode is automatically changed on all data units to match. See [Set the Firewall Mode](#). In multiple context mode, you set the firewall mode per context.
- Change to multiple context mode—To change to multiple context mode after you deploy, change the mode on the control unit; the mode is automatically changed on all data units to match. See [Enable Multiple Context Mode](#).

ASA: Configure Data Interfaces

This procedure configures basic parameters for each data interface that you assigned to the cluster when you deployed it in FXOS. For clustering on multiple chassis, data interfaces are always Spanned EtherChannel interfaces.



Note The management interface was pre-configured when you deployed the cluster. You can also change the management interface parameters in ASA, but this procedure focuses on data interfaces. The management interface is an individual interface, as opposed to a Spanned interface. See [Management Interface, on page 5](#) for more information.

Before you begin

- For multiple context mode, start this procedure in the system execution space. If you are not already in the System configuration mode, enter the **changeto system** command.
- For transparent mode, configure the bridge group. See [Configure the Bridge Virtual Interface \(BVI\)](#).

- When using Spanned EtherChannels for a cluster with multiple chassis, the port-channel interface will not come up until clustering is fully enabled. This requirement prevents traffic from being forwarded to a node that is not an active node in the cluster.

Procedure

Step 1 Specify the interface ID.

interface *id*

Refer to the FXOS chassis for the interfaces assigned to this cluster. The interface ID can be:

- **port-channel** *integer*
- **ethernet** *slot/port*

Example:

```
ciscoasa(config)# interface port-channel 1
```

Step 2 Enable the interface:

no shutdown

Step 3 (Optional) If you are creating VLAN subinterfaces on this interface, do so now.

Example:

```
ciscoasa(config)# interface port-channel 1.10
ciscoasa(config-if)# vlan 10
```

The rest of this procedure applies to the subinterfaces.

Step 4 (Multiple Context Mode) Allocate the interface to a context, then change to the context and enter interface mode.

Example:

```
ciscoasa(config)# context admin
ciscoasa(config)# allocate-interface port-channel1
ciscoasa(config)# changeto context admin
ciscoasa(config-if)# interface port-channel 1
```

For multiple context mode, the rest of the interface configuration occurs within each context.

Step 5 Name the interface:

nameif *name*

Example:

```
ciscoasa(config-if)# nameif inside
```

The *name* is a text string up to 48 characters, and is not case-sensitive. You can change the name by reentering this command with a new value.

Step 6 Perform one of the following, depending on the firewall mode.

- Routed Mode—Set the IPv4 and/or IPv6 address:

(IPv4)

ip address *ip_address* [*mask*]

(IPv6)

ipv6 address *ipv6-prefix/prefix-length*

Example:

```
ciscoasa(config-if)# ip address 10.1.1.1 255.255.255.0
ciscoasa(config-if)# ipv6 address 2001:DB8::1001/32
```

DHCP, PPPoE, and IPv6 autoconfig are not supported. For point-to-point connections, you can specify a 31-bit subnet mask (255.255.255.254). In this case, no IP addresses are reserved for the network or broadcast addresses. Manually configuring the link-local address is also not supported.

- Transparent Mode—Assign the interface to a bridge group:

bridge-group *number*

Example:

```
ciscoasa(config-if)# bridge-group 1
```

Where *number* is an integer between 1 and 100. You can assign up to 64 interfaces to a bridge group. You cannot assign the same interface to more than one bridge group. Note that the BVI configuration includes the IP address.

Step 7 Set the security level:

security-level *number*

Example:

```
ciscoasa(config-if)# security-level 50
```

Where *number* is an integer between 0 (lowest) and 100 (highest).

Step 8 (Clustering on multiple chassis) Configure a unique, manual global MAC address for a Spanned EtherChannel to avoid potential network connectivity problems.

mac-address *mac_address*

- *mac_address*—The MAC address is in H.H.H format, where H is a 16-bit hexadecimal digit. For example, the MAC address 00-0C-F1-42-4C-DE is entered as 000C.F142.4CDE. The first two bytes of a manual MAC address cannot be A2 if you also want to use auto-generated MAC addresses.

You must configure a unique MAC address not currently in use on your network. With a manually-configured MAC address, the MAC address stays with the current control unit. If you do not configure a MAC address,

then if the control unit changes, the new control unit uses a new MAC address for the interface, which can cause a temporary network outage.

In multiple context mode, if you share an interface between contexts, you should instead enable auto-generation of MAC addresses so you do not need to set the MAC address manually. Note that you must manually configure the MAC address using this command for *non-shared* interfaces.

Example:

```
ciscoasa(config-if)# mac-address 000C.F142.4CDE
```

- Step 9** (Inter-site clustering) Configure a site-specific MAC address and, for routed mode, an IP address for each site:

mac-address *mac_address* **site-id** *number* **site-ip** *ip_address*

Example:

```
ciscoasa(config-if)# mac-address aaaa.1111.1234
ciscoasa(config-if)# mac-address aaaa.1111.aaaa site-id 1 site-ip 10.9.9.1
ciscoasa(config-if)# mac-address aaaa.1111.bbbb site-id 2 site-ip 10.9.9.2
ciscoasa(config-if)# mac-address aaaa.1111.cccc site-id 3 site-ip 10.9.9.3
ciscoasa(config-if)# mac-address aaaa.1111.dddd site-id 4 site-ip 10.9.9.4
```

The site-specific IP addresses must be on the same subnet as the global IP address. The site-specific MAC address and IP address used by a unit depends on the site ID you specify in each unit's bootstrap configuration.

ASA: Customize the Cluster Configuration

If you want to change bootstrap settings after you deploy the cluster or configure additional options, such as clustering health monitoring, TCP connection replication delay, flow mobility, and other optimizations, you can do so on the control unit.

Configure Basic ASA Cluster Parameters

You can customize cluster settings on the control node.

Before you begin

- For multiple context mode, complete this procedure in the system execution space on the control unit. To change from the context to the system execution space, enter the **changeto system** command.
- The local-unit name and several other options can only be set on the FXOS chassis, or they can only be changed on the ASA if you disable clustering, so they are not included in the following procedure.

Procedure

- Step 1** Confirm that this unit is the control unit:
- show cluster info**

Example:

```
asa(config)# show cluster info
Cluster cluster1: On
  Interface mode: spanned
  This is "unit-1-2" in state CONTROL_NODE
    ID      : 2
    Version : 9.5(2)
    Serial No.: FCH183770GD
    CCL IP   : 127.2.1.2
    CCL MAC  : 0015.c500.019f
    Last join : 01:18:34 UTC Nov 4 2015
    Last leave: N/A
  Other members in the cluster:
    Unit "unit-1-3" in state DATA_NODE
      ID      : 4
      Version : 9.5(2)
      Serial No.: FCH19057ML0
      CCL IP   : 127.2.1.3
      CCL MAC  : 0015.c500.018f
      Last join : 20:29:57 UTC Nov 4 2015
      Last leave: 20:24:55 UTC Nov 4 2015
    Unit "unit-1-1" in state DATA_NODE
      ID      : 1
      Version : 9.5(2)
      Serial No.: FCH19057ML0
      CCL IP   : 127.2.1.1
      CCL MAC  : 0015.c500.017f
      Last join : 20:20:53 UTC Nov 4 2015
      Last leave: 20:18:15 UTC Nov 4 2015
    Unit "unit-2-1" in state DATA_NODE
      ID      : 3
      Version : 9.5(2)
      Serial No.: FCH19057ML0
      CCL IP   : 127.2.2.1
      CCL MAC  : 0015.c500.020f
      Last join : 20:19:57 UTC Nov 4 2015
      Last leave: 20:24:55 UTC Nov 4 2015
```

If a different unit is the control unit, exit the connection and connect to the correct unit.

- Step 2** Specify the maximum transmission unit for the cluster control link interface to be at least 100 bytes higher than the highest MTU of the data interfaces.

mtu cluster bytes**Example:**

```
ciscoasa(config)# mtu cluster 9184
```

We suggest setting the cluster control link MTU to the maximum; the minimum value is 1400 bytes. In addition, we do not recommend setting the cluster control link MTU between 2561 and 8362; due to block pool handling, this MTU size is not optimal for system operation. Because the cluster control link traffic includes data packet forwarding, the cluster control link needs to accommodate the entire size of a data packet plus cluster traffic overhead. For example, because the maximum MTU is 9184, then the highest data interface MTU can be 9084, while the cluster control link can be set to 9184. When a node joins the cluster, it checks MTU compatibility by sending a ping to the control node with a packet size matching the cluster control link

MTU. If the ping fails, a notification is generated so you can fix the MTU mismatch on connecting switches and try again.

Step 3 Enter cluster configuration mode:

cluster group *name*

Step 4 (Optional) Enable console replication from data units to the control unit:

console-replicate

This feature is disabled by default. The ASA prints out some messages directly to the console for certain critical events. If you enable console replication, data units send the console messages to the control unit so that you only need to monitor one console port for the cluster.

Step 5 Set the minimum trace level for clustering events:

trace-level *level*

Set the minimum level as desired:

- **critical**—Critical events (severity=1)
- **warning**—Warnings (severity=2)
- **informational**—Informational events (severity=3)
- **debug**—Debugging events (severity=4)

Step 6 (Optional) (Firepower 9300 only) Ensure that the security modules in a chassis join the cluster simultaneously, so that traffic is evenly distributed between the modules. If a module joins very much in advance of other modules, it can receive more traffic than desired, because the other modules cannot yet share the load.

unit parallel-join *num_of_units* **max-bundle-delay** *max_delay_time*

- *num_of_units*—Specifies the minimum number of modules in the same chassis required to be ready before a module can join the cluster, between 1 and 3. The default is 1, meaning that a module will not wait for other modules to be ready before it joins the cluster. If you set the value to 3, for example, then each module will wait the *max_delay_time* or until all 3 modules are ready before joining the cluster. All 3 modules will request to join the cluster roughly simultaneously, and will all start receiving traffic around the same time.
- *max_delay_time*—Specifies the maximum delay time in minutes before a module stops waiting for other modules to be ready before it joins the cluster, between 0 and 30 minutes. The default is 0, meaning the module will not wait for other modules to be ready before it joins the cluster. If you set the *num_of_units* to 1, then this value must be 0. If you set the *num_of_units* to 2 or 3, then this value must be 1 or more. This timer is per module, but when the first module joins the cluster, then all other module timers end, and the remaining modules join the cluster.

For example, you set the *num_of_units* to 3, and the *max_delay_time* to 5 minutes. When module 1 comes up, it starts its 5 minute timer. Module 2 comes up 2 minutes later and starts its 5 minute timer. Module 3 comes up 1 minute later, therefore all modules will now join the cluster at the 4 minute mark; they will not wait for the timers to complete. If module 3 never comes up, then Module 1 will join the cluster at the end of its 5 minute timer, and Module 2 will also join, even though its timer still has 2 minutes remaining; it will not wait for its timer to complete.

Step 7 Configure the maximum number of cluster members.

cluster-member-limit *number*

- *number*—2 to 16. The default is 16.

If you know that your cluster will be fewer than the maximum of 16 units, then we recommend that you set the actual planned number of units. Setting the maximum units lets the cluster manage resources better. For example, if you use port address translation (PAT), then the control unit can allocate port blocks to the planned number of members, and it will not have to reserve ports for extra units you don't plan to use.

- Step 8** Set the keepalive interval for flow state refresh messages (clu_heartbeat and clu_update messages) from the flow owner to the director and backup owner.

clu-keepalive-interval *seconds*

- *seconds*—15 to 55. The default is 15.

You may want to set the interval to be longer than the default to reduce the amount of traffic on the cluster control link.

Configure Health Monitoring and Auto-Rejoin Settings

This procedure configures unit and interface health monitoring.

You might want to disable health monitoring of non-essential interfaces, for example, the management interface. You can monitor any port-channel ID or single physical interface ID. Health monitoring is not performed on VLAN subinterfaces or virtual interfaces such as VNIs or BVIs. You cannot configure monitoring for the cluster control link; it is always monitored.

Procedure

-
- Step 1** Enter cluster configuration mode:

cluster group *name*

- Step 2** Customize the cluster unit health check feature:

health-check [**holdtime** *timeout*]

The **holdtime** determines the amount of time between unit heartbeat status messages, between .3 and 45 seconds; The default is 3 seconds.

To determine unit health, the ASA cluster units send heartbeat messages on the cluster control link to other units. If a unit does not receive any heartbeat messages from a peer unit within the holdtime period, the peer unit is considered unresponsive or dead.

When any topology changes occur (such as adding or removing a data interface, enabling or disabling an interface on the ASA, Firepower 4100/9300 chassis, or the switch, or adding an additional switch to form a VSS, vPC, StackWise, or StackWise Virtual) you should disable the health check feature and also disable interface monitoring for the disabled interfaces (**no health-check monitor-interface**). When the topology change is complete, and the configuration change is synced to all units, you can re-enable the health check feature.

Example:

```
ciscoasa(cfg-cluster)# health-check holdtime 5
```

Step 3 Disable the interface health check on an interface:

no health-check monitor-interface [*interface_id* | **service-application**]

The interface health check monitors for link failures. If all physical ports for a given logical interface fail on a particular unit, but there are active ports under the same logical interface on other units, then the unit is removed from the cluster. The amount of time before the ASA removes a member from the cluster depends on the type of interface and whether the unit is an established member or is joining the cluster.

Health check is enabled by default for all interfaces. You can disable it per interface using the **no** form of this command. You might want to disable health monitoring of non-essential interfaces, for example, the management interface. Health monitoring is not performed on VLAN subinterfaces or virtual interfaces such as VNIs or BVIs. You cannot configure monitoring for the cluster control link; it is always monitored. Specify the **service-application** to disable monitoring of a decorator application.

When any topology changes occur (such as adding or removing a data interface, enabling or disabling an interface on the ASA, Firepower 4100/9300 chassis, or the switch, or adding an additional switch to form a VSS, vPC, StackWise, or StackWise Virtual) you should disable the health check feature (**no health-check**) and also disable interface monitoring for the disabled interfaces. When the topology change is complete, and the configuration change is synced to all units, you can re-enable the health check feature.

Example:

```
ciscoasa(cfg-cluster)# no health-check monitor-interface port-channel1
```

Step 4 Customize the auto-rejoin cluster settings after a health check failure:

health-check {**data-interface** | **cluster-interface** | **system**} **auto-rejoin** [**unlimited** | *auto_rejoin_max*]
auto_rejoin_interval *auto_rejoin_interval_variation*

- **system**—Specifies the auto-rejoin settings for internal errors. Internal failures include: application sync timeout; inconsistent application statuses; and so on.
- **unlimited**—(Default for the **cluster-interface**) Does not limit the number of rejoin attempts.
- *auto-rejoin-max*—Sets the number of rejoin attempts, between 0 and 65535. **0** disables auto-rejoining. The default for the **data-interface** and **system** is 3.
- *auto_rejoin_interval*—Defines the interval duration in minutes between rejoin attempts, between 2 and 60. The default value is 5 minutes. The maximum total time that the unit attempts to rejoin the cluster is limited to 14400 minutes (10 days) from the time of last failure.
- *auto_rejoin_interval_variation*—Defines if the interval duration increases. Set the value between 1 and 3: **1** (no change); **2** (2 x the previous duration), or **3** (3 x the previous duration). For example, if you set the interval duration to 5 minutes, and set the variation to 2, then the first attempt is after 5 minutes; the 2nd attempt is 10 minutes (2 x 5); the 3rd attempt 20 minutes (2 x 10), and so on. The default value is **1** for the **cluster-interface** and **2** for the **data-interface** and **system**.

Example:

```
ciscoasa(cfg-cluster)# health-check data-interface auto-rejoin 10 3 3
```

Step 5 Set the chassis rejoin to match the **health-check system auto-rejoin** command for chassis heartbeat failures.

health-check chassis-heartbeat-delay-rejoin

By default, if the chassis heartbeat fails and then recovers, the node rejoins the cluster immediately. However, if you configure the **health-check chassis-heartbeat-delay-rejoin** command, it will rejoin according to the settings of the **health-check system auto-rejoin** command.

Step 6 Configure the debounce time before the ASA considers an interface to be failed and the unit is removed from the cluster.

health-check monitor-interface debounce-time *ms*

Set the debounce time between 300 and 9000 ms. The default is 500 ms. Lower values allow for faster detection of interface failures. Note that configuring a lower debounce time increases the chances of false-positives. When an interface status update occurs, the ASA waits the number of milliseconds specified before marking the interface as failed and the unit is removed from the cluster. In the case of an EtherChannel that transitions from a down state to an up state (for example, the switch reloaded, or the switch enabled an EtherChannel), a longer debounce time can prevent the interface from appearing to be failed on a cluster unit just because another cluster unit was faster at bundling the ports.

Example:

```
ciscoasa(cfg-cluster)# health-check monitor-interface debounce-time 300
```

Step 7 Configure the chassis health check interval:

app-agent heartbeat [*interval ms*] [*retry-count number*]

- **interval *ms***—Set the amount of time between heartbeats, between 100 and 6000 ms, in multiples of 100. The default is 1000 ms.
- **retry-count *number***—Set the number of retries, between 1 and 30. The default is 3 retries.

The ASA checks whether it can communicate over the backplane with the host chassis.

The minimum combined time (*interval x retry-count*) cannot be less than 600 ms. For example, if you set the interval to 100, and the retry count to 3, then the total combined time is 300 ms, which is not supported. For example, you can set the interval to 100, and the retry count to 6 to meet the minimum time (600 ms).

Example:

```
ciscoasa(cfg-cluster)# app-agent heartbeat interval 1000 retry-count 10
```

Step 8 (Optional) Configure traffic load monitoring.

load-monitor [*frequency seconds*] [*intervals intervals*]

- **frequency *seconds***—Sets the time in seconds between monitoring messages, between 10 and 360 seconds. The default is 20 seconds.
- **intervals *intervals***—Sets the number of intervals for which the ASA maintains data, between 1 and 60. The default is 30.

You can monitor the traffic load for cluster members, including total connection count, CPU and memory usage, and buffer drops. If the load is too high, you can choose to manually disable clustering on the unit if the remaining units can handle the load, or adjust the load balancing on the external switch. This feature is enabled by default. For example, for inter-chassis clustering on the Firepower 9300 with 3 security modules

in each chassis, if 2 security modules in a chassis leave the cluster, then the same amount of traffic to the chassis will be sent to the remaining module and potentially overwhelm it. You can periodically monitor the traffic load. If the load is too high, you can choose to manually disable clustering on the unit.

Use the **show cluster info load-monitor** command to view the traffic load.

Example:

```
ciscoasa(cfg-cluster)# load-monitor frequency 50 intervals 25
ciscoasa(cfg-cluster)# show cluster info load-monitor
ID  Unit Name
0   B
1   A_1
Information from all units with 50 second interval:
Unit      Connections    Buffer Drops    Memory Used    CPU Used
Average from last 1 interval:
0          0              0              14             25
1          0              0              16             20
Average from last 25 interval:
0          0              0              12             28
1          0              0              13             27
```

Configure Connection Rebalancing and the Cluster TCP Replication Delay

You can configure connection rebalancing. You can enable the cluster replication delay for TCP connections to help eliminate the “unnecessary work” related to short-lived flows by delaying the director/backup flow creation. Note that if a unit fails before the director/backup flow is created, then those flows cannot be recovered. Similarly, if traffic is rebalanced to a different unit before the flow is created, then the flow cannot be recovered. You should not enable the TCP replication delay for traffic on which you disable TCP randomization.

Procedure

Step 1 Enter cluster configuration mode:

cluster group name

Step 2 (Optional) Enable connection rebalancing for TCP traffic:

conn-rebalance [frequency seconds]

Example:

```
ciscoasa(cfg-cluster)# conn-rebalance frequency 60
```

This command is disabled by default. If enabled, ASAs exchange information about the connections per second periodically, and offload new connections from devices with more connections per second to less loaded devices. Existing connections are never moved. Moreover, because this command only rebalances based on connections per second, the total number of established connections on each node is not considered, and the total number of connections may not be equal. The frequency, between 1 and 360 seconds, specifies how often the load information is exchanged. The default is 5 seconds.

Once a connection is offloaded to a different node, it becomes an asymmetric connection.

Do not configure connection rebalancing for inter-site topologies; you do not want new connections rebalanced to cluster members at a different site.

Step 3 Enable the cluster replication delay for TCP connections:

cluster replication delay *seconds* {**http** | **match tcp** {**host** *ip_address* | *ip_address mask* | **any** | **any4** | **any6**} [{**eq** | **lt** | **gt**} *port*] {**host** *ip_address* | *ip_address mask* | **any** | **any4** | **any6**} [{**eq** | **lt** | **gt**} *port*]}

Example:

```
ciscoasa(config)# cluster replication delay 15 match tcp any any eq ftp
ciscoasa(config)# cluster replication delay 15 http
```

Set the *seconds* between 1 and 15. The **http** delay is enabled by default for 5 seconds.

Configure Inter-Site Features

For inter-site clustering, you can customize your configuration to enhance redundancy and stability.

Enable Director Localization

To improve performance and reduce round-trip time latency for inter-site clustering for data centers, you can enable director localization. New connections are typically load-balanced and owned by cluster members within a given site. However, the ASA assigns the director role to a member at *any* site. Director localization enables additional director roles: a local director at the same site as the owner, and a global director that can be at any site. Keeping the owner and director at the same site improves performance. Also, if the original owner fails, the local director chooses a new connection owner at the same site. The global director is used if a cluster member receives packets for a connection that is owned on a different site.

Before you begin

- Set the site ID for the chassis on Firepower 4100/9300 chassis supervisor.
- The following traffic types do not support localization: NAT or PAT traffic; SCTP-inspected traffic; Fragmentation owner query.

Procedure

Step 1 Enter cluster configuration mode:

cluster group *name*

Example:

```
ciscoasa(config)# cluster group cluster1
ciscoasa(cfg-cluster)#
```

Step 2 Enable director localization:

director-localization

Enable Site Redundancy

To protect flows from a site failure, you can enable site redundancy. If the connection backup owner is at the same site as the owner, then an additional backup owner will be chosen from another site to protect flows from a site failure.

Before you begin

- Set the site ID for the chassis on Firepower 4100/9300 chassis supervisor.

Procedure

Step 1 Enter cluster configuration mode:

cluster group *name*

Example:

```
ciscoasa(config)# cluster group cluster1  
ciscoasa(cfg-cluster)#
```

Step 2 Enable site redundancy.

site-redundancy

Configure Per-Site Gratuitous ARP

The ASA generates gratuitous ARP (GARP) packets to keep the switching infrastructure up to date: the highest priority member at each site periodically generates GARP traffic for the global MAC/IP addresses.

When using per-site MAC and IP addresses, packets sourced from the cluster use a site-specific MAC address and IP address, while packets received by the cluster use a global MAC address and IP address. If traffic is not generated from the global MAC address periodically, you could experience a MAC address timeout on your switches for the global MAC address. After a timeout, traffic destined for the global MAC address will be flooded across the entire switching infrastructure, which can cause performance and security concerns.

GARP is enabled by default when you set the site ID for each unit and the site MAC address for each Spanned EtherChannel. You can customize the GARP interval, or you can disable GARP.

Before you begin

- Set the site ID for the cluster member in the bootstrap configuration.
- Set the per-site MAC address for the Spanned EtherChannel in the control unit configuration.

Procedure

Step 1 Enter cluster configuration mode.

cluster group *name*

Example:

```
ciscoasa(config)# cluster group cluster1
ciscoasa(cfg-cluster)#
```

Step 2 Customize the GARP interval.

site-periodic-garp interval *seconds*

- *seconds*—Sets the time in seconds between GARP generation, between 1 and 1000000 seconds. The default is 290 seconds.

To disable GARP, enter **no site-periodic-garp interval**.

Example:

```
ciscoasa(cfg-cluster)# site-periodic-garp interval 500
```

Configure Cluster Flow Mobility

You can inspect LISP traffic to enable flow mobility when a server moves between sites.

About LISP Inspection

You can inspect LISP traffic to enable flow mobility between sites.

About LISP

Data center virtual machine mobility such as VMware VMotion enables servers to migrate between data centers while maintaining connections to clients. To support such data center server mobility, routers need to be able to update the ingress route towards the server when it moves. Cisco Locator/ID Separation Protocol (LISP) architecture separates the device identity, or endpoint identifier (EID), from its location, or routing locator (RLOC), into two different numbering spaces, making server migration transparent to clients. For example, when a server moves to a new site and a client sends traffic to the server, the router redirects traffic to the new location.

LISP requires routers and servers in certain roles, such as the LISP egress tunnel router (ETR), ingress tunnel router (ITR), first hop routers, map resolver (MR), and map server (MS). When the first hop router for the server senses that the server is connected to a different router, it updates all of the other routers and databases so that the ITR connected to the client can intercept, encapsulate, and send traffic to the new server location.

Secure Firewall ASA LISP Support

The ASA does not run LISP itself; it can, however, inspect LISP traffic for location changes and then use this information for seamless clustering operation. Without LISP integration, when a server moves to a new site, traffic comes to an ASA cluster member at the new site instead of to the original flow owner. The new ASA

forwards traffic to the ASA at the old site, and then the old ASA has to send traffic back to the new site to reach the server. This traffic flow is sub-optimal and is known as “tromboning” or “hair-pinning.”

With LISP integration, the ASA cluster members can inspect LISP traffic passing between the first hop router and the ETR or ITR, and can then change the flow owner to be at the new site.

LISP Guidelines

- The ASA cluster members must reside between the first hop router and the ITR or ETR for the site. The ASA cluster itself cannot be the first hop router for an extended segment.
- Only fully-distributed flows are supported; centralized flows, semi-distributed flows, or flows belonging to individual nodes are not moved to new owners. Semi-distributed flows include applications, such as SIP, where all child flows are owned by the same ASA that owns the parent flow.
- The cluster only moves Layer 3 and 4 flow states; some application data might be lost.
- For short-lived flows or non-business-critical flows, moving the owner may not be worthwhile. You can control the types of traffic that are supported with this feature when you configure the inspection policy, and should limit flow mobility to essential traffic.

ASA LISP Implementation

This feature includes several inter-related configurations (all of which are described in this chapter):

1. (Optional) Limit inspected EIDs based on the host or server IP address—The first hop router might send EID-notify messages for hosts or networks the ASA cluster is not involved with, so you can limit the EIDs to only those servers or networks relevant to your cluster. For example, if the cluster is only involved with 2 sites, but LISP is running on 3 sites, you should only include EIDs for the 2 sites involved with the cluster.
2. LISP traffic inspection—The ASA inspects LISP traffic on UDP port 4342 for the EID-notify message sent between the first hop router and the ITR or ETR. The ASA maintains an EID table that correlates the EID and the site ID. For example, you should inspect LISP traffic with a source IP address of the first hop router and a destination address of the ITR or ETR. Note that LISP traffic is not assigned a director, and LISP traffic itself does not participate in cluster state sharing.
3. Service Policy to enable flow mobility on specified traffic—You should enable flow mobility on business-critical traffic. For example, you can limit flow mobility to only HTTPS traffic, and/or to traffic to specific servers.
4. Site IDs—The ASA uses the site ID for each cluster node to determine the new owner.
5. Cluster-level configuration to enable flow mobility—You must also enable flow mobility at the cluster level. This on/off toggle lets you easily enable or disable flow mobility for a particular class of traffic or applications.

Configure LISP Inspection

You can inspect LISP traffic to enable flow mobility when a server moves between sites.

Before you begin

- Set the site ID for the chassis on Firepower 4100/9300 chassis supervisor.
- LISP traffic is not included in the default-inspection-traffic class, so you must configure a separate class for LISP traffic as part of this procedure.

Procedure

- Step 1** (Optional) Configure a LISP inspection map to limit inspected EIDs based on IP address, and to configure the LISP pre-shared key:
- Create an extended ACL; only the destination IP address is matched to the EID embedded address:
access list *eid_acl_name* **extended permit ip** *source_address mask destination_address mask*
 Both IPv4 and IPv6 ACLs are accepted. See the command reference for exact **access-list extended** syntax.
 - Create the LISP inspection map, and enter parameters mode:
policy-map type inspect lisp *inspect_map_name*
parameters
 - Define the allowed EIDs by identifying the ACL you created:
allowed-eid access-list *eid_acl_name*
 The first hop router or ITR/ETR might send EID-notify messages for hosts or networks the ASA cluster is not involved with, so you can limit the EIDs to only those servers or networks relevant to your cluster. For example, if the cluster is only involved with 2 sites, but LISP is running on 3 sites, you should only include EIDs for the 2 sites involved with the cluster.
 - If necessary, enter the pre-shared key:
validate-key *key*

Example:

```
ciscoasa(config)# access-list TRACKED_EID_LISP extended permit ip any 10.10.10.0 255.255.255.0
ciscoasa(config)# policy-map type inspect lisp LISP_EID_INSPECT
ciscoasa(config-pmap)# parameters
ciscoasa(config-pmap-p)# allowed-eid access-list TRACKED_EID_LISP
ciscoasa(config-pmap-p)# validate-key MadMaxShinyandChrome
```

- Step 2** Configure LISP inspection for UDP traffic between the first hop router and the ITR or ETR on port 4342:
- Configure the extended ACL to identify LISP traffic:
access list *inspect_acl_name* **extended permit udp** *source_address mask destination_address mask eq 4342*
 You *must* specify UDP port 4342. Both IPv4 and IPv6 ACLs are accepted. See the command reference for exact **access-list extended** syntax.
 - Create a class map for the ACL:
class-map *inspect_class_name*
match access-list *inspect_acl_name*
 - Specify the policy map, the class map, enable inspection using the optional LISP inspection map, and apply the service policy to an interface (if new):
policy-map *policy_map_name*

```

class inspect_class_name
inspect lisp [inspect_map_name]
service-policy policy_map_name {global | interface ifc_name}

```

If you have an existing service policy, specify the existing policy map name. By default, the ASA includes a global policy called **global_policy**, so for a global policy, specify that name. You can also create one service policy per interface if you do not want to apply the policy globally. LISP inspection is applied to traffic bidirectionally so you do not need to apply the service policy on both the source and destination interfaces; all traffic that enters or exits the interface to which you apply the policy map is affected if the traffic matches the class map for both directions.

Example:

```

ciscoasa(config)# access-list LISP_ACL extended permit udp host 192.168.50.89 host
192.168.10.8 eq 4342
ciscoasa(config)# class-map LISP_CLASS
ciscoasa(config-cmap)# match access-list LISP_ACL
ciscoasa(config-cmap)# policy-map INSIDE_POLICY
ciscoasa(config-pmap)# class LISP_CLASS
ciscoasa(config-pmap-c)# inspect lisp LISP_EID_INSPECT
ciscoasa(config)# service-policy INSIDE_POLICY interface inside

```

The ASA inspects LISP traffic for the EID-notify message sent between the first hop router and the ITR or ETR. The ASA maintains an EID table that correlates the EID and the site ID.

Step 3

Enable Flow Mobility for a traffic class:

- a) Configure the extended ACL to identify business critical traffic that you want to re-assign to the most optimal site when servers change sites:

```
access list flow_acl_name extended permit udp source_address mask destination_address mask eq port
```

Both IPv4 and IPv6 ACLs are accepted. See the command reference for exact **access-list extended** syntax. You should enable flow mobility on business-critical traffic. For example, you can limit flow mobility to only HTTPS traffic, and/or to traffic to specific servers.

- b) Create a class map for the ACL:

```
class-map flow_map_name
match access-list flow_acl_name

```

- c) Specify the same policy map on which you enabled LISP inspection, the flow class map, and enable flow mobility:

```
policy-map policy_map_name
class flow_map_name
cluster flow-mobility lisp

```

Example:

```

ciscoasa(config)# access-list IMPORTANT-FLOWS extended permit tcp any 10.10.10.0 255.255.255.0
eq https
ciscoasa(config)# class-map IMPORTANT-FLOWS-MAP
ciscoasa(config)# match access-list IMPORTANT-FLOWS
ciscoasa(config-cmap)# policy-map INSIDE_POLICY

```

```
ciscoasa(config-pmap)# class IMPORTANT-FLOWS-MAP
ciscoasa(config-pmap-c)# cluster flow-mobility lisp
```

Step 4 Enter cluster group configuration mode, and enable flow mobility for the cluster:

cluster group name

flow-mobility lisp

This on/off toggle lets you easily enable or disable flow mobility.

Examples

The following example:

- Limits EIDs to those on the 10.10.10.0/24 network
- Inspects LISP traffic (UDP 4342) between a LISP router at 192.168.50.89 (on inside) and an ITR or ETR router (on another ASA interface) at 192.168.10.8
- Enables flow mobility for all inside traffic going to a server on 10.10.10.0/24 using HTTPS.
- Enables flow mobility for the cluster.

```
access-list TRACKED_EID_LISP extended permit ip any 10.10.10.0 255.255.255.0
policy-map type inspect lisp LISP_EID_INSPECT
  parameters
    allowed-eid access-list TRACKED_EID_LISP
    validate-key MadMaxShinyandChrome
!
access-list LISP_ACL extended permit udp host 192.168.50.89 host 192.168.10.8 eq 4342
class-map LISP_CLASS
  match access-list LISP_ACL
policy-map INSIDE_POLICY
  class LISP_CLASS
    inspect lisp LISP_EID_INSPECT
service-policy INSIDE_POLICY interface inside
!
access-list IMPORTANT-FLOWS extended permit tcp any 10.10.10.0 255.255.255.0 eq https
class-map IMPORTANT-FLOWS-MAP
  match access-list IMPORTANT-FLOWS
policy-map INSIDE_POLICY
  class IMPORTANT-FLOWS-MAP
    cluster flow-mobility lisp
!
cluster group cluster1
  flow-mobility lisp
```

Configure Distributed Site-to-Site VPN

By default, the cluster uses centralized site-to-site VPN mode. To take advantage of the scalability of clustering, you can enable distributed site-to-site VPN mode.

About Distributed Site-to-Site VPN

In distributed mode, site-to-site IPsec IKEv2 VPN connections are distributed across nodes of a cluster. Distributing VPN connections across the nodes of a cluster allows both the capacity and throughput of the cluster to be fully utilized, significantly scaling VPN support beyond centralized VPN capabilities.

Distributed VPN Connection Roles

When running in distributed VPN mode, the following roles are assigned to the cluster nodes:

- **Active Session Owner**—The node that initially receives the connection or that has transitioned a backup session to an active session. The owner maintains the state and processes packets for the complete session, including the IKE and IPsec tunnels and all traffic associated with them.
- **Backup Session Owner**—The node that is handling the backup session for an existing active session. If the active session owner fails, the backup session owner becomes the active session owner, and a new backup session is established on a different node.
- **Forwarder**—If traffic associated with a VPN session is sent to a node that does not own the VPN session, that node will use the cluster control link to forward the traffic to the node that owns the VPN session.
- **Orchestrator**—The orchestrator (always the control node of the cluster) is responsible for calculating which sessions will move and where they will move to when executing an Active Session Redistribution (ASR). It sends a request to the owner node X to move N sessions to node Y. Node X will respond back to the orchestrator when complete, specifying how many sessions it was able to move.

Distributed VPN Session Characteristics

Distributed site-to-site VPN Sessions have the following characteristics. Otherwise, VPN connections behave as they normally do if not on a cluster.

- VPN sessions are distributed across the cluster at the session level. Meaning the same cluster node handles the IKE and IPsec tunnels and all their traffic for a VPN connection. If VPN session traffic is sent to a cluster node that does not own that VPN session, traffic is forwarded to the cluster node that owns the VPN session.
- VPN sessions have a Session ID that is unique across the cluster. Using the session ID, traffic is validated, forwarding decisions are made, and IKE negotiation is completed.
- In a site-to-site VPN hub and spoke configuration, when clients connect through the cluster (called hair-pinning), the session traffic flowing in and the session traffic flowing out may be on different cluster nodes.
- You can require that the backup session to be allocated on a security module in another chassis; this provides protection against chassis failure. Or, you can choose to allocate backup sessions on any node in the cluster; this provides protection against node failure only. When there are two chassis in the cluster, remote-chassis backup is strongly recommended.

Distributed VPN Handling of Cluster Events

Event	Distributed VPN
Node failure	For all active sessions on this failed node, the backup sessions (on another node) become active, and backup sessions are reallocated on another node according to the backup strategy.

Event	Distributed VPN
Chassis failure	When a remote-chassis backup strategy is being used, for all active sessions on the failed chassis, the backup sessions (on a node in the other chassis) become active. When the nodes are replaced, backup sessions for these now active sessions will be reallocated on nodes in the replaced chassis. When a flat backup strategy is being used, if both the active and backup sessions are on the failed chassis, the connection will drop. All active sessions with backup sessions on a node in the other chassis fall back to these sessions. New backup sessions will be allocated on another node in the surviving chassis.
Inactivate a cluster node	For all active sessions on the cluster node being inactivated, backup sessions (on another node) become active and reallocate backup sessions on another node according to the backup strategy.
Cluster node join	If the VPN cluster mode on the new node is not set to distributed, the control node will request a mode change. After the VPN mode is compatible, the cluster node will be assigned active and backup sessions in the flow of normal operations.

IPsec IKEv2 Modifications

IKEv2 is modified while in distributed site-to-site VPN mode in the following ways:

- An identity is used in place of IP/port tuples. This allows for proper forwarding decisions on the packets, and cleanup of previous connections that may be on other cluster members.
- The (SPI) identifiers that identify a single IKEv2 session are locally generated, random 8-byte values that are unique across the cluster. An SPI embeds a time stamp and a cluster node ID. Upon receipt of an IKE negotiation packet, if the time stamp or cluster node ID check fails, the packet is dropped and a message is logged indicating the reason.
- IKEv2 processing has been modified to prevent NAT-T negotiations from failing by being split across cluster members. A new ASP classify domain, *cluster_isakmp_redirect*, and rules are added when IKEv2 is enabled on an interface. Use the **show asp table classify domain cluster_isakmp_redirect** command to see the rules.

High Availability for Distributed Site-to-Site VPN Within the Cluster

The following capabilities provide resiliency against single failure of a security module or chassis:

- VPN Sessions that are backed up on another security module in the cluster, on any chassis, withstand security module failures.
- VPN Sessions that are backed up on another chassis withstand chassis failures.
- The control node can change without losing VPN site-to-site sessions.

If an additional failure occurs before the cluster has stabilized, connections may be lost if the both active and backup sessions are on the failed nodes.

All attempts are made to ensure no sessions are lost when a node leaves the cluster in a graceful manner such as disabling the VPN cluster mode, reloading a cluster node, and other anticipated chassis changes. During these types of operations, sessions will not be lost as long as the cluster is given time to re-establish session

backups between operations. If a graceful exit is triggered on the last cluster node, it will gracefully tear down existing sessions.

CMPv2

The CMPv2 ID certificate and key pairs are synchronized across the cluster nodes. However, only the control node in the cluster automatically renews and rekeys the CMPv2 certificate. The control node synchronizes these new ID certificates and keys to all cluster nodes on a renewal. In this way, all nodes in the cluster utilize the CMPv2 certificates for authentication, and also any node is capable of taking over as the control node.

Licenses for Distributed Site-to-Site VPN

A Carrier license is required for distributed site-to-site VPN, on each member of the cluster.

Each VPN connection requires two *Other VPN* licensed sessions (the *Other VPN* license is part of the *Essentials* license), one for the active session and one for the backup session. The maximum VPN session capacity of the cluster can be no more than half of the licensed capacity due to using two licenses for each session.

Prerequisites for Distributed Site-to-Site VPN

Model Support

- Firepower 9300
- Maximum of 6 modules on up to 2 chassis. You can have different quantities of installed security modules in each chassis, although we recommend an equal distribution.

Maximum VPN Sessions

Each security module supports up to 6K VPN sessions for a maximum of approximately 36K sessions across 6 nodes.

The actual number of sessions supported on a cluster node is determined by platform capacity, allocated licenses, and per context resource allocation. When utilization is close to the limit, there may be cases where session creation fails, even though the maximum capacity has not been reached on each cluster node. This is because active session allocation is determined by external switching, and backup session allocation is determined by an internal cluster algorithm. Customers are encouraged to size their utilization accordingly and allow room for uneven distribution.

Guidelines for Distributed Site-to-Site VPN

Firewall Mode

Distributed site-to-site VPN is supported in routed mode only.

Context Mode

Distributed site-to-site VPN operates in both single and multiple context modes. However, in multiple context mode, active session redistribution is done at the system level, not at the context level. This prevents an active session associated with a context from moving to a cluster member that contains active sessions associated with a different context, unknowingly creating an unsupportable load.

Unsupported Inspections

The following types of inspections are not supported or are disabled in distributed site-to-site VPN mode:

- CTIQBE
- DCERPC
- H323, H225, and RAS
- IPsec pass-through
- MGCP
- MMP
- NetBIOS
- PPTP
- RADIUS
- RSH
- RTSP
- SCCP (Skinny)
- SUNRPC
- TFTP
- WAAS
- WCCP
- XDMCP

Additional Guidelines

- Only IKEv2 IPsec site-to-site VPN is supported in distributed site-to-site VPN mode. IKEv1 is not supported. IKEv1 site-to-site is supported in centralized VPN mode.
- Inter-site clustering is not supported.
- Dynamic PAT is not available while in distributed site-to-site VPN mode.

Enable Distributed Site-to-Site VPN

Enable distributed site-to-site VPN to take advantage of the scalability of clustering for VPN sessions.



Note Changing the VPN mode between centralized and distributed requires all nodes in the cluster to reload. Changing the backup mode is dynamic and will not terminate sessions.

Before you begin

Configure site-to-site VPN according to the VPN configuration guide.

Procedure

Step 1 On the control node of the cluster, enter cluster configuration mode.

cluster group *name*

Example:

```
ciscoasa(config)# cluster group cluster1
ciscoasa(cfg-cluster)#
```

Step 2 Enable distributed site-to-site VPN.

vpn-mode distributed backup flat

or

vpn-mode distributed backup remote-chassis

In **flat** backup mode, standby sessions are established on any other cluster node. This will protect users from module failures; however, chassis failure protection is not guaranteed.

In **remote-chassis** backup mode, standby sessions are established on a node of another chassis in the cluster. This will protect users from both module failures and chassis failures.

If **remote-chassis** is configured in a single chassis environment (intentionally configured or the result of a failure), no backups will be created until another chassis joins.

Example:

```
ciscoasa(cfg-cluster)# vpn-mode distributed backup remote-chassis
WARNING: Do you want to proceed with changing the vpn-mode, save the device configuration,
and initiate a reboot? [confirm]
```

Redistribute Distributed S2S VPN Sessions

Active Session Redistribution (ASR) redistributes the active VPN session load across the cluster members. Due to the dynamic nature of beginning and ending sessions, ASR is a best effort balancing of the sessions across all cluster members. Repeated redistribution actions will optimize the balance.

Redistribution can be run at any time, should be run after any topology change in the cluster, and is recommended after a new member joins the cluster. The goal of redistribution is to create a stable VPN cluster. A stable VPN cluster has an almost equal number of active and backup sessions across the nodes.

To move a session, the backup session becomes the active one and another node is selected to host a new backup session. Moving sessions is dependent on the location of the active session's backup and the number of active sessions already on that particular backup node. If the backup session node is unable to host the active session for some reason, the original node remains owner of the session.

In multiple-context mode, active session redistribution is done at the system level, not the individual context level. It is not done at the context level because an active session in one context could be moved a member that contains many more active sessions in a different context, creating more load on that cluster member.

Before you begin

- Enable system logs if you would like to monitor redistribution activity.
- This procedure must be carried out on the control unit of the cluster.

Procedure

Step 1 On the control node, view how active and backup sessions are distributed across the cluster.

show cluster vpn-sessiondb distribution**Example:**

Distribution information displays as follows:

```
Member 0 (unit-1-1): active: 209; backups at: 1(111), 2(98)
Member 1 (unit-1-3): active: 204; backups at: 0(108), 2(96)
Member 2 (unit-1-2): active: 0
```

Each row contains the member ID, member name, number of active sessions, and on which members the backup sessions reside. For the example above, one would read the information as:

- Member 0 has 209 active sessions, 111 sessions are backed up on member 1, 98 sessions are backed up on member 2
- Member 1 has 204 active sessions, 108 sessions are backed up on member 0, 96 sessions are backed up on member 2
- Member 2 has NO active sessions; therefore, no cluster members are backing up sessions for this node. This member has recently joined the cluster.

Step 2 Redistribute sessions.

cluster redistribute vpn-sessiondb

This command returns immediately (with no message) while it executes in the background.

Depending on the number of sessions to redistribute and the load on the cluster, this may take some time. Syslogs containing the following phrases (and other system details not shown here) are provided as redistribution activity occurs:

Syslog Phrase	Notes
VPN session redistribution started	Control node only
Sent request to move <i>number</i> sessions from <i>orig-member-name</i> to <i>dest-member-name</i>	Control node only
Failed to send session redistribution message to <i>member-name</i>	Control node only
Received request to move <i>number</i> sessions from <i>orig-member-name</i> to <i>dest-member-name</i>	Data node only
Moved <i>number</i> sessions to <i>member-name</i>	The number of active sessions moved to the named cluster.

Syslog Phrase	Notes
Failed to receive session move response from <i>dest-member-name</i>	Control node only
VPN session completed	Control node only
Cluster topology change detected. VPN session redistribution aborted.	

Step 3 Re-enter the **show cluster vpn-sessiondb distribution** command to view the results.

FXOS: Remove a Cluster Node

The following sections describe how to remove nodes temporarily or permanently from the cluster.

Temporary Removal

A cluster node will be automatically removed from the cluster due to a hardware or network failure, for example. This removal is temporary until the conditions are rectified, and it can rejoin the cluster. You can also manually disable clustering.

To check whether a device is currently in the cluster, check the cluster status within the application using the **show cluster info** command:

```
ciscoasa# show cluster info
Clustering is not enabled
```

- Disable clustering in the application—You can disable clustering using the application CLI. Enter the **cluster remove unit name** command to remove any node other than the one you are logged into. The bootstrap configuration remains intact, as well as the last configuration synced from the control node, so you can later re-add the node without losing your configuration. If you enter this command on a data node to remove the control node, a new control node is elected.

When a device becomes inactive, all data interfaces are shut down; only the Management interface can send and receive traffic. To resume traffic flow, re-enable clustering. The Management interface remains up using the IP address the node received from the bootstrap configuration. However if you reload, and the node is still inactive in the cluster (for example, you saved the configuration with clustering disabled), the Management interface is disabled.

To reenabling clustering, on the ASA enter **cluster group name** and then **enable**.

- Disable the application instance—At the FXOS CLI, see the following example:

```
Firepower-chassis# scope ssa
Firepower-chassis /ssa # scope slot 1
Firepower-chassis /ssa/slot # scope app-instance asa asa1
Firepower-chassis /ssa/slot/app-instance # disable
Firepower-chassis /ssa/slot/app-instance* # commit-buffer
Firepower-chassis /ssa/slot/app-instance #
```

To reenabling:

```
Firepower-chassis /ssa/slot/app-instance # enable
```

```
Firepower-chassis /ssa/slot/app-instance* # commit-buffer
Firepower-chassis /ssa/slot/app-instance #
```

- Shut down the security module/engine—At the FXOS CLI, see the following example:

```
Firepower-chassis# scope service-profile server 1/1
Firepower-chassis /org/service-profile # power down soft-shut-down
Firepower-chassis /org/service-profile* # commit-buffer
Firepower-chassis /org/service-profile #
```

To power up:

```
Firepower-chassis /org/service-profile # power up
Firepower-chassis /org/service-profile* # commit-buffer
Firepower-chassis /org/service-profile #
```

- Shut down the chassis—At the FXOS CLI, see the following example:

```
Firepower-chassis# scope chassis 1
Firepower-chassis /chassis # shutdown no-prompt
```

Permanent Removal

You can permanently remove a cluster node using the following methods.

- Delete the logical device—At the FXOS CLI, see the following example:

```
Firepower-chassis# scope ssa
Firepower-chassis /ssa # delete logical-device cluster1
Firepower-chassis /ssa* # commit-buffer
Firepower-chassis /ssa #
```

- Remove the chassis or security module from service—If you remove a device from service, you can add replacement hardware as a new node of the cluster.

ASA: Manage Cluster Members

After you deploy the cluster, you can change the configuration and manage cluster members.

Become an Inactive Member

To become an inactive member of the cluster, disable clustering on the node while leaving the clustering configuration intact.



Note When an ASA becomes inactive (either manually or through a health check failure), all data interfaces are shut down; only the management-only interface can send and receive traffic. To resume traffic flow, re-enable clustering; or you can remove the node altogether from the cluster. The management interface remains up using the IP address the node received from the cluster IP pool. However if you reload, and the node is still inactive in the cluster (for example, you saved the configuration with clustering disabled), then the management interface is disabled. You must use the console port for any further configuration.

Before you begin

- You must use the console port; you cannot enable or disable clustering from a remote CLI connection.
- For multiple context mode, perform this procedure in the system execution space. If you are not already in the System configuration mode, enter the **changeto system** command.

Procedure

Step 1 Enter cluster configuration mode:

cluster group *name*

Example:

```
ciscoasa(config)# cluster group pod1
```

Step 2 Disable clustering:

no enable

If this node was the control node, a new control election takes place, and a different member becomes the control node.

The cluster configuration is maintained, so that you can enable clustering again later.

Deactivate a Unit

To deactivate a member other than the node you are logged into, perform the following steps.



Note When an ASA becomes inactive, all data interfaces are shut down; only the management-only interface can send and receive traffic. To resume traffic flow, re-enable clustering. The management interface remains up using the IP address the node received from the cluster IP pool. However if you reload, and the node is still inactive in the cluster (for example, if you saved the configuration with clustering disabled), the management interface is disabled. You must use the console port for any further configuration.

Before you begin

For multiple context mode, perform this procedure in the system execution space. If you are not already in the System configuration mode, enter the **changeto system** command.

Procedure

Remove the node from the cluster.

cluster remove unit *node_name*

The bootstrap configuration remains intact, as well as the last configuration synched from the control node, so that you can later re-add the node without losing your configuration. If you enter this command on a data node to remove the control node, a new control node is elected.

To view member names, enter **cluster remove unit ?**, or enter the **show cluster info** command.

Example:

```
ciscoasa(config)# cluster remove unit ?
Current active units in the cluster:
asa2

ciscoasa(config)# cluster remove unit asa2
WARNING: Clustering will be disabled on unit asa2. To bring it back
to the cluster please logon to that unit and re-enable clustering
```

Rejoin the Cluster

If a node was removed from the cluster, for example for a failed interface or if you manually deactivated a member, you must manually rejoin the cluster.

Before you begin

- You must use the console port to reenabling clustering. Other interfaces are shut down.
- For multiple context mode, perform this procedure in the system execution space. If you are not already in the System configuration mode, enter the **changeto system** command.
- Make sure the failure is resolved before you try to rejoin the cluster.

Procedure

Step 1 At the console, enter cluster configuration mode:

cluster group *name*

Example:

```
ciscoasa(config)# cluster group pod1
```

- Step 2** Enable clustering.
enable
-

Change the Control Unit

**Caution**

The best method to change the control node is to disable clustering on the control node, wait for a new control election, and then re-enable clustering. If you must specify the exact node you want to become the control node, use the procedure in this section. Note, however, that for centralized features, if you force a control node change using this procedure, then all connections are dropped, and you have to re-establish the connections on the new control node.

To change the control node, perform the following steps.

Before you begin

For multiple context mode, perform this procedure in the system execution space. If you are not already in the System configuration mode, enter the **changeto system** command.

Procedure

Set a new node as the control node:

cluster control-node unit*node_name*

Example:

```
ciscoasa(config)# cluster control-node unit asa2
```

You will need to reconnect to the Main cluster IP address.

To view member names, enter **cluster control-node unit ?** (to see all names except the current node), or enter the **show cluster info** command.

Execute a Command Cluster-Wide

To send a command to all members in the cluster, or to a specific member, perform the following steps. Sending a **show** command to all members collects all output and displays it on the console of the current unit. (Note that alternatively there are show commands that you can enter on the control unit to view cluster-wide statistics.) Other commands, such as **capture** and **copy**, can also take advantage of cluster-wide execution.

Procedure

Send a command to all members, or if you specify the unit name, a specific member:

cluster exec [*unit unit_name*] *command*

Example:

```
ciscoasa# cluster exec show xlate
```

To view member names, enter **cluster exec unit ?** (to see all names except the current unit), or enter the **show cluster info** command.

Examples

To copy the same capture file from all units in the cluster at the same time to a TFTP server, enter the following command on the control unit:

```
ciscoasa# cluster exec copy /pcap capture: tftp://10.1.1.56/capture1.pcap
```

Multiple PCAP files, one from each unit, are copied to the TFTP server. The destination capture file name is automatically attached with the unit name, such as capture1_asa1.pcap, capture1_asa2.pcap, and so on. In this example, asa1 and asa2 are cluster unit names.

The following sample output for the **cluster exec show memory** command shows memory information for each member in the cluster:

```
ciscoasa# cluster exec show memory
unit-1-1 (LOCAL):*****
Free memory:      108724634538 bytes (92%)
Used memory:      9410087158 bytes ( 8%)
-----
Total memory:     118111600640 bytes (100%)

unit-1-3:*****
Free memory:      108749922170 bytes (92%)
Used memory:      9371097334 bytes ( 8%)
-----
Total memory:     118111600640 bytes (100%)

unit-1-2:*****
Free memory:      108426753537 bytes (92%)
Used memory:      9697869087 bytes ( 8%)
-----
Total memory:     118111600640 bytes (100%)
```

ASA: Monitoring the ASA Cluster on the Firepower 4100/9300 chassis

You can monitor and troubleshoot cluster status and connections.

Monitoring Cluster Status

See the following commands for monitoring cluster status:

- **show cluster info [health], show cluster chassis info**

With no keywords, the **show cluster info** command shows the status of all members of the cluster.

The **show cluster info health** command shows the current health of interfaces, units, and the cluster overall.

See the following output for the **show cluster info** command:

```
asa(config)# show cluster info
Cluster cluster1: On
  Interface mode: spanned
  This is "unit-1-2" in state MASTER
    ID      : 2
    Version  : 9.5(2)
    Serial No.: FCH183770GD
    CCL IP   : 127.2.1.2
    CCL MAC  : 0015.c500.019f
    Last join : 01:18:34 UTC Nov 4 2015
    Last leave: N/A
Other members in the cluster:
  Unit "unit-1-3" in state SLAVE
    ID      : 4
    Version  : 9.5(2)
    Serial No.: FCH19057ML0
    CCL IP   : 127.2.1.3
    CCL MAC  : 0015.c500.018f
    Last join : 20:29:57 UTC Nov 4 2015
    Last leave: 20:24:55 UTC Nov 4 2015
  Unit "unit-1-1" in state SLAVE
    ID      : 1
    Version  : 9.5(2)
    Serial No.: FCH19057ML0
    CCL IP   : 127.2.1.1
    CCL MAC  : 0015.c500.017f
    Last join : 20:20:53 UTC Nov 4 2015
    Last leave: 20:18:15 UTC Nov 4 2015
  Unit "unit-2-1" in state SLAVE
    ID      : 3
    Version  : 9.5(2)
    Serial No.: FCH19057ML0
    CCL IP   : 127.2.2.1
    CCL MAC  : 0015.c500.020f
    Last join : 20:19:57 UTC Nov 4 2015
    Last leave: 20:24:55 UTC Nov 4 2015
```

- **show cluster info auto-join**

Shows whether the cluster unit will automatically rejoin the cluster after a time delay and if the failure conditions (such as waiting for the license, chassis health check failure, and so on) are cleared. If the unit is permanently disabled, or if the unit is already in the cluster, then this command will not show any output.

See the following outputs for the **show cluster info auto-join** command:

```
ciscoasa(cfg-cluster)# show cluster info auto-join
Unit will try to join cluster in 253 seconds.
Quit reason: Received control message DISABLE

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit will try to join cluster when quit reason is cleared.
Quit reason: Master has application down that slave has up.

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit will try to join cluster when quit reason is cleared.
Quit reason: Chassis-blade health check failed.

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit will try to join cluster when quit reason is cleared.
Quit reason: Service chain application became down.

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit will try to join cluster when quit reason is cleared.
Quit reason: Unit is kicked out from cluster because of Application health check failure.

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit join is pending (waiting for the smart license entitlement: ent1)

ciscoasa(cfg-cluster)# show cluster info auto-join
Unit join is pending (waiting for the smart license export control flag)
```

- **show cluster info transport {asp | cp [detail]}**

Shows transport related statistics for the following:

- **asp** —Data plane transport statistics.
- **cp** —Control plane transport statistics.

If you enter the **detail** keyword, you can view cluster reliable transport protocol usage so you can identify packet drop issues when the buffer is full in the control plane. See the following output for the **show cluster info transport cp detail** command:

```
ciscoasa# show cluster info transport cp detail
Member ID to name mapping:
  0 - unit-1-1    2 - unit-4-1    3 - unit-2-1

Legend:
  U    - unreliable messages
  UE   - unreliable messages error
  SN   - sequence number
  ESN  - expecting sequence number
  R    - reliable messages
  RE   - reliable messages error
  RDC  - reliable message deliveries confirmed
  RA   - reliable ack packets received
  RFR  - reliable fast retransmits
  RTR  - reliable timer-based retransmits
```

RDP - reliable message dropped
 RDPR - reliable message drops reported
 RI - reliable message with old sequence number
 RO - reliable message with out of order sequence number
 ROW - reliable message with out of window sequence number
 ROB - out of order reliable messages buffered
 RAS - reliable ack packets sent

This unit as a sender

```

-----
      all      0      2      3
U    123301    3867966  3230662  3850381
UE    0        0        0        0
SN    1656a4ce acb26fe  5f839f76  7b680831
R     733840    1042168  852285   867311
RE    0        0        0        0
RDC   699789    934969   740874   756490
RA    385525    281198   204021   205384
RFR   27626     56397    0        0
RTR   34051     107199   111411   110821
RDP    0        0        0        0
RDPR   0        0        0        0
  
```

This unit as a receiver of broadcast messages

```

-----
      0      2      3
U    111847    121862  120029
R     7503     665700  749288
ESN   5d75b4b3 6d81d23  365ddd50
RI     630     34278   40291
RO     0       582     850
ROW    0       566     850
ROB    0       16      0
RAS   1571     123289  142256
  
```

This unit as a receiver of unicast messages

```

-----
      0      2      3
U     1     3308122  4370233
R    513846   879979   1009492
ESN   4458903a 6d841a84  7b4e7fa7
RI    66024    108924   102114
RO     0       0       0
ROW    0       0       0
ROB    0       0       0
RAS   130258   218924   228303
  
```

Gated Tx Buffered Message Statistics

```

-----
current sequence number: 0

total:                    0
current:                  0
high watermark:          0

delivered:                0
deliver failures:         0

buffer full drops:        0
message truncate drops:   0

gate close ref count:     0

num of supported clients:45
  
```

MRT Tx of broadcast messages

=====

Message high watermark: 3%

Total messages buffered at high watermark: 5677

[Per-client message usage at high watermark]

Client name	Total messages	Percentage
Cluster Redirect Client	4153	73%
Route Cluster Client	419	7%
RRI Cluster Client	1105	19%

Current MRT buffer usage: 0%

Total messages buffered in real-time: 1

[Per-client message usage in real-time]

Legend:

F - MRT messages sending when buffer is full

L - MRT messages sending when cluster node leave

R - MRT messages sending in Rx thread

Client name	Total messages	Percentage	F	L	R
VPN Clustering HA Client	1	100%	0	0	0

MRT Tx of unitcast messages(to member_id:0)

=====

Message high watermark: 31%

Total messages buffered at high watermark: 4059

[Per-client message usage at high watermark]

Client name	Total messages	Percentage
Cluster Redirect Client	3731	91%
RRI Cluster Client	328	8%

Current MRT buffer usage: 29%

Total messages buffered in real-time: 3924

[Per-client message usage in real-time]

Legend:

F - MRT messages sending when buffer is full

L - MRT messages sending when cluster node leave

R - MRT messages sending in Rx thread

Client name	Total messages	Percentage	F	L	R
Cluster Redirect Client	3607	91%	0	0	0
RRI Cluster Client	317	8%	0	0	0

MRT Tx of unitcast messages(to member_id:2)

=====

Message high watermark: 14%

Total messages buffered at high watermark: 578

[Per-client message usage at high watermark]

Client name	Total messages	Percentage
VPN Clustering HA Client	578	100%

Current MRT buffer usage: 0%

Total messages buffered in real-time: 0

MRT Tx of unitcast messages(to member_id:3)

=====

Message high watermark: 12%

Total messages buffered at high watermark: 573

[Per-client message usage at high watermark]

Client name	Total messages	Percentage
-------------	----------------	------------

```

VPN Clustering HA Client          572      99%
Cluster VPN Unique ID Client      1        0%

Current MRT buffer usage: 0%
Total messages buffered in real-time: 0

```

- **show cluster history**

Shows the cluster history, as well as error messages about why a cluster unit failed to join or why a unit left the cluster.

Capturing Packets Cluster-Wide

See the following command for capturing packets in a cluster:

cluster exec capture

To support cluster-wide troubleshooting, you can enable capture of cluster-specific traffic on the control node using the **cluster exec capture** command, which is then automatically enabled on all of the data nodes in the cluster.

Monitoring Cluster Resources

See the following command for monitoring cluster resources:

show cluster {cpu | memory | resource} [options], show cluster chassis [cpu | memory | resource usage]

Displays aggregated data for the entire cluster. The options available depends on the data type.

Monitoring Cluster Traffic

See the following command for monitoring cluster traffic:

- **show conn [detail | count], cluster exec show conn**

The **show conn** command shows whether a flow is a director, backup, or forwarder flow. Use the **cluster exec show conn** command on any unit to view all connections. This command can show how traffic for a single flow arrives at different ASAs in the cluster. The throughput of the cluster is dependent on the efficiency and configuration of load balancing. This command provides an easy way to view how traffic for a connection is flowing through the cluster, and can help you understand how a load balancer might affect the performance of a flow.

The following is sample output for the **show conn detail** command:

```

ciscoasa/ASA2/slave# show conn detail
15 in use, 21 most used
Cluster:
    fwd connections: 0 in use, 0 most used
    dir connections: 0 in use, 0 most used
    centralized connections: 0 in use, 44 most used
Flags: A - awaiting inside ACK to SYN, a - awaiting outside ACK to SYN,
      B - initial SYN from outside, b - TCP state-bypass or nailed,
      C - CTIQBE media, c - cluster centralized,
      D - DNS, d - dump, E - outside back connection, e - semi-distributed,

```

F - outside FIN, f - inside FIN,
 G - group, g - MGCP, H - H.323, h - H.225.0, I - inbound data,
 i - incomplete, J - GTP, j - GTP data, K - GTP t3-response
 k - Skinny media, L - LISP triggered flow owner mobility
 M - SMTP data, m - SIP media, n - GUP
 N - inspected by Snort
 O - outbound data, o - offloaded,
 P - inside back connection,
 Q - Diameter, q - SQL*Net data,
 R - outside acknowledged FIN,
 R - UDP SUNRPC, r - inside acknowledged FIN, S - awaiting inside SYN,
 s - awaiting outside SYN, T - SIP, t - SIP transient, U - up,
 V - VPN orphan, W - WAAS,
 w - secondary domain backup,
 X - inspected by service module,
 x - per session, Y - director stub flow, y - backup stub flow,
 Z - Scansafe redirection, z - forwarding stub flow

Cluster units to ID mappings:

ID 0: unit-2-1
 ID 1: unit-1-1
 ID 2: unit-1-2
 ID 3: unit-2-2
 ID 4: unit-2-3
 ID 255: The default cluster member ID which indicates no ownership or affiliation
 with an existing cluster member

• **show cluster info [conn-distribution | packet-distribution | loadbalance]**

The **show cluster info conn-distribution** and **show cluster info packet-distribution** commands show traffic distribution across all cluster units. These commands can help you to evaluate and adjust the external load balancer.

The **show cluster info loadbalance** command shows connection rebalance statistics.

• **show cluster info load-monitor [details]**

The **show cluster info load-monitor** command shows the traffic load for cluster members for the last interval and also the average over total number of intervals configured (30 by default). Use the **details** keyword to view the value for each measure at each interval.

```

ciscoasa(cfg-cluster)# show cluster info load-monitor
ID  Unit Name
0   B
1   A_1
Information from all units with 20 second interval:
Unit      Connections    Buffer Drops    Memory Used    CPU Used
Average from last 1 interval:
0          0              0              14             25
1          0              0              16             20
Average from last 30 interval:
0          0              0              12             28
1          0              0              13             27
  
```

```

ciscoasa(cfg-cluster)# show cluster info load-monitor details
  
```

```

ID  Unit Name
  
```

```

0   B
  
```

1 A_1

Information from all units with 20 second interval

Connection count captured over 30 intervals:

Unit ID 0

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

Unit ID 1

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

Buffer drops captured over 30 intervals:

Unit ID 0

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

Unit ID 1

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

Memory usage(%) captured over 30 intervals:

```

Unit ID  0
          25      25      30      30      30      35
          25      25      35      30      30      30
          25      25      30      25      25      35
          30      30      30      25      25      25
          25      20      30      30      30      30

Unit ID  1
          30      25      35      25      30      30
          25      25      35      25      30      35
          30      30      35      30      30      30
          25      20      30      25      25      30
          20      30      35      30      30      35

```

CPU usage(%) captured over 30 intervals:

```

Unit ID  0
          25      25      30      30      30      35
          25      25      35      30      30      30
          25      25      30      25      25      35
          30      30      30      25      25      25
          25      20      30      30      30      30

Unit ID  1
          30      25      35      25      30      30
          25      25      35      25      30      35
          30      30      35      30      30      30
          25      20      30      25      25      30
          20      30      35      30      30      35

```

- **show cluster {access-list | conn [count] | traffic | user-identity | xlate} [options], show cluster chassis {access-list | conn | traffic | user-identity | xlate count}**

Displays aggregated data for the entire cluster. The options available depends on the data type.

See the following output for the **show cluster access-list** command:

```

ciscoasa# show cluster access-list
hitcnt display order: cluster-wide aggregated result, unit-A, unit-B, unit-C, unit-D
access-list cached ACL log flows: total 0, denied 0 (deny-flow-max 4096) alert-interval
300

```

```

access-list 101; 122 elements; name hash: 0xe7d586b5
access-list 101 line 1 extended permit tcp 192.168.143.0 255.255.255.0 any eq www
(hitcnt=0, 0, 0, 0, 0) 0x207a2b7d
access-list 101 line 2 extended permit tcp any 192.168.143.0 255.255.255.0 (hitcnt=0,
0, 0, 0, 0) 0xfe4f4947
access-list 101 line 3 extended permit tcp host 192.168.1.183 host 192.168.43.238
(hitcnt=1, 0, 0, 0, 1) 0x7b521307
access-list 101 line 4 extended permit tcp host 192.168.1.116 host 192.168.43.238
(hitcnt=0, 0, 0, 0, 0) 0x5795c069
access-list 101 line 5 extended permit tcp host 192.168.1.177 host 192.168.43.238
(hitcnt=1, 0, 0, 1, 0) 0x51bde7ee
access-list 101 line 6 extended permit tcp host 192.168.1.177 host 192.168.43.13
(hitcnt=0, 0, 0, 0, 0) 0x1e68697c
access-list 101 line 7 extended permit tcp host 192.168.1.177 host 192.168.43.132
(hitcnt=2, 0, 0, 1, 1) 0xc1ce5c49
access-list 101 line 8 extended permit tcp host 192.168.1.177 host 192.168.43.192
(hitcnt=3, 0, 1, 1, 1) 0xb6f59512
access-list 101 line 9 extended permit tcp host 192.168.1.177 host 192.168.43.44
(hitcnt=0, 0, 0, 0, 0) 0xdc104200
access-list 101 line 10 extended permit tcp host 192.168.1.112 host 192.168.43.44
(hitcnt=429, 109, 107, 109, 104)
0xce4f281d
access-list 101 line 11 extended permit tcp host 192.168.1.170 host 192.168.43.238
(hitcnt=3, 1, 0, 0, 2) 0x4143a818
access-list 101 line 12 extended permit tcp host 192.168.1.170 host 192.168.43.169
(hitcnt=2, 0, 1, 0, 1) 0xb18dfea4
access-list 101 line 13 extended permit tcp host 192.168.1.170 host 192.168.43.229
(hitcnt=1, 1, 0, 0, 0) 0x21557d71
access-list 101 line 14 extended permit tcp host 192.168.1.170 host 192.168.43.106
(hitcnt=0, 0, 0, 0, 0) 0x7316e016
access-list 101 line 15 extended permit tcp host 192.168.1.170 host 192.168.43.196
(hitcnt=0, 0, 0, 0, 0) 0x013fd5b8
access-list 101 line 16 extended permit tcp host 192.168.1.170 host 192.168.43.75
(hitcnt=0, 0, 0, 0, 0) 0x2c7dba0d

```

To display the aggregated count of in-use connections for all units, enter:

```

ciscoasa# show cluster conn count
Usage Summary In Cluster:*****
124 in use, fwd connection 0 in use, dir connection 0 in use, centralized connection
0 in use (Cluster-wide aggregated)

unit-1-1(LOCAL):*****
40 in use, 48 most used, fwd connection 0 in use, 0 most used, dir connection 0 in use,
0 most used, centralized connection 0 in use, 46 most used

unit-2-2:*****
18 in use, 40 most used, fwd connection 0 in use, 0 most used, dir connection 0 in use,
0 most used, centralized connection 0 in use, 45 most used

```

- **show asp cluster counter**

This command is useful for datapath troubleshooting.

Monitoring Cluster Routing

See the following commands for cluster routing:

- **show route cluster**

- **debug route cluster**

Shows cluster information for routing.

- **show lisp eid**

Shows the ASA EID table showing EIDs and site IDs.

See the following output from the **cluster exec show lisp eid** command.

```
ciscoasa# cluster exec show lisp eid
L1 (LOCAL) :*****
      LISP EID      Site ID
      33.44.33.105      2
      33.44.33.201      2
      11.22.11.1         4
      11.22.11.2         4
L2:*****
      LISP EID      Site ID
      33.44.33.105      2
      33.44.33.201      2
      11.22.11.1      4
      11.22.11.2      4
```

- **show asp table classify domain inspect-lisp**

This command is useful for troubleshooting.

Monitoring Distributed S2S VPN

Use the following commands to monitor status and distribution of the VPN sessions:

- The overall distribution of sessions is provided using **show cluster vpn-sessiondb distribution**. If running in a multiple context environment, this command must be run in the system execution space. This **show** command provides a quick view of the sessions, rather than having to execute **show vpn-sessiondb summary** on each node.
- A unified view of the VPN connections on the cluster using the **show cluster vpn-sessiondb summary** command is also available.
- Individual device monitoring using the **show vpn-sessiondb** command shows the number of active and backup sessions on a device in addition to the usual VPN information.

Configuring Logging for Clustering

See the following command for configuring logging for clustering:

logging device-id

Each node in the cluster generates syslog messages independently. You can use the **logging device-id** command to generate syslog messages with identical or different device IDs to make messages appear to come from the same or different nodes in the cluster.

Debugging Clustering

See the following commands for debugging clustering:

- **debug cluster [ccp | datapath | fsm | general | hc | license | rpc | service-module | transport]**

Shows debug messages for clustering.

- **debug service-module**

Shows debug messages for blade level issues including health check issues between the supervisor and the application.

- **show cluster info trace**

The **show cluster info trace** command shows the debug information for further troubleshooting.

See the following output for the **show cluster info trace** command:

```
ciscoasa# show cluster info trace
Feb 02 14:19:47.456 [DEBUG]Receive CCP message: CCP_MSG_LOAD_BALANCE
Feb 02 14:19:47.456 [DEBUG]Receive CCP message: CCP_MSG_LOAD_BALANCE
Feb 02 14:19:47.456 [DEBUG]Send CCP message to all: CCP_MSG_KEEPALIVE from 80-1 at
MASTER
```

For example, if you see the following messages showing that two nodes with the same **local-unit** name are acting as the control node, it could mean that either two nodes have the same **local-unit** name (check your configuration), or a node is receiving its own broadcast messages (check your network).

```
ciscoasa# show cluster info trace
May 23 07:27:23.113 [CRIT]Received datapath event 'multi control_nodes' with parameter
1.
May 23 07:27:23.113 [CRIT]Found both unit-9-1 and unit-9-1 as control_node units.
Control_node role retained by unit-9-1, unit-9-1 will leave then join as a Data_node
May 23 07:27:23.113 [DEBUG]Send event (DISABLE, RESTART | INTERNAL-EVENT, 5000 msecs,
Detected another Control_node, leave and re-join as Data_node) to FSM. Current state
CONTROL_NODE
May 23 07:27:23.113 [INFO]State machine changed from state CONTROL_NODE to DISABLED
```

Troubleshooting Distributed S2S VPN

Distributed VPN Notifications

You will be notified with messages containing the identified phrases when the following error situations occur on a cluster running distributed VPN:

Situation	Notification
If an existing or joining cluster data node is not in distributed VPN mode when attempting to join the cluster:	New cluster member (<i>member-name</i>) rejected due to vpn mode mismatch. and Control node (<i>control-name</i>) rejects enrollment request from unit (<i>unit-name</i>) for the reason: the vpn mode capabilities are not compatible with the control node configuration
If licensing is not properly configured on a cluster member for Distributed VPN:	ERROR: Control node requested cluster vpn-mode change to distributed. Unable to change mode due to missing Carrier License.
If the time stamp or member ID is invalid in the SPI of a received IKEv2 packet:	Expired SPI received or Corrupted SPI detected
If the cluster is unable to create a backup session:	Failed to create the backup for an IKEv2 session.
IKEv2 Initial Contact (IC) processing error:	IKEv2 Negotiation aborted due to ERROR: Stale backup session found on backup
Redistribution problems:	Failed to send session redistribution message to <i>member-name</i> Failed to receive session move response from <i>member-name</i> (control node only)
If the topology changes during redistribution of the sessions:	Cluster topology change detected. VPN session redistribution aborted.

You may be encountering one of the following situations:

- Site-to-site VPN sessions are being distributed to only one of the chassis in a cluster when the Nexus 7K switch is configured with a layer 4 port as a load-balancing algorithm using the **port-channel load-balance src-dst l4port** command. An example of the cluster session allocation looks like below:

```

SSP-Cluster/data node(cfg-cluster)# show cluster vpn-sessiondb distribution
Member 0 (unit-1-3): active: 0
Member 1 (unit-2-2): active: 13295; backups at: 0(2536), 2(2769), 3(2495), 4(2835),
5(2660)
Member 2 (unit-2-3): active: 12174; backups at: 0(2074), 1(2687), 3(2207), 4(3084),
5(2122)
Member 3 (unit-2-1): active: 13416; backups at: 0(2419), 1(3013), 2(2712), 4(2771),
5(2501)
Member 4 (unit-1-1): active: 0
Member 5 (unit-1-2): active: 0

```

Since site-to-site IKEv2 VPN uses port 500 for both source and destination ports, IKE packets are only sent to one of the links in the port channel connected between the Nexus 7K and the chassis.

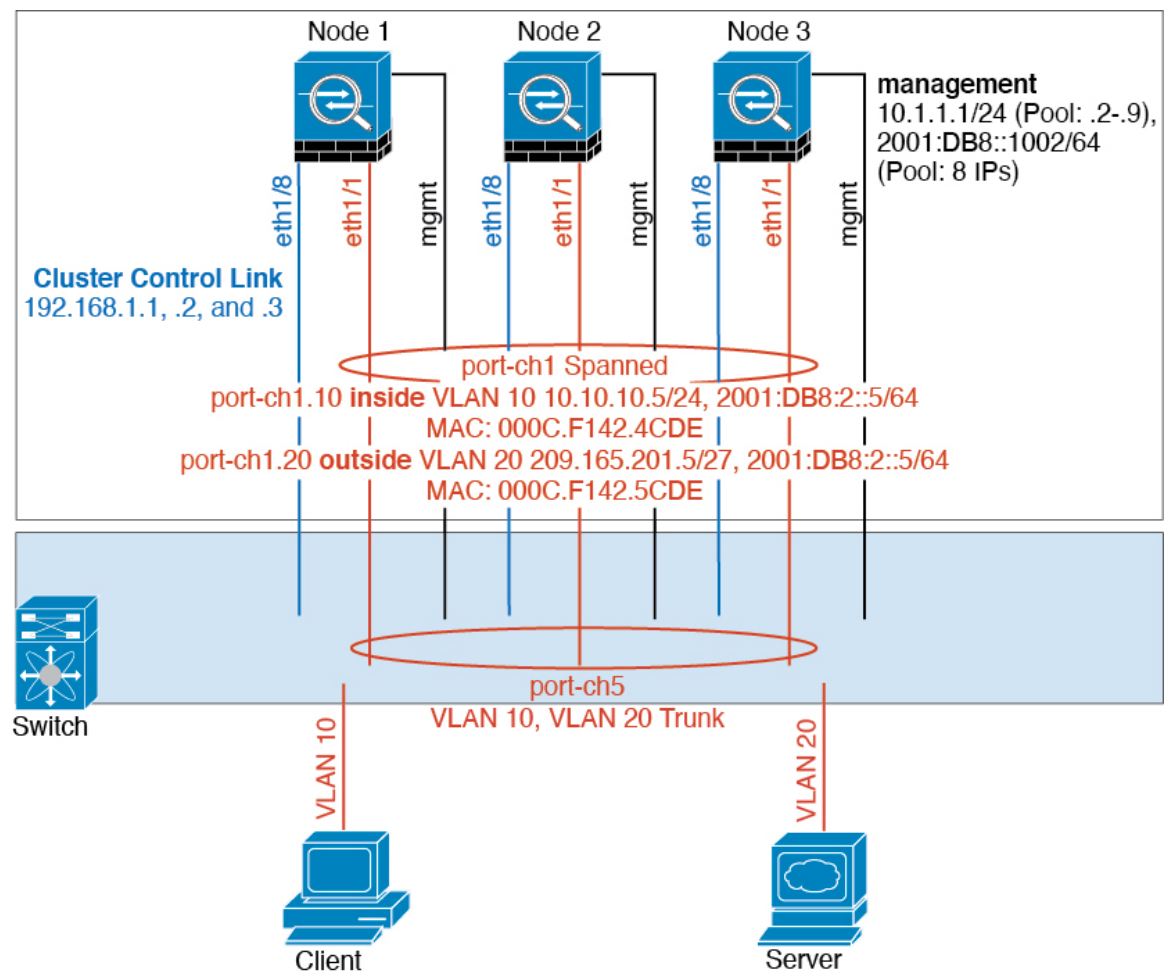
Change the Nexus 7K load balancing algorithm to IP and Layer 4 port using the **port-channel load-balance src-dst ip-l4port**. Then the IKE packets are sent to all the links and thus all nodes.

For a more immediate adjustment, on the control node of the cluster, execute: **cluster redistribute vpn-sessiondb** to redistribute active VPN sessions to the cluster nodes of the other chassis.

Examples for ASA Clustering

These examples include typical deployments.

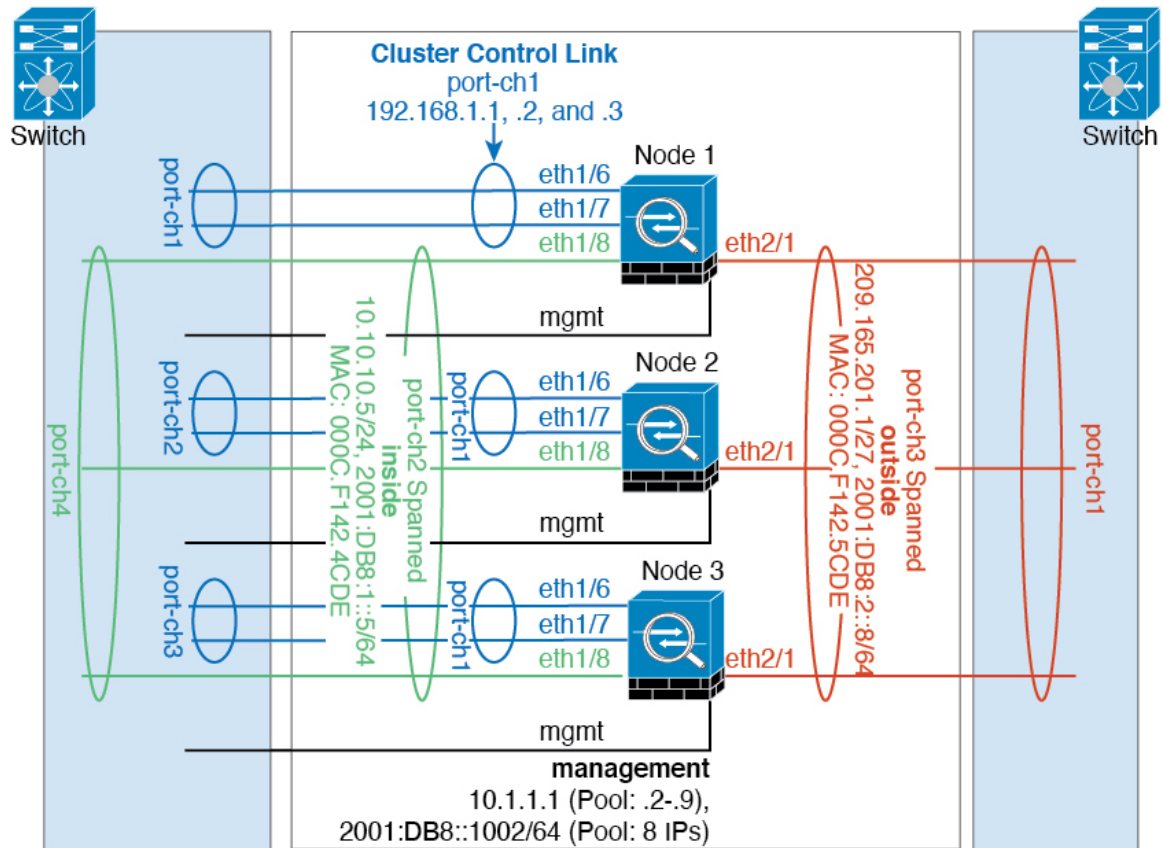
Firewall on a Stick



Data traffic from different security domains are associated with different VLANs, for example, VLAN 10 for the inside network and VLAN 20 for the outside network. Each ASA has a single physical port connected to the external switch or router. Trunking is enabled so that all packets on the physical link are 802.1q encapsulated. The ASA is the firewall between VLAN 10 and VLAN 20.

When using Spanned EtherChannels, all data links are grouped into one EtherChannel on the switch side. If the ASA becomes unavailable, the switch will rebalance traffic between the remaining units.

Traffic Segregation



You may prefer physical separation of traffic between the inside and outside network.

As shown in the diagram above, there is one Spanned EtherChannel on the left side that connects to the inside switch, and the other on the right side to outside switch. You can also create VLAN subinterfaces on each EtherChannel if desired.

OTV Configuration for Routed Mode Inter-Site Clustering

The success of inter-site clustering for routed mode with Spanned EtherChannels depends on the proper configuration and monitoring of OTV. OTV plays a major role by forwarding the packets across the DCI. OTV forwards unicast packets across the DCI only when it learns the MAC address in its forwarding table. If the MAC address is not learned in the OTV forwarding table, it will drop the unicast packets.

Sample OTV Configuration

```
//Sample OTV config:
//3151 - Inside VLAN, 3152 - Outside VLAN, 202 - CCL VLAN
//aaaa.1111.1234 - ASA inside interface global vMAC
```

```
//0050.56A8.3D22 - Server MAC

feature ospf
feature otv

mac access-list ALL_MACs
  10 permit any any
mac access-list HSRP_VMAC
  10 permit aaaa.1111.1234 0000.0000.0000 any
  20 permit aaaa.2222.1234 0000.0000.0000 any
  30 permit any aaaa.1111.1234 0000.0000.0000
  40 permit any aaaa.2222.1234 0000.0000.0000
vlan access-map Local 10
  match mac address HSRP_VMAC
  action drop
vlan access-map Local 20
  match mac address ALL_MACs
  action forward
vlan filter Local vlan-list 3151-3152

//To block global MAC with ARP inspection:
arp access-list HSRP_VMAC_ARP
  10 deny aaaa.1111.1234 0000.0000.0000 any
  20 deny aaaa.2222.1234 0000.0000.0000 any
  30 deny any aaaa.1111.1234 0000.0000.0000
  40 deny any aaaa.2222.1234 0000.0000.0000
  50 permit ip any mac
ip arp inspection filter HSRP_VMAC_ARP 3151-3152

no ip igmp snooping optimise-multicast-flood
vlan 1,202,1111,2222,3151-3152

otv site-vlan 2222
mac-list GMAC_DENY seq 10 deny aaaa.aaaa.aaaa ffff.ffff.ffff
mac-list GMAC_DENY seq 20 deny aaaa.bbbb.bbbb ffff.ffff.ffff
mac-list GMAC_DENY seq 30 permit 0000.0000.0000 0000.0000.0000
route-map stop-GMAC permit 10
  match mac-list GMAC_DENY

interface Overlay1
  otv join-interface Ethernet8/1
  otv control-group 239.1.1.1
  otv data-group 232.1.1.0/28
  otv extend-vlan 202, 3151
  otv arp-nd timeout 60
  no shutdown

interface Ethernet8/1
  description uplink_to_OTV_cloud
  mtu 9198
  ip address 10.4.0.18/24
  ip igmp version 3
  no shutdown

interface Ethernet8/2

interface Ethernet8/3
  description back_to_default_vdc_e6/39
  switchport
  switchport mode trunk
  switchport trunk allowed vlan 202,2222,3151-3152
  mac packet-classify
  no shutdown
```

```

otv-isis default
  vpn Overlay1
    redistribute filter route-map stop-GMAC
otv site-identifier 0x2
//OTV flood not required for ARP inspection:
otv flood mac 0050.56A8.3D22 vlan 3151

```

OTV Filter Modifications Required Because of Site Failure

If a site goes down, the filters need to be removed from OTV because you do not want to block the global MAC address anymore. There are some additional configurations required.

You need to add a static entry for the ASA global MAC address on the OTV switch in the site that is functional. This entry will let the OTV at the other end add these entries on the overlay interface. This step is required because if the server and client already have the ARP entry for the ASA, which is the case for existing connections, then they will not send the ARP again. Therefore, OTV will not get a chance to learn the ASA global MAC address in its forwarding table. Because OTV does not have the global MAC address in its forwarding table, and per OTV design it will not flood unicast packets over the overlay interface, then it will drop the unicast packets to the global MAC address from the server, and the existing connections will break.

```

//OTV filter configs when one of the sites is down

mac-list GMAC_A seq 10 permit 0000.0000.0000 0000.0000.0000
route-map a-GMAC permit 10
  match mac-list GMAC_A

otv-isis default
  vpn Overlay1
    redistribute filter route-map a-GMAC

no vlan filter Local vlan-list 3151

//For ARP inspection, allow global MAC:
arp access-list HSRP_VMAC_ARP_Allow
  50 permit ip any mac
ip arp inspection filter HSRP_VMAC_ARP_Allow 3151-3152

mac address-table static aaaa.1111.1234 vlan 3151 interface Ethernet8/3
//Static entry required only in the OTV in the functioning Site

```

When the other site is restored, you need to add the filters back again and remove this static entry on the OTV. It is very important to clear the dynamic MAC address table on both the OTVs to clear the overlay entry for the global MAC address.

MAC Address Table Clearing

When a site goes down, and a static entry for the global MAC address is added to OTV, you need to let the other OTV learn the global MAC address on the overlay interface. After the other site comes up, these entries should be cleared. Make sure to clear the mac address table to make sure OTV does not have these entries in its forwarding table.

```

cluster-N7k6-OTV# show mac address-table
Legend:
* - primary entry, G - Gateway MAC, (R) - Routed MAC, O - Overlay MAC
age - seconds since last seen,+ - primary entry using vPC Peer-Link,

```

```
(T) - True, (F) - False
VLAN MAC Address Type age Secure NTFY Ports/SWID.SSID.LID
-----+-----+-----+-----+-----+-----+-----+-----
G -      d867.d900.2e42 static - F F sup-eth1 (R)
O 202    885a.92f6.44a5 dynamic - F F Overlay1
* 202    885a.92f6.4b8f dynamic 5  F F Eth8/3
O 3151   0050.5660.9412 dynamic - F F Overlay1
* 3151   aaaa.1111.1234 dynamic 50 F F Eth8/3
```

OTV ARP Cache Monitoring

OTV maintains an ARP cache to proxy ARP for IP addresses that it learned across the OTV interface.

```
cluster-N7k6-OTV# show otv arp-nd-cache
OTV ARP/ND L3->L2 Address Mapping Cache

Overlay Interface Overlay1
VLAN MAC Address Layer-3 Address Age Expires In
3151 0050.5660.9412 10.0.0.2 1w0d 00:00:31
cluster-N7k6-OTV#
```

Examples for Inter-Site Clustering

The following examples show supported cluster deployments.

Spanned EtherChannel Routed Mode Example with Site-Specific MAC and IP Addresses

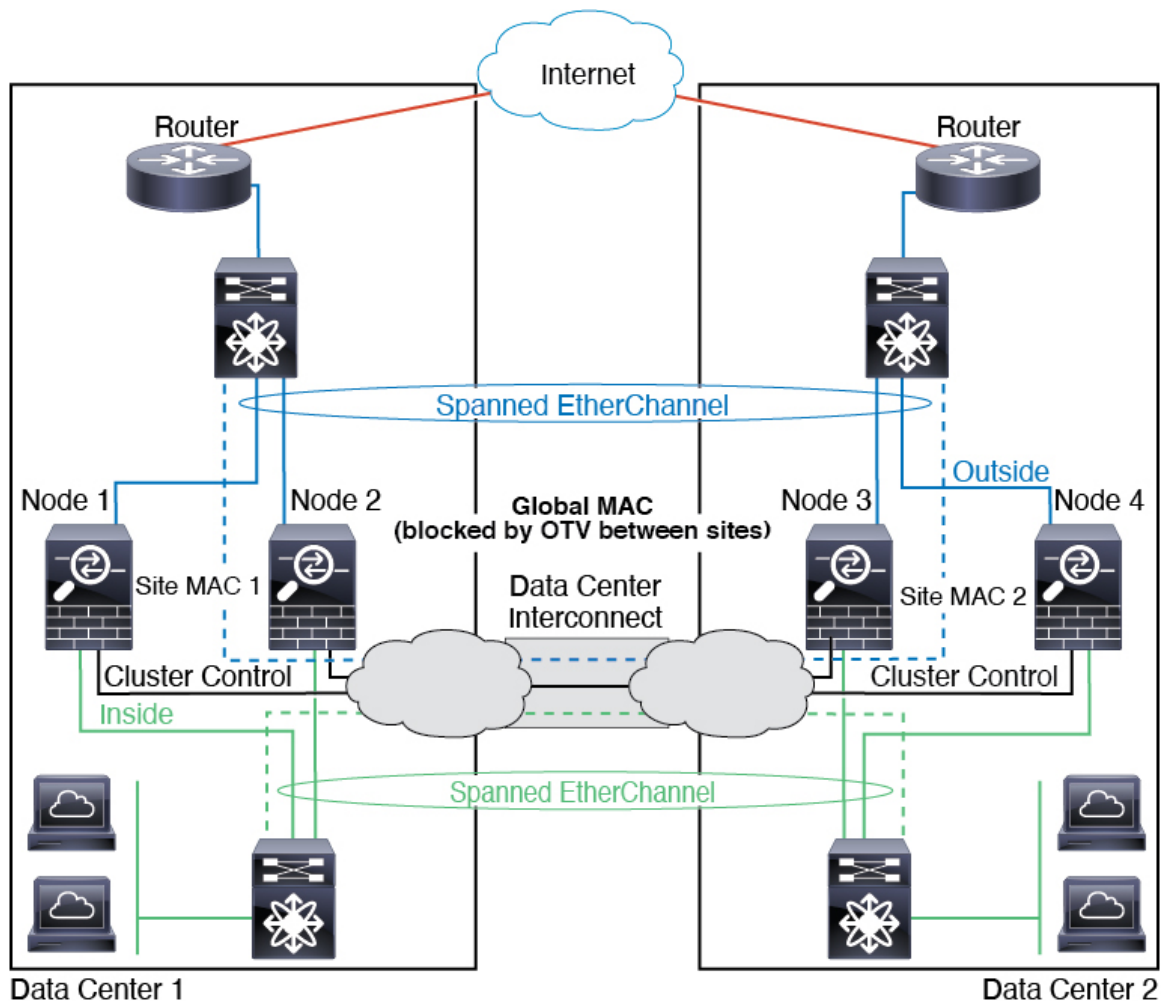
The following example shows 2 cluster members at each of 2 data centers placed between the gateway router and an inside network at each site (East-West insertion). The cluster members are connected by the cluster control link over the DCI. The cluster members at each site connect to the local switches using spanned EtherChannels for both the inside and outside networks. Each EtherChannel is spanned across all chassis in the cluster.

The data VLANs are extended between the sites using Overlay Transport Virtualization (OTV) (or something similar). You must add filters blocking the global MAC address to prevent traffic from traversing the DCI to the other site when the traffic is destined for the cluster. If the cluster nodes at one site become unreachable, you must remove the filters so traffic can be sent to the other site's cluster nodes. You should use VACLs to filter the global MAC address. For some switches, such as Nexus with the F3-series line card, you must also use ARP inspection to block ARP packets from the global MAC address. ARP inspection requires you to set both the site MAC address and the site IP address on the ASA. If you only configure the site MAC address be sure to disable ARP inspection.

The cluster acts as the gateway for the inside networks. The global virtual MAC, which is shared across all cluster nodes, is used only to receive packets. Outgoing packets use a site-specific MAC address from each DC cluster. This feature prevents the switches from learning the same global MAC address from both sites on two different ports, which causes MAC flapping; instead, they only learn the site MAC address.

In this scenario:

- All egress packets sent from the cluster use the site MAC address and are localized at the data center.
- All ingress packets to the cluster are sent using the global MAC address, so they can be received by any of the nodes at both sites; filters at the OTV localize the traffic within the data center.



Spanned EtherChannel Transparent Mode North-South Inter-Site Example

The following example shows 2 cluster members at each of 2 data centers placed between inside and outside routers (North-South insertion). The cluster members are connected by the cluster control link over the DCI. The cluster members at each site connect to the local switches using spanned EtherChannels for the inside and outside. Each EtherChannel is spanned across all chassis in the cluster.

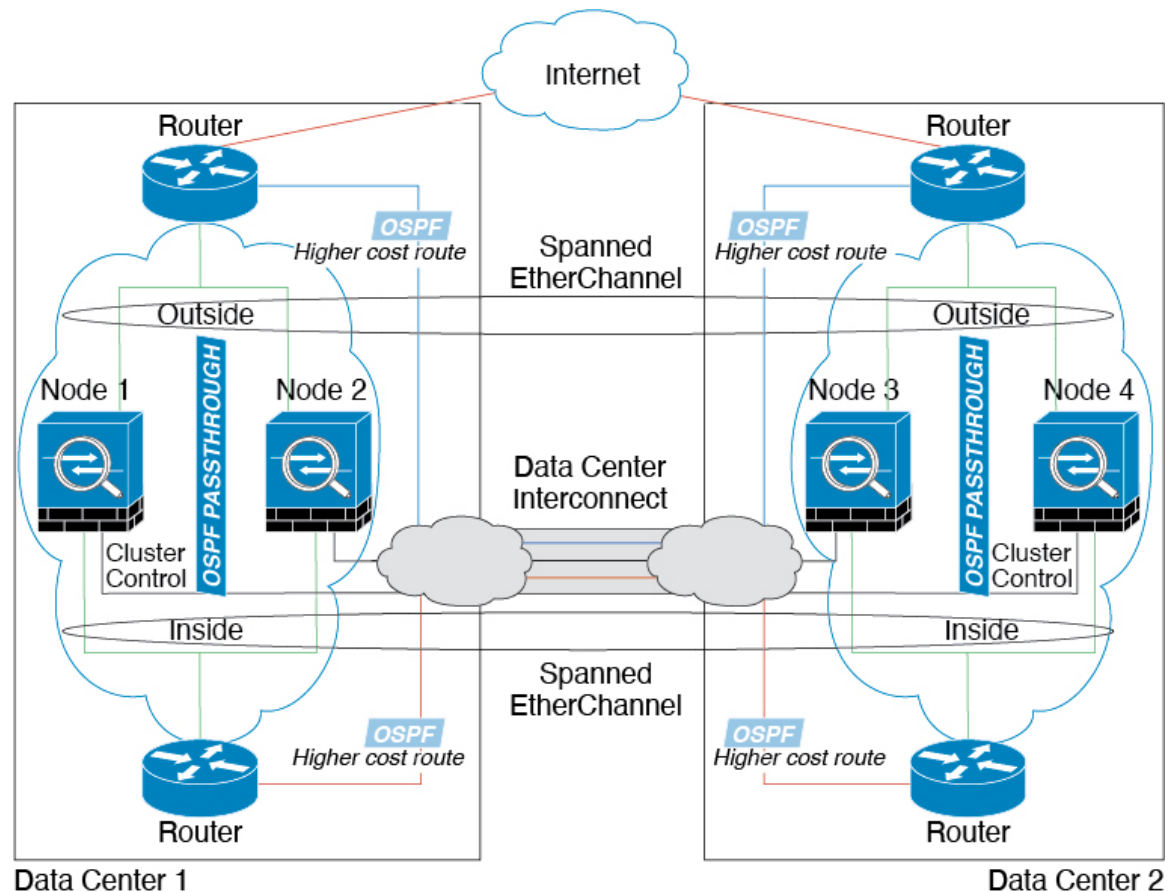
The inside and outside routers at each data center use OSPF, which is passed through the transparent ASAs. Unlike MACs, router IPs are unique on all routers. By assigning a higher cost route across the DCI, traffic stays within each data center unless all cluster members at a given site go down. The lower cost route through the ASAs must traverse the same bridge group at each site for the cluster to maintain asymmetric connections. In the event of a failure of all cluster members at one site, traffic goes from each router over the DCI to the cluster members at the other site.

The implementation of the switches at each site can include:

- Inter-site VSS, vPC, StackWise, or StackWise Virtual—In this scenario, you install one switch at Data Center 1, and the other at Data Center 2. One option is for the cluster nodes at each Data Center to only connect to the local switch, while the redundant switch traffic goes across the DCI. In this case, connections are for the most part kept local to each datacenter. You can optionally connect each node to both switches

across the DCI if the DCI can handle the extra traffic. In this case, traffic is distributed across the data centers, so it is essential for the DCI to be very robust.

- Local VSS, vPC, StackWise, or StackWise Virtual at each site—For better switch redundancy, you can install 2 separate redundant switch pairs at each site. In this case, although the cluster nodes still have a spanned EtherChannel with Data Center 1 chassis connected only to both local switches, and Data Center 2 chassis connected to those local switches, the spanned EtherChannel is essentially “split.” Each local redundant switch system sees the spanned EtherChannel as a site-local EtherChannel.

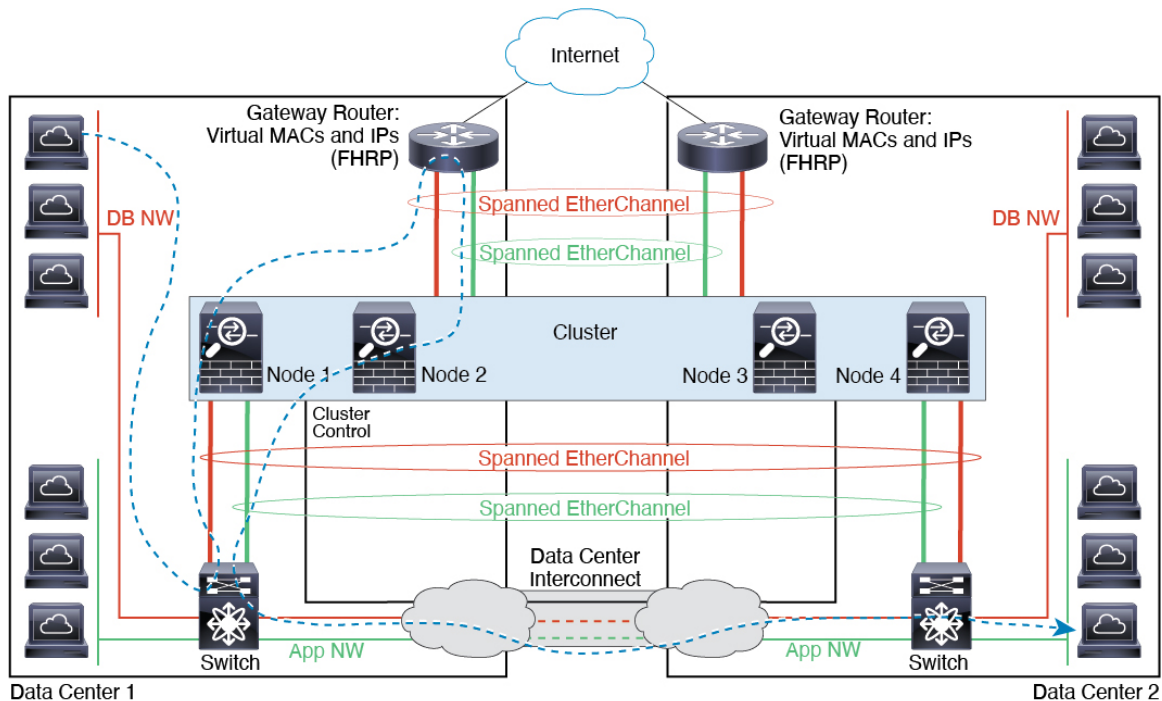


Spanned EtherChannel Transparent Mode East-West Inter-Site Example

The following example shows 2 cluster members at each of 2 data centers placed between the gateway router and two inside networks at each site, the App network and the DB network (East-West insertion). The cluster members are connected by the cluster control link over the DCI. The cluster members at each site connect to the local switches using spanned EtherChannels for both the App and DB networks on the inside and outside. Each EtherChannel is spanned across all chassis in the cluster.

The gateway router at each site uses an FHRP such as HSRP to provide the same destination virtual MAC and IP addresses at each site. A good practice to avoid unintended MAC address flapping is to statically add the gateway routers real MAC addresses to the ASA MAC address table using the **mac-address-table static outside_interface mac_address** command. Without these entries, if the gateway at site 1 communicates with the gateway at site 2, that traffic might pass through the ASA and attempt to reach site 2 from the inside interface and cause problems. The data VLANs are extended between the sites using Overlay Transport

Virtualization (OTV) (or something similar). You must add filters to prevent traffic from traversing the DCI to the other site when the traffic is destined for the gateway router. If the gateway router at one site becomes unreachable, you must remove the filters so traffic can be sent to the other site's gateway router.



Reference for Clustering

This section includes more information about how clustering operates.

ASA Features and Clustering

Some ASA features are not supported with ASA clustering, and some are only supported on the control node. Other features might have caveats for proper usage.

Unsupported Features with Clustering

These features cannot be configured with clustering enabled, and the commands will be rejected.

- Unified Communication features that rely on TLS Proxy
- Remote access VPN (SSL VPN and IPsec VPN)
- Virtual Tunnel Interfaces (VTIs)
- IS-IS routing
- The following application inspections:
 - CTIQBE

- H323, H225, and RAS
 - IPsec passthrough
 - MGCP
 - MMP
 - RTSP
 - SCCP (Skinny)
 - WAAS
 - WCCP
-
- Botnet Traffic Filter
 - Auto Update Server
 - DHCP client, server, and proxy. DHCP relay is supported.
 - VPN load balancing
 - Failover
 - Integrated Routing and Bridging
 - Dead Connection Detection (DCD)
 - FIPS mode

Centralized Features for Clustering

The following features are only supported on the control node, and are not scaled for the cluster.



Note

Traffic for centralized features is forwarded from member nodes to the control node over the cluster control link.

If you use the rebalancing feature, traffic for centralized features may be rebalanced to non-control nodes before the traffic is classified as a centralized feature; if this occurs, the traffic is then sent back to the control node.

For centralized features, if the control node fails, all connections are dropped, and you have to re-establish the connections on the new control node.

- The following application inspections:
 - DCERPC
 - ESMTP
 - IM
 - NetBIOS
 - PPTP

- RADIUS
 - RSH
 - SNMP
 - SQLNET
 - SUNRPC
 - TFTP
 - XDMCP
- Static route monitoring
 - Authentication and Authorization for network access. Accounting is decentralized.
 - Filtering Services
 - Site-to-site VPN

In centralized mode, VPN connections are established with the control node of the cluster only. This is the default mode for VPN clustering. Site-to-site VPN can also be deployed in Distributed VPN Mode, where S2S IKEv2 VPN connections are distributed across nodes.
 - IGMP multicast control plane protocol processing (data plane forwarding is distributed across the cluster)
 - PIM multicast control plane protocol processing (data plane forwarding is distributed across the cluster)
 - Dynamic routing

Features Applied to Individual Units

These features are applied to each ASA node, instead of the cluster as a whole or to the control node.

- QoS—The QoS policy is synced across the cluster as part of configuration replication. However, the policy is enforced on each node independently. For example, if you configure policing on output, then the conform rate and conform burst values are enforced on traffic exiting a particular ASA. In a cluster with 3 nodes and with traffic evenly distributed, the conform rate actually becomes 3 times the rate for the cluster.
- Threat detection—Threat detection works on each node independently; for example, the top statistics is node-specific. Port scanning detection, for example, does not work because scanning traffic will be load-balanced between all nodes, and one node will not see all traffic.
- Resource management—Resource management in multiple context mode is enforced separately on each node based on local usage.
- LISP traffic—LISP traffic on UDP port 4342 is inspected by each receiving node, but is not assigned a director. Each node adds to the EID table that is shared across the cluster, but the LISP traffic itself does not participate in cluster state sharing.

AAA for Network Access and Clustering

AAA for network access consists of three components: authentication, authorization, and accounting. Authentication and authorization are implemented as centralized features on the clustering control node with

replication of the data structures to the cluster data nodes. If a control node is elected, the new control node will have all the information it needs to continue uninterrupted operation of the established authenticated users and their associated authorizations. Idle and absolute timeouts for user authentications are preserved when a control node change occurs.

Accounting is implemented as a distributed feature in a cluster. Accounting is done on a per-flow basis, so the cluster node owning a flow will send accounting start and stop messages to the AAA server when accounting is configured for a flow.

Connection Settings

Connection limits are enforced cluster-wide (see the **set connection conn-max**, **set connection embryonic-conn-max**, **set connection per-client-embryonic-max**, and **set connection per-client-max** commands). Each node has an estimate of the cluster-wide counter values based on broadcast messages. Due to efficiency considerations, the configured connection limit across the cluster might not be enforced exactly at the limit number. Each node may overestimate or underestimate the cluster-wide counter value at any given time. However, the information will get updated over time in a load-balanced cluster.

FTP and Clustering

- If FTP data channel and control channel flows are owned by different cluster members, then the data channel owner will periodically send idle timeout updates to the control channel owner and update the idle timeout value. However, if the control flow owner is reloaded, and the control flow is re-hosted, the parent/child flow relationship will not longer be maintained; the control flow idle timeout will not be updated.
- If you use AAA for FTP access, then the control channel flow is centralized on the control node.

ICMP Inspection

The flow of ICMP and ICMP error packets through the cluster varies depending on whether ICMP/ICMP error inspection is enabled. Without ICMP inspection, ICMP is a one-direction flow, and there is no director flow support. With ICMP inspection, the ICMP flow becomes two-directional and is backed up by a director/backup flow. One difference for an inspected ICMP flow is in the director handling of a forwarded packet: the director will forward the ICMP echo reply packet to the flow owner instead of returning the packet to the forwarder.

Multicast Routing and Clustering

The control unit handles all multicast routing packets and data packets until fast-path forwarding is established. After the connection is established, each data unit can forward multicast data packets.

NAT and Clustering

NAT can affect the overall throughput of the cluster. Inbound and outbound NAT packets can be sent to different ASAs in the cluster, because the load balancing algorithm relies on IP addresses and ports, and NAT causes inbound and outbound packets to have different IP addresses and/or ports. When a packet arrives at the ASA that is not the NAT owner, it is forwarded over the cluster control link to the owner, causing large amounts of traffic on the cluster control link. Note that the receiving node does not create a forwarding flow to the owner, because the NAT owner may not end up creating a connection for the packet depending on the results of security and policy checks.

If you still want to use NAT in clustering, then consider the following guidelines:

- PAT with Port Block Allocation—See the following guidelines for this feature:
 - Maximum-per-host limit is not a cluster-wide limit, and is enforced on each node individually. Thus, in a 3-node cluster with the maximum-per-host limit configured as 1, if the traffic from a host is load-balanced across all 3 nodes, then it can get allocated 3 blocks with 1 in each node.
 - Port blocks created on the backup node from the backup pools are not accounted for when enforcing the maximum-per-host limit.
 - On-the-fly PAT rule modifications, where the PAT pool is modified with a completely new range of IP addresses, will result in xlate backup creation failures for the xlate backup requests that were still in transit while the new pool became effective. This behavior is not specific to the port block allocation feature, and is a transient PAT pool issue seen only in cluster deployments where the pool is distributed and traffic is load-balanced across the cluster nodes.
 - When operating in a cluster, you cannot simply change the block allocation size. The new size is effective only after you reload each device in the cluster. To avoid having to reload each device, we recommend that you delete all block allocation rules and clear all xlates related to those rules. You can then change the block size and recreate the block allocation rules.
- NAT pool address distribution for dynamic PAT—When you configure a PAT pool, the cluster divides each IP address in the pool into port blocks. By default, each block is 512 ports, but if you configure port block allocation rules, your block setting is used instead. These blocks are distributed evenly among the nodes in the cluster, so that each node has one or more blocks for each IP address in the PAT pool. Thus, you could have as few as one IP address in a PAT pool for a cluster, if that is sufficient for the number of PAT'ed connections you expect. Port blocks cover the 1024-65535 port range, unless you configure the option to include the reserved ports, 1-1023, on the PAT pool NAT rule.
- Reusing a PAT pool in multiple rules—To use the same PAT pool in multiple rules, you must be careful about the interface selection in the rules. You must either use specific interfaces in all rules, or "any" in all rules. You cannot mix specific interfaces and "any" across the rules, or the system might not be able to match return traffic to the right node in the cluster. Using unique PAT pools per rule is the most reliable option.
- No round-robin—Round-robin for a PAT pool is not supported with clustering.
- No extended PAT—Extended PAT is not supported with clustering.
- Dynamic NAT xlates managed by the control node—The control node maintains and replicates the xlate table to data nodes. When a data node receives a connection that requires dynamic NAT, and the xlate is not in the table, it requests the xlate from the control node. The data node owns the connection.
- Stale xlates—The xlate idle time on the connection owner does not get updated. Thus, the idle time might exceed the idle timeout. An idle timer value higher than the configured timeout with a refcnt of 0 is an indication of a stale xlate.
- Per-session PAT feature—Although not exclusive to clustering, the per-session PAT feature improves the scalability of PAT and, for clustering, allows each data node to own PAT connections; by contrast, multi-session PAT connections have to be forwarded to and owned by the control node. By default, all TCP traffic and UDP DNS traffic use a per-session PAT xlate, whereas ICMP and all other UDP traffic uses multi-session. You can configure per-session NAT rules to change these defaults for TCP and UDP, but you cannot configure per-session PAT for ICMP. For example, with an increasing use of Quic protocol over UDP/443 as a best performance alternative compared to HTTPS TLS over TCP/443, you should enable per-session PAT for UDP/443. For traffic that benefits from multi-session PAT, such as H.323, SIP, or Skinny, you can disable per-session PAT for the associated TCP ports (the UDP ports for those

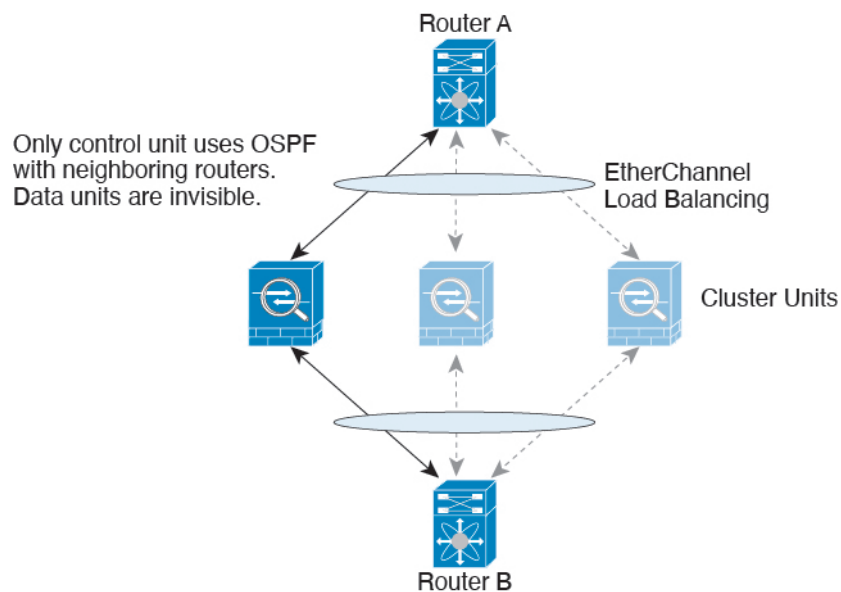
H.323 and SIP are already multi-session by default). For more information about per-session PAT, see the firewall configuration guide.

- No static PAT for the following inspections—
 - FTP
 - PPTP
 - RSH
 - SQLNET
 - TFTP
 - XDMCP
 - SIP
- If you have an extremely large number of NAT rules, over ten thousand, you should enable the transactional commit model using the **asp rule-engine transactional-commit nat** command in the device CLI. Otherwise, the node might not be able to join the cluster.

Dynamic Routing and Clustering

The routing process only runs on the control unit, and routes are learned through the control unit and replicated to secondaries. If a routing packet arrives at a data unit, it is redirected to the control unit.

Figure 1: Dynamic Routing



After the data units learn the routes from the control unit, each unit makes forwarding decisions independently.

The OSPF LSA database is not synchronized from the control unit to data units. If there is a control unit switchover, the neighboring router will detect a restart; the switchover is not transparent. The OSPF process picks an IP address as its router ID. Although not required, you can assign a static router ID to ensure a

consistent router ID is used across the cluster. See the OSPF Non-Stop Forwarding feature to address the interruption.

SCTP and Clustering

An SCTP association can be created on any node (due to load balancing); its multi-homing connections must reside on the same node.

SIP Inspection and Clustering

A control flow can be created on any node (due to load balancing); its child data flows must reside on the same node.

TLS Proxy configuration is not supported.

SNMP and Clustering

An SNMP agent polls each individual ASA by its Local IP address. You cannot poll consolidated data for the cluster.

You should always use the Local address, and not the Main cluster IP address for SNMP polling. If the SNMP agent polls the Main cluster IP address, if a new control node is elected, the poll to the new control node will fail.

When using SNMPv3 with clustering, if you add a new cluster node after the initial cluster formation, then SNMPv3 users are not replicated to the new node. You must re-add them on the control node to force the users to replicate to the new node, or directly on the data node.

STUN and Clustering

STUN inspection is supported in failover and cluster modes, as pinholes are replicated. However, the transaction ID is not replicated among nodes. In the case where a node fails after receiving a STUN Request and another node received the STUN Response, the STUN Response will be dropped.

Syslog and NetFlow and Clustering

- **Syslog**—Each node in the cluster generates its own syslog messages. You can configure logging so that each node uses either the same or a different device ID in the syslog message header field. For example, the hostname configuration is replicated and shared by all nodes in the cluster. If you configure logging to use the hostname as the device ID, syslog messages generated by all nodes look as if they come from a single node. If you configure logging to use the local-node name that is assigned in the cluster bootstrap configuration as the device ID, syslog messages look as if they come from different nodes.
- **NetFlow**—Each node in the cluster generates its own NetFlow stream. The NetFlow collector can only treat each ASA as a separate NetFlow exporter.

Cisco TrustSec and Clustering

Only the control node learns security group tag (SGT) information. The control node then populates the SGT to data nodes, and data nodes can make a match decision for SGT based on the security policy.

VPN and Clustering on the Secure Firewall eXtensible Operating System (FXOS) Chassis

An ASA FXOS Cluster supports one of two mutually exclusive modes for S2S VPN, centralized or distributed:

- **Centralized VPN Mode.** The default mode. In centralized mode, VPN connections are established with the control unit of the cluster only.

VPN functionality is limited to the control unit and does not take advantage of the cluster high availability capabilities. If the control unit fails, all existing VPN connections are lost, and VPN connected users see a disruption in service. When a new control unit is elected, you must reestablish the VPN connections.

When you connect a VPN tunnel to a Spanned interface address, connections are automatically forwarded to the control unit. VPN-related keys and certificates are replicated to all units.

- **Distributed VPN Mode.** In this mode, S2S IPsec IKEv2 VPN connections are distributed across members of an ASA cluster providing scalability. Distributing VPN connections across the members of a cluster allows both the capacity and throughput of the cluster to be fully utilized, significantly scaling VPN support beyond Centralized VPN capabilities.

**Note**

Centralized VPN clustering mode supports S2S IKEv1 and S2S IKEv2.

Distributed VPN clustering mode supports S2S IKEv2 only.

Distributed VPN clustering mode is supported on the Firepower 9300 only.

Remote access VPN is not supported in centralized or distributed VPN clustering mode.

Performance Scaling Factor

When you combine multiple units into a cluster, you can expect the total cluster performance to be approximately 80% of the maximum combined throughput.

For example, for TCP throughput, the Firepower 9300 with 3 SM-40 modules can handle approximately 135 Gbps of real world firewall traffic when running alone. For 2 chassis, the maximum combined throughput will be approximately 80% of 270 Gbps (2 chassis x 135 Gbps): 216 Gbps.

Control Unit Election

Members of the cluster communicate over the cluster control link to elect a control unit as follows:

1. When you deploy the cluster, each unit broadcasts an election request every 3 seconds.
2. Any other units with a higher priority respond to the election request; the priority is set when you deploy the cluster and is not configurable.
3. If after 45 seconds, a unit does not receive a response from another unit with a higher priority, then it becomes the control unit.

**Note**

If multiple units tie for the highest priority, the cluster unit name and then the serial number is used to determine the control unit.

4. If a unit later joins the cluster with a higher priority, it does not automatically become the control unit; the existing control unit always remains as the control unit unless it stops responding, at which point a new control unit is elected.

5. In a "split brain" scenario when there are temporarily multiple control units, then the unit with highest priority retains the role while the other units return to data unit roles.



Note You can manually force a unit to become the control unit. For centralized features, if you force a control unit change, then all connections are dropped, and you have to re-establish the connections on the new control unit.

High Availability Within the Cluster

Clustering provides high availability by monitoring chassis, unit, and interface health and by replicating connection states between units.

Chassis-Application Monitoring

Chassis-application health monitoring is always enabled. The Firepower 4100/9300 chassis supervisor checks the ASA application periodically (every second). If the ASA is up and cannot communicate with the Firepower 4100/9300 chassis supervisor for 3 seconds, the ASA generates a syslog message and leaves the cluster.

If the Firepower 4100/9300 chassis supervisor cannot communicate with the application after 45 seconds, it reloads the ASA. If the ASA cannot communicate with the supervisor, it removes itself from the cluster.

Unit Health Monitoring

Each unit periodically sends a broadcast keepaliveheartbeat packet over the cluster control link. If the control node does not receive any keepaliveheartbeat packets or other packets from a data node within the configurable timeout period, then the control node removes the data node from the cluster. If the data nodes do not receive packets from the control node, then a new control node is elected from the remaining node.

If nodes cannot reach each other over the cluster control link because of a network failure and not because a node has actually failed, then the cluster may go into a "split brain" scenario where isolated data nodes will elect their own control nodes. For example, if a router fails between two cluster locations, then the original control node at location 1 will remove the location 2 data nodes from the cluster. Meanwhile, the nodes at location 2 will elect their own control node and form their own cluster. Note that asymmetric traffic may fail in this scenario. After the cluster control link is restored, then the control node that has the higher priority will keep the control node's role. See [Control Unit Election, on page 81](#) for more information.

Interface Monitoring

Each node monitors the link status of all hardware interfaces in use, and reports status changes to the control node. For clustering on multiple chassis, Spanned EtherChannels use the cluster Link Aggregation Control Protocol (cLACP). Each chassis monitors the link status and the cLACP protocol messages to determine if the port is still active in the EtherChannel, and informs the ASA application if the interface is down. When you enable health monitoring, all physical interfaces are monitored by default (including the main EtherChannel for EtherChannel interfaces). Only named interfaces that are in an Up state can be monitored. For example, all member ports of an EtherChannel must fail before a *named* EtherChannel is removed from the cluster (depending on your minimum port bundling setting). You can optionally disable monitoring per interface.

If a monitored interface fails on a particular node, but it is active on other nodes, then the node is removed from the cluster. The amount of time before the ASA removes a node from the cluster depends on whether the node is an established member or is joining the cluster. The ASA does not monitor interfaces for the first

90 seconds that a node joins the cluster. Interface status changes during this time will not cause the ASA to be removed from the cluster. For an established member, the node is removed after 500 ms.

For clustering on multiple chassis, if you add or delete an EtherChannel from the cluster, interface health-monitoring is suspended for 95 seconds to ensure that you have time to make the changes on each chassis.

Decorator Application Monitoring

When you install a decorator application on an interface, such as the Radware DefensePro application, then both the ASA and the decorator application must be operational to remain in the cluster. The unit does not join the cluster until both applications are operational. Once in the cluster, the unit monitors the decorator application health every 3 seconds. If the decorator application is down, the unit is removed from the cluster.

Status After Failure

When a node in the cluster fails, the connections hosted by that node are seamlessly transferred to other nodes; state information for traffic flows is shared over the control node's cluster control link.

If the control node fails, then another member of the cluster with the highest priority (lowest number) becomes the control node.

The ASA automatically tries to rejoin the cluster, depending on the failure event.



Note

When the ASA becomes inactive and fails to automatically rejoin the cluster, all data interfaces are shut down; only the management-only interface can send and receive traffic. The management interface remains up using the IP address the node received from the cluster IP pool. However if you reload, and the node is still inactive in the cluster, the management interface is disabled. You must use the console port for any further configuration.

Rejoining the Cluster

After a cluster member is removed from the cluster, how it can rejoin the cluster depends on why it was removed:

- Failed cluster control link when initially joining—After you resolve the problem with the cluster control link, you must manually rejoin the cluster by re-enabling clustering at the ASA console port by entering **cluster group** *name*, and then **enable**.
- Failed cluster control link after joining the cluster—The ASA automatically tries to rejoin every 5 minutes, indefinitely. This behavior is configurable.
- Failed data interface—The ASA automatically tries to rejoin at 5 minutes, then at 10 minutes, and finally at 20 minutes. If the join is not successful after 20 minutes, then the ASA disables clustering. After you resolve the problem with the data interface, you have to manually enable clustering at the ASA console port by entering **cluster group** *name*, and then **enable**. This behavior is configurable.
- Failed unit—If the unit was removed from the cluster because of a unit health check failure, then rejoining the cluster depends on the source of the failure. For example, a temporary power failure means the unit will rejoin the cluster when it starts up again as long as the cluster control link is up. The unit attempts to rejoin the cluster every 5 seconds.

- **Failed Chassis-Application Communication**—When the ASA detects that the chassis-application health has recovered, the ASA tries to rejoin the cluster immediately. Alternatively, you can configure the ASA to use the same rejoin settings as for internal errors (below).
- **Failed decorator application**—The ASA rejoins the cluster when it senses that the decorator application is back up.
- **Internal error**—Internal failures include: application sync timeout; inconsistent application statuses; and so on. A unit will attempt to rejoin the cluster automatically at the following intervals: 5 minutes, 10 minutes, and then 20 minutes. This behavior is configurable.

Data Path Connection State Replication

Every connection has one owner and at least one backup owner in the cluster. The backup owner does not take over the connection in the event of a failure; instead, it stores TCP/UDP state information, so that the connection can be seamlessly transferred to a new owner in case of a failure. The backup owner is usually also the director.

Some traffic requires state information above the TCP or UDP layer. See the following table for clustering support or lack of support for this kind of traffic.

Table 1: Features Replicated Across the Cluster

Traffic	State Support	Notes
Up time	Yes	Keeps track of the system up time.
ARP Table	Yes	—
MAC address table	Yes	—
User Identity	Yes	Includes AAA rules (uauth).
IPv6 Neighbor database	Yes	—
Dynamic routing	Yes	—
SNMP Engine ID	No	—
Distributed VPN (Site-to-Site) for Firepower 4100/9300	Yes	Backup session becomes the active session, then a new backup session is created.

How the Cluster Manages Connections

Connections can be load-balanced to multiple nodes of the cluster. Connection roles determine how connections are handled in both normal operation and in a high availability situation.

Connection Roles

See the following roles defined for each connection:

- **Owner**—Usually, the node that initially receives the connection. The owner maintains the TCP state and processes packets. A connection has only one owner. If the original owner fails, then when new nodes receive packets from the connection, the director chooses a new owner from those nodes.

- **Backup owner**—The node that stores TCP/UDP state information received from the owner, so that the connection can be seamlessly transferred to a new owner in case of a failure. The backup owner does not take over the connection in the event of a failure. If the owner becomes unavailable, then the first node to receive packets from the connection (based on load balancing) contacts the backup owner for the relevant state information so it can become the new owner.

As long as the director (see below) is not the same node as the owner, then the director is also the backup owner. If the owner chooses itself as the director, then a separate backup owner is chosen.

For clustering on the Firepower 9300, which can include up to 3 cluster nodes in one chassis, if the backup owner is on the same chassis as the owner, then an additional backup owner will be chosen from another chassis to protect flows from a chassis failure.

If you enable director localization for inter-site clustering, then there are two backup owner roles: the local backup and the global backup. The owner always chooses a local backup at the same site as itself (based on site ID). The global backup can be at any site, and might even be the same node as the local backup. The owner sends connection state information to both backups.

If you enable site redundancy, and the backup owner is at the same site as the owner, then an additional backup owner will be chosen from another site to protect flows from a site failure. Chassis backup and site backup are independent, so in some cases a flow will have both a chassis backup and a site backup.

- **Director**—The node that handles owner lookup requests from forwarders. When the owner receives a new connection, it chooses a director based on a hash of the source/destination IP address and ports (see below for ICMP hash details), and sends a message to the director to register the new connection. If packets arrive at any node other than the owner, the node queries the director about which node is the owner so it can forward the packets. A connection has only one director. If a director fails, the owner chooses a new director.

As long as the director is not the same node as the owner, then the director is also the backup owner (see above). If the owner chooses itself as the director, then a separate backup owner is chosen.

If you enable director localization for inter-site clustering, then there are two director roles: the local director and the global director. The owner always chooses a local director at the same site as itself (based on site ID). The global director can be at any site, and might even be the same node as the local director. If the original owner fails, then the local director chooses a new connection owner at the same site.

ICMP/ICMPv6 hash details:

- For Echo packets, the source port is the ICMP identifier, and the destination port is 0.
- For Reply packets, the source port is 0, and the destination port is the ICMP identifier.
- For other packets, both source and destination ports are 0.

- **Forwarder**—A node that forwards packets to the owner. If a forwarder receives a packet for a connection it does not own, it queries the director for the owner, and then establishes a flow to the owner for any other packets it receives for this connection. The director can also be a forwarder. If you enable director localization, then the forwarder always queries the local director. The forwarder only queries the global director if the local director does not know the owner, for example, if a cluster member receives packets for a connection that is owned on a different site. Note that if a forwarder receives the SYN-ACK packet, it can derive the owner directly from a SYN cookie in the packet, so it does not need to query the director. (If you disable TCP sequence randomization, the SYN cookie is not used; a query to the director is required.) For short-lived flows such as DNS and ICMP, instead of querying, the forwarder immediately sends the packet to the director, which then sends them to the owner. A connection can have multiple

forwarders; the most efficient throughput is achieved by a good load-balancing method where there are no forwarders and all packets of a connection are received by the owner.



Note We do not recommend disabling TCP sequence randomization when using clustering. There is a small chance that some TCP sessions won't be established, because the SYN/ACK packet might be dropped.

- **Fragment Owner**—For fragmented packets, cluster nodes that receive a fragment determine a fragment owner using a hash of the fragment source IP address, destination IP address, and the packet ID. All fragments are then forwarded to the fragment owner over the cluster control link. Fragments may be load-balanced to different cluster nodes, because only the first fragment includes the 5-tuple used in the switch load balance hash. Other fragments do not contain the source and destination ports and may be load-balanced to other cluster nodes. The fragment owner temporarily reassembles the packet so it can determine the director based on a hash of the source/destination IP address and ports. If it is a new connection, the fragment owner will register to be the connection owner. If it is an existing connection, the fragment owner forwards all fragments to the provided connection owner over the cluster control link. The connection owner will then reassemble all fragments.

When a connection uses Port Address Translation (PAT), then the PAT type (per-session or multi-session) influences which member of the cluster becomes the owner of a new connection:

- **Per-session PAT**—The owner is the node that receives the initial packet in the connection.
By default, TCP and DNS UDP traffic use per-session PAT.
- **Multi-session PAT**—The owner is always the control node. If a multi-session PAT connection is initially received by a data node, then the data node forwards the connection to the control node.
By default, UDP (except for DNS UDP) and ICMP traffic use multi-session PAT, so these connections are always owned by the control node.

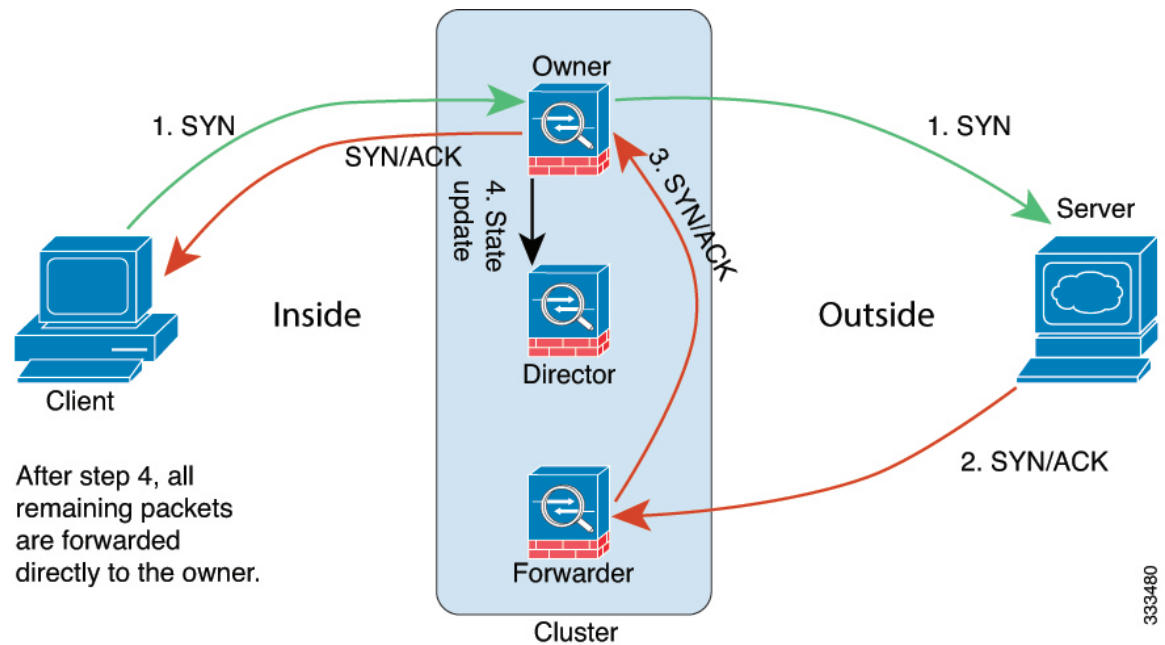
You can change the per-session PAT defaults for TCP and UDP so connections for these protocols are handled per-session or multi-session depending on the configuration. For ICMP, you cannot change from the default multi-session PAT. For more information about per-session PAT, see the firewall configuration guide.

New Connection Ownership

When a new connection is directed to a node of the cluster via load balancing, that node owns both directions of the connection. If any connection packets arrive at a different node, they are forwarded to the owner node over the cluster control link. If a reverse flow arrives at a different node, it is redirected back to the original node.

Sample Data Flow for TCP

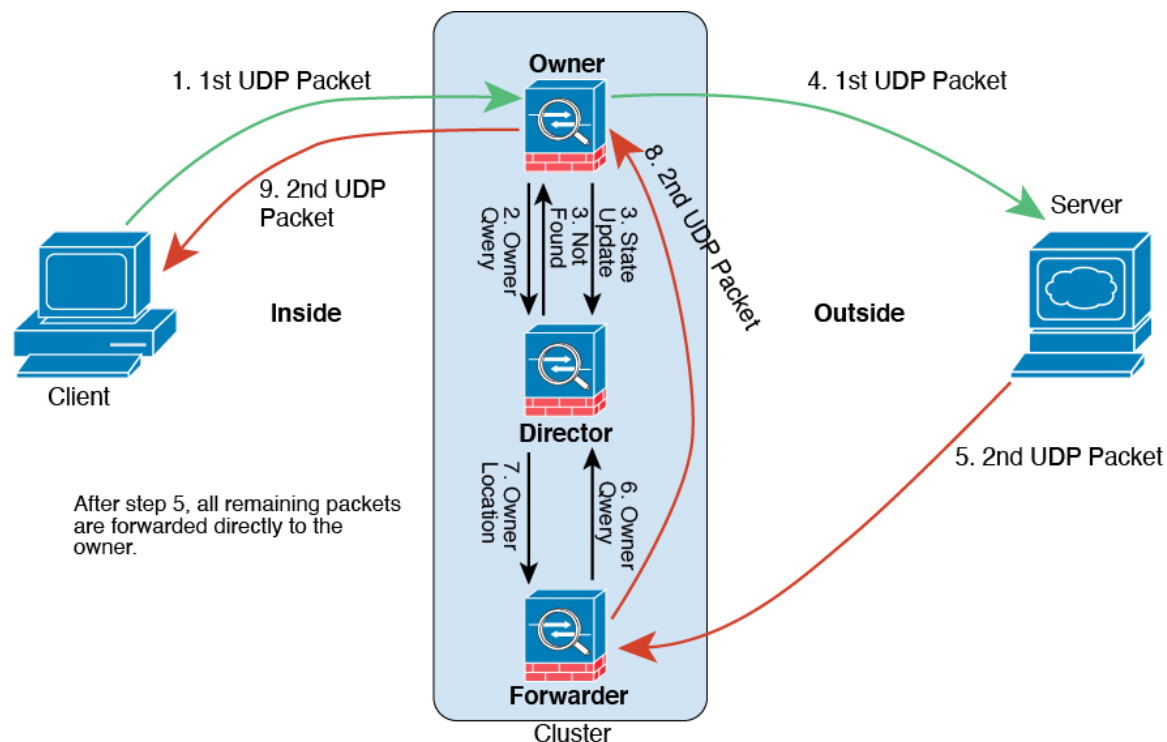
The following example shows the establishment of a new connection.



1. The SYN packet originates from the client and is delivered to one ASA (based on the load balancing method), which becomes the owner. The owner creates a flow, encodes owner information into a SYN cookie, and forwards the packet to the server.
2. The SYN-ACK packet originates from the server and is delivered to a different ASA (based on the load balancing method). This ASA is the forwarder.
3. Because the forwarder does not own the connection, it decodes owner information from the SYN cookie, creates a forwarding flow to the owner, and forwards the SYN-ACK to the owner.
4. The owner sends a state update to the director, and forwards the SYN-ACK to the client.
5. The director receives the state update from the owner, creates a flow to the owner, and records the TCP state information as well as the owner. The director acts as the backup owner for the connection.
6. Any subsequent packets delivered to the forwarder will be forwarded to the owner.
7. If packets are delivered to any additional nodes, it will query the director for the owner and establish a flow.
8. Any state change for the flow results in a state update from the owner to the director.

Sample Data Flow for ICMP and UDP

The following example shows the establishment of a new connection.

1. *Figure 2: ICMP and UDP Data Flow*

The first UDP packet originates from the client and is delivered to one ASA (based on the load balancing method).

2. The node that received the first packet queries the director node that is chosen based on a hash of the source/destination IP address and ports.
3. The director finds no existing flow, creates a director flow and forwards the packet back to the previous node. In other words, the director has elected an owner for this flow.
4. The owner creates the flow, sends a state update to the director, and forwards the packet to the server.
5. The second UDP packet originates from the server and is delivered to the forwarder.
6. The forwarder queries the director for ownership information. For short-lived flows such as DNS, instead of querying, the forwarder immediately sends the packet to the director, which then sends it to the owner.
7. The director replies to the forwarder with ownership information.
8. The forwarder creates a forwarding flow to record owner information and forwards the packet to the owner.
9. The owner forwards the packet to the client.

Rebalancing New TCP Connections Across the Cluster

If the load balancing capabilities of the upstream or downstream routers result in unbalanced flow distribution, you can configure new connection rebalancing so nodes with higher new connections per second will redirect new TCP flows to other nodes. No existing flows will be moved to other nodes.

Because this command only rebalances based on connections per second, the total number of established connections on each node is not considered, and the total number of connections may not be equal.

Once a connection is offloaded to a different node, it becomes an asymmetric connection.

Do not configure connection rebalancing for inter-site topologies; you do not want new connections rebalanced to cluster members at a different site.

History for ASA Clustering on the Firepower 4100/9300

Feature Name	Version	Feature Information
MTU ping test on node join	9.23(1)	When a node joins the cluster, it checks MTU compatibility by sending a ping to the control node with a packet size matching the cluster control link MTU. If the ping fails, a notification is generated so you can fix the MTU mismatch on connecting switches and try again.
Configurable delay to rejoin cluster after chassis heartbeat failure (Firepower 4100/9300)	9.20(2)	By default, if the chassis heartbeat fails and then recovers, the node rejoins the cluster immediately. However, if you configure the health-check chassis-heartbeat-delay-rejoin command, it will rejoin according to the settings of the health-check system auto-rejoin command. New/Modified commands: health-check chassis-heartbeat-delay-rejoin
Configurable cluster keepalive interval for flow status	9.20(1)	The flow owner sends keepalives (clu_heartbeat messages) and updates (clu_update messages) to the director and backup owner to refresh the flow state. You can now set the keepalive interval. The default is 15 seconds, and you can set the interval between 15 and 55 seconds. You may want to set the interval to be longer to reduce the amount of traffic on the cluster control link. New/Modified commands: clu-keepalive-interval
Removal of biased language	9.19(1)	Commands, command output, and syslog messages that contained the terms "Master" and "Slave" have been changed to "Control" and "Data." New/Modified commands: cluster control-node, enable as-data-node, prompt, show cluster history, show cluster info
Improved PAT port block allocation for clustering on the Firepower 4100/9300	9.16(1)	The improved PAT port block allocation ensures that the control unit keeps ports in reserve for joining nodes, and proactively reclaims unused ports. To best optimize the allocation, you can set the maximum nodes you plan to have in the cluster using the cluster-member-limit command. The control unit can then allocate port blocks to the planned number of nodes, and it will not have to reserve ports for extra nodes you don't plan to use. The default is 16 nodes. You can also monitor syslog 747046 to ensure that there are enough ports available for a new node. New/Modified commands: cluster-member-limit, show nat pool cluster [summary], show nat pool ip detail
show cluster history command improvements	9.16(1)	We have added additional outputs for the show cluster history command. New/Modified commands: show cluster history brief, show cluster history latest, show cluster history reverse, show cluster history time
Configuration sync to data units in parallel	9.14(1)	The control unit now syncs configuration changes with data units in parallel by default. Formerly, syncing occurred sequentially. New/Modified commands: config-replicate-parallel

Feature Name	Version	Feature Information
Messages for cluster join failure or eviction added to show cluster history	9.14(1)	New messages were added to the show cluster history command for when a cluster unit either fails to join the cluster or leaves the cluster. New/Modified commands: show cluster history
Initiator and responder information for Dead Connection Detection (DCD), and DCD support in a cluster.	9.13(1)	If you enable Dead Connection Detection (DCD), you can use the show conn detail command to get information about the initiator and responder. Dead Connection Detection allows you to maintain an inactive connection, and the show conn output tells you how often the endpoints have been probed. In addition, DCD is now supported in a cluster. New/Modified commands: show conn (output only).
Monitor the traffic load for a cluster	9.13(1)	You can now monitor the traffic load for cluster members, including total connection count, CPU and memory usage, and buffer drops. If the load is too high, you can choose to manually disable clustering on the unit if the remaining units can handle the load, or adjust the load balancing on the external switch. This feature is enabled by default. New/Modified commands: debug cluster load-monitor , load-monitor , show cluster info load-monitor
Accelerated cluster joining	9.13(1)	When a data unit has the same configuration as the control unit, it will skip syncing the configuration and will join faster. This feature is enabled by default. This feature is configured on each unit, and is not replicated from the control unit to the data unit. Note Some configuration commands are not compatible with accelerated cluster joining; if these commands are present on the unit, even if accelerated cluster joining is enabled, configuration syncing will always occur. You must remove the incompatible configuration for accelerated cluster joining to work. Use the show cluster info unit-join-acceleration incompatible-config to view incompatible configuration. New/Modified commands: unit join-acceleration , show cluster info unit-join-acceleration incompatible-config
Per-site gratuitous ARP for clustering	9.12(1)	The ASA now generates gratuitous ARP (GARP) packets to keep the switching infrastructure up to date: the highest priority member at each site periodically generates GARP traffic for the global MAC/IP addresses. When using per-site MAC and IP addresses, packets sourced from the cluster use a site-specific MAC address and IP address, while packets received by the cluster use a global MAC address and IP address. If traffic is not generated from the global MAC address periodically, you could experience a MAC address timeout on your switches for the global MAC address. After a timeout, traffic destined for the global MAC address will be flooded across the entire switching infrastructure, which can cause performance and security concerns. GARP is enabled by default when you set the site ID for each unit and the site MAC address for each Spanned EtherChannel. New/Modified commands: site-periodic-garp interval
Parallel cluster joining of units per Firepower 9300 chassis	9.10(1)	For the Firepower 9300, this feature ensures that the security modules in a chassis join the cluster simultaneously, so that traffic is evenly distributed between the modules. If a module joins very much in advance of other modules, it can receive more traffic than desired, because the other modules cannot yet share the load. New/Modified commands: unit parallel-join

Feature Name	Version	Feature Information
Cluster control link customizable IP Address for the Firepower 4100/9300	9.10(1)	By default, the cluster control link uses the 127.2.0.0/16 network. You can now set the network when you deploy the cluster in FXOS. The chassis auto-generates the cluster control link interface IP address for each unit based on the chassis ID and slot ID: 127.2.<i>chassis_id.slot_id</i>. However, some networking deployments do not allow 127.2.0.0/16 traffic to pass. Therefore, you can now set a custom /16 subnet for the cluster control link in FXOS except for loopback (127.0.0.0/8) and multicast (224.0.0.0/4) addresses. New/Modified FXOS commands: set cluster-control-link network
Cluster interface debounce time now applies to interfaces changing from a down state to an up state	9.10(1)	When an interface status update occurs, the ASA waits the number of milliseconds specified in the health-check monitor-interface debounce-time command or the ASDM Configuration > Device Management > High Availability and Scalability > ASA Cluster screen before marking the interface as failed and the unit is removed from the cluster. This feature now applies to interfaces changing from a down state to an up state. For example, in the case of an EtherChannel that transitions from a down state to an up state (for example, the switch reloaded, or the switch enabled an EtherChannel), a longer debounce time can prevent the interface from appearing to be failed on a cluster unit just because another cluster unit was faster at bundling the ports. We did not modify any commands.
Automatically rejoin the cluster after an internal failure	9.9(2)	Formerly, many error conditions caused a cluster unit to be removed from the cluster, and you were required to manually rejoin the cluster after resolving the issue. Now, a unit will attempt to rejoin the cluster automatically at the following intervals by default: 5 minutes, 10 minutes, and then 20 minutes. These values are configurable. Internal failures include: application sync timeout; inconsistent application statuses; and so on. New or Modified commands: health-check system auto-rejoin, show cluster info auto-join
Show transport related statistics for cluster reliable transport protocol messages	9.9(2)	You can now view per-unit cluster reliable transport buffer usage so you can identify packet drop issues when the buffer is full in the control plane. New or modified command: show cluster info transport cp detail
cluster remove unit command behavior matches no enable behavior	9.9(1)	The cluster remove unit command now removes a unit from the cluster until you manually reenables clustering or reload, similar to the no enable command. Previously, if you redeployed the bootstrap configuration from FXOS, clustering would be reenabled. Now, the disabled status persists even in the case of a bootstrap configuration redeployment. Reloading the ASA, however, will reenables clustering. New/Modified command: cluster remove unit
Improved chassis health check failure detection for the chassis	9.9(1)	You can now configure a lower holdtime for the chassis health check: 100 ms. The previous minimum was 300 ms. Note that the minimum combined time (<i>interval x retry-count</i>) cannot be less than 600 ms. New or modified command: app-agent heartbeat interval
Inter-site redundancy for clustering	9.9(1)	Inter-site redundancy ensures that a backup owner for a traffic flow will always be at the other site from the owner. This feature guards against site failure. New or modified command: site-redundancy, show asp cluster counter change, show asp table cluster chash-table, show conn flag

Feature Name	Version	Feature Information
Distributed Site-to-Site VPN with clustering on the Firepower 9300	9.9(1)	An ASA cluster on the Firepower 9300 supports Site-to-Site VPN in distributed mode. Distributed mode provides the ability to have many Site-to-Site IPsec IKEv2 VPN connections distributed across members of an ASA cluster, not just on the control unit (as in centralized mode). This significantly scales VPN support beyond Centralized VPN capabilities and provides high availability. Distributed S2S VPN runs on a cluster of up to two chassis, each containing up to three modules (six total cluster members), each module supporting up to 6K active sessions (12K total), for a maximum of approximately 36K active sessions (72K total). New or modified commands: cluster redistribute vpn-sessiondb, show cluster vpn-sessiondb, vpn mode, show cluster resource usage, show vpn-sessiondb, show connection detail, show crypto ikev2
Improved cluster unit health-check failure detection	9.8(1)	You can now configure a lower holdtime for the unit health check: .3 seconds minimum. The previous minimum was .8 seconds. This feature changes the unit health check messaging scheme to <i>heartbeats</i> in the data plane from <i>keepalives</i> in the control plane. Using heartbeats improves the reliability and the responsiveness of clustering by not being susceptible to control plane CPU hogging and scheduling delays. Note that configuring a lower holdtime increases cluster control link messaging activity. We suggest that you analyze your network before you configure a low holdtime; for example, make sure a ping from one unit to another over the cluster control link returns within the <i>holdtime</i>/3, because there will be three heartbeat messages during one holdtime interval. If you downgrade your ASA software after setting the hold time to .3 - .7, this setting will revert to the default of 3 seconds because the new setting is unsupported. We modified the following commands: health-check holdtime, show asp drop cluster counter, show cluster info health details
Configurable debounce time to mark an interface as failed for the Firepower 4100/9300 chassis	9.8(1)	You can now configure the debounce time before the ASA considers an interface to be failed, and the unit is removed from the cluster. This feature allows for faster detection of interface failures. Note that configuring a lower debounce time increases the chances of false-positives. When an interface status update occurs, the ASA waits the number of milliseconds specified before marking the interface as failed and the unit is removed from the cluster. The default debounce time is 500 ms, with a range of 300 ms to 9 seconds. New or modified command: health-check monitor-interface debounce-time
Inter-site clustering improvement for the ASA on the Firepower 4100/9300 chassis	9.7(1)	You can now configure the site ID for each Firepower 4100/9300 chassis when you deploy the ASA cluster. Previously, you had to configure the site ID within the ASA application; this new feature eases initial deployment. Note that you can no longer set the site ID within the ASA configuration. Also, for best compatibility with inter-site clustering, we recommend that you upgrade to ASA 9.7(1) and FXOS 2.1.1, which includes several improvements to stability and performance. We modified the following command: site-id

Feature Name	Version	Feature Information
Director localization: inter-site clustering improvement for data centers	9.7(1)	To improve performance and keep traffic within a site for inter-site clustering for data centers, you can enable director localization. New connections are typically load-balanced and owned by cluster members within a given site. However, the ASA assigns the director role to a member at <i>any</i> site. Director localization enables additional director roles: a local director at the same site as the owner, and a global director that can be at any site. Keeping the owner and director at the same site improves performance. Also, if the original owner fails, the local director chooses a new connection owner at the same site. The global director is used if a cluster member receives packets for a connection that is owned on a different site. We introduced or modified the following commands: director-localization, show asp table cluster chash, show conn, show conn detail
Support for 16 chassis for the Firepower 4100 series	9.6(2)	You can now add up to 16 chassis to the cluster for the Firepower 4100 series. We did not modify any commands.
Support for the Firepower 4100 series	9.6(1)	With FXOS 1.1.4, the ASA supports inter-chassis clustering on the Firepower 4100 series for up to 6 chassis. We did not modify any commands.
Support for site-specific IP addresses in Routed, Spanned EtherChannel mode	9.6(1)	For inter-site clustering in routed mode with Spanned EtherChannels, you can now configure site-specific IP addresses in addition to site-specific MAC addresses. The addition of site IP addresses allows you to use ARP inspection on the Overlay Transport Virtualization (OTV) devices to prevent ARP responses from the global MAC address from traveling over the Data Center Interconnect (DCI), which can cause routing problems. ARP inspection is required for some switches that cannot use VACLs to filter MAC addresses. We modified the following commands: mac-address, show interface
Inter-chassis clustering for 16 modules, and inter-site clustering for the Firepower 9300 ASA application	9.5(2.1)	With FXOS 1.1.3, you can now enable inter-chassis, and by extension inter-site clustering. You can include up to 16 modules. For example, you can use 1 module in 16 chassis, or 2 modules in 8 chassis, or any combination that provides a maximum of 16 modules. We did not modify any commands.
Site-specific MAC addresses for inter-site clustering support for Spanned EtherChannel in Routed firewall mode	9.5(2)	You can now use inter-site clustering for Spanned EtherChannels in routed mode. To avoid MAC address flapping, configure a site ID for each cluster member so that a site-specific MAC address for each interface can be shared among a site's units. We introduced or modified the following commands: site-id, mac-address site-id, show cluster info, show interface
ASA cluster customization of the auto-rejoin behavior when an interface or the cluster control link fails	9.5(2)	You can now customize the auto-rejoin behavior when an interface or the cluster control link fails. We introduced the following command: health-check auto-rejoin
The ASA cluster supports GTPv1 and GTPv2	9.5(2)	The ASA cluster now supports GTPv1 and GTPv2 inspection. We did not modify any commands.

Feature Name	Version	Feature Information
Cluster replication delay for TCP connections	9.5(2)	This feature helps eliminate the “unnecessary work” related to short-lived flows by delaying the director/backup flow creation. We introduced the following command: cluster replication delay
LISP Inspection for Inter-Site Flow Mobility	9.5(2)	Cisco Locator/ID Separation Protocol (LISP) architecture separates the device identity from its location into two different numbering spaces, making server migration transparent to clients. The ASA can inspect LISP traffic for location changes and then use this information for seamless clustering operation; the ASA cluster members inspect LISP traffic passing between the first hop router and the egress tunnel router (ETR) or ingress tunnel router (ITR), and then change the flow owner to be at the new site. We introduced or modified the following commands: allowed-eid, clear cluster info flow-mobility counters, clear lisp eid, cluster flow-mobility lisp, debug cluster flow-mobility, debug lisp eid-notify-intercept, flow-mobility lisp, inspect lisp, policy-map type inspect lisp, site-id, show asp table classify domain inspect-lisp, show cluster info flow-mobility counters, show conn, show lisp eid, show service-policy, validate-key
Carrier Grade NAT enhancements now supported in failover and ASA clustering	9.5(2)	For carrier-grade or large-scale PAT, you can allocate a block of ports for each host, rather than have NAT allocate one port translation at a time (see RFC 6888). This feature is now supported in failover and ASA cluster deployments. We modified the following command: show local-host
Configurable level for clustering trace entries	9.5(2)	By default, all levels of clustering events are included in the trace buffer, including many low level events. To limit the trace to higher level events, you can set the minimum trace level for the cluster. We introduced the following command: trace-level
Intra-chassis ASA Clustering for the Firepower 9300	9.4(1.150)	You can cluster up to 3 security modules within the Firepower 9300 chassis. All modules in the chassis must belong to the cluster. We introduced the following commands: cluster replication delay, debug service-module, management-only individual, show cluster chassis