



Grappe ASA pour le Firepower 4100/9300

La mise en grappe vous permet de regrouper plusieurs Châssis Firepower 4100/9300 en un seul périphérique logique. La série Châssis Firepower 4100/9300 comprend les Firepower 9300 et Firepower 4100. Une grappe offre toute la commodité d'un seul appareil (gestion, intégration dans un réseau), tout en offrant le débit accru et la redondance de plusieurs périphériques.



Remarque

Certaines fonctionnalités ne sont pas prises en charge lors de l'utilisation de la mise en grappe. Consultez [Fonctionnalités non prises en charge par la mise en grappe](#), à la page 78.

- [À propos de la mise en grappe sur les châssis Firepower 4100/9300](#), à la page 1
- [Exigences et conditions préalables à la mise en grappe sur Châssis Firepower 4100/9300](#), à la page 8
- [Licences pour la mise en grappe sur le Châssis Firepower 4100/9300](#), à la page 10
- [Lignes directrices et limites de la mise en grappe](#), à la page 11
- [Configurer la mise en grappe sur le Châssis Firepower 4100/9300](#), à la page 16
- [FXOS : supprimer un nœud de la grappe](#), à la page 51
- [ASA : Gérer les membres de la grappe](#), à la page 53
- [ASA : supervision de la grappe ASA sur le Châssis Firepower 4100/9300](#), à la page 57
- [Dépannage du VPN S2S distribué](#), à la page 68
- [Exemples de mise en grappe d'ASA](#), à la page 70
- [Référence pour la mise en grappe](#), à la page 78
- [Historique de la mise en grappe d'ASA sur le Firepower 4100/9300](#), à la page 94

À propos de la mise en grappe sur les châssis Firepower 4100/9300

Lorsque vous déployez une grappe sur le Châssis Firepower 4100/9300, elle effectue les opérations suivantes :

- crée une *liaison de commande de grappe* (par défaut, le canal de port 48) pour la communication de nœud à nœud.

Pour une grappe isolée des modules de sécurité dans un châssis Firepower 9300, ce lien utilise le fond de panier Firepower 9300 pour les communications de la grappe.

Pour la mise en grappe avec plusieurs châssis, vous devez affecter manuellement une ou plusieurs interfaces physiques à cet EtherChannel pour les communications entre les châssis.

- Crée la configuration de démarrage de grappe dans l'application.

Lorsque vous déployez la grappe, le superviseur de châssis envoie une configuration de démarrage minimale à chaque unité, qui comprend le nom de la grappe, l'interface de liaison de commande de grappe et d'autres paramètres de la grappe. Certaines parties de la configuration de démarrage peuvent être configurées par l'utilisateur dans l'application si vous souhaitez personnaliser votre environnement de mise en grappe.

- Affecte des interfaces de données à la grappe en tant *qu'interfaces étendues*.

Pour une grappe isolée des modules de sécurité dans un châssis Firepower 9300, les interfaces étendues ne se limitent pas aux EtherChannels, comme c'est le cas pour la mise en grappe avec plusieurs châssis. Le superviseur Firepower 9300 utilise la technologie EtherChannel en interne pour équilibrer la charge du trafic vers plusieurs modules sur une interface partagée, de sorte que tout type d'interface de données fonctionne pour le mode étendu. Pour la mise en grappe avec plusieurs châssis, vous devez utiliser des EtherChannels étendus pour toutes les interfaces de données.

**Remarque**

Les interfaces individuelles ne sont pas prises en charge, à l'exception d'une interface de gestion.

- Attribue une interface de gestion à toutes les unités de la grappe.

Voir l'une des sections suivantes pour plus d'informations sur la mise en grappe.

Configuration du démarrage

Lorsque vous déployez la grappe, le superviseur de châssis Firepower 4100/9300 envoie une configuration de démarrage minimale à chaque unité, qui comprend le nom de la grappe, l'interface de liaison de commande de grappe et d'autres paramètres de la grappe. Certaines parties de la configuration de démarrage sont configurables par l'utilisateur si vous souhaitez personnaliser votre environnement de mise en grappe.

Membres de la grappe

Les membres de la grappe collaborent pour partager la politique de sécurité et les flux de trafic.

L'unité de **contrôle** est l'un des membres de la grappe. L'unité de contrôle est déterminée automatiquement. Tous les autres membres sont des unités de **données**.

Vous devez effectuer toute la configuration sur l'unité de contrôle uniquement; la configuration est ensuite reproduite dans les unités de données.

Certaines fonctionnalités ne sont pas évolutives dans une grappe, et l'unité de contrôle gère tout le trafic pour ces fonctionnalités. Voir [Fonctionnalités centralisées pour la mise en grappe, à la page 79](#).

Liaison de commande de grappe

La liaison de commande de grappe est un EtherChannel (canal de port 48) pour la communication d'unité à unité. Pour la mise en grappe intra-châssis, cette liaison utilise le fond de panier Firepower 9300 pour les

communications de la grappe. Pour la mise en grappe inter-châssis, vous devez affecter manuellement une ou plusieurs interfaces physiques à cet EtherChannel sur le Châssis Firepower 4100/9300 pour les communications entre les châssis.

Dans le cas d'une grappe inter-châssis à deux châssis, ne connectez pas directement la liaison de commande de grappe d'un châssis à l'autre. Si vous connectez directement les interfaces, lorsqu'une unité tombe en panne, la liaison de commande de grappe tombe en panne, et donc l'unité intègre restante. Si vous connectez la liaison de commande de grappe par l'intermédiaire d'un commutateur, cette dernière reste active pour l'unité intègre.

Le trafic de liaison de commande de grappe comprend à la fois un trafic de contrôle et un trafic de données.

Le trafic de contrôle comprend :

- Choix du nœud de contrôle.
- Duplication de la configuration.
- Surveillance de l'intégrité

Le trafic de données comprend :

- Duplication de l'état.
- Requêtes de propriété de connexion et transfert de paquets de données.

Consultez les sections suivantes pour en savoir plus sur la liaison de commande de grappe.

Dimensionner la liaison de commande de grappe

Si possible, vous devez dimensionner la liaison de commande de grappe en fonction du débit attendu de chaque châssis afin que la liaison de commande de grappe puisse gérer les scénarios les plus défavorables.

Le trafic de liaison de commande de grappe est principalement composé de mises à jour d'état et de paquets transférés. Le volume de trafic varie à un moment donné sur la liaison de commande de grappe. La quantité de trafic transféré dépend de l'efficacité de l'équilibrage de la charge et de l'importance du trafic pour les fonctionnalités centralisées. Par exemple :

- La NAT entraîne un mauvais équilibrage de la charge des connexions et la nécessité de rééquilibrer tout le trafic de retour vers les bonnes unités.
- L'AAA pour l'accès réseau est une fonctionnalité centralisée, de sorte que tout le trafic est acheminé à l'unité de contrôle.
- Lorsque les membres changent, la grappe doit rééquilibrer un grand nombre de connexions, utilisant ainsi temporairement une grande quantité de bande passante de la liaison de commande de grappe.

Une liaison de commande de grappe à bande passante plus élevée aide la grappe à converger plus rapidement lorsque les membres changent et empêche les goulots d'étranglement.



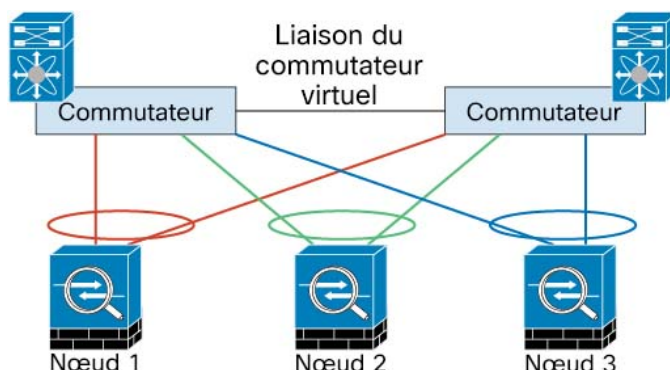
Remarque

Si votre grappe génère un trafic asymétrique (rééquilibrer) important, vous devez augmenter la taille du lien de commande de grappe.

Redondance de la liaison de commande de la grappe

Nous vous recommandons d'utiliser un EtherChannel pour la liaison de commande de la grappe, afin de pouvoir transférer le trafic sur plusieurs liaisons de l'EtherChannel tout en assurant la redondance.

Le diagramme suivant montre comment utiliser un EtherChannel comme liaison de commande de la grappe dans un système de commutation virtuelle (VSS), un canal de port virtuel (vPC), un StackWise ou un environnement StackWise Virtual. Tous les liens de l'EtherChannel sont actifs. Lorsque le commutateur fait partie d'un système redondant, vous pouvez connecter des interfaces de pare-feu dans le même EtherChannel pour séparer les commutateurs du système redondant. Les interfaces des commutateurs sont membres de la même interface de canal de port EtherChannel, car les commutateurs distincts se comportent comme un seul commutateur. Notez qu'il s'agit d'un EtherChannel local au périphérique et non d'un EtherChannel étendu.



Fiabilité de la liaison de commande de grappe pour la mise en grappe

Pour assurer la fonctionnalité de la liaison de commande de grappe, vérifiez que le temps aller-retour (RTT) entre les unités est inférieur à 20 ms. Cette latence maximale améliore la compatibilité avec les membres de la grappe installés à différents sites géographiques. Pour vérifier votre latence, envoyez un message Ping sur la liaison de commande de grappe entre les unités.

La liaison de commande de grappe doit être fiable, sans paquets en désordre ou abandonnés; par exemple, pour un déploiement intersite, vous devez utiliser un lien dédié.

Réseau de liaison de commande de grappe

Le Châssis Firepower 4100/9300 génère automatiquement l'adresse IP de l'interface de liaison de commande de grappe pour chaque unité en fonction de l'ID de châssis et de l'ID d'emplacement : `127.2.chassis_id.slot-id`. Vous pouvez personnaliser cette adresse IP lorsque vous déployez la grappe. Le réseau de liaison de commande de grappe ne peut pas comprendre de routeurs entre les unités; seule la commutation de couche 2 est autorisée. Pour le trafic intersites, Cisco recommande d'utiliser la virtualisation du transport par superposition (OTV).

Interfaces de la grappe

Pour une grappe isolée des modules de sécurité dans un châssis Firepower 9300, vous pouvez affecter des interfaces physiques ou des EtherChannels (également appelés canaux de port) à la grappe. Les interfaces affectées à la grappe sont des interfaces étendues qui équilibrent la charge du trafic entre tous les membres de la grappe.

Pour la mise en grappe avec plusieurs châssis, vous pouvez uniquement affecter des EtherChannels de données à la grappe. Ces EtherChannels étendus comprennent les mêmes interfaces membre sur chaque châssis; Sur le commutateur en amont, toutes ces interfaces sont incluses dans un seul EtherChannel, de sorte que le commutateur ne sache pas qu'il est connecté à plusieurs périphériques.

Les interfaces individuelles ne sont pas prises en charge, à l'exception d'une interface de gestion.

Connexion à un système de commutateurs redondants

Nous vous recommandons de connecter les EtherChannels à un système de commutation redondant tel qu'un VSS, vPC, StackWise ou un système virtuel StackWise pour assurer la redondance de vos interfaces.

Réplication de la configuration

Tous les nœuds de la grappe partagent une configuration unique. Vous pouvez uniquement apporter des modifications à la configuration sur le nœud de contrôle (à l'exception de la configuration de démarrage) et les modifications sont automatiquement synchronisées avec tous les autres nœuds de la grappe.

Gestion des grappes Cisco Secure Firewall ASA

L'un des avantages de l'utilisation de la mise en grappe ASA est la facilité de gestion. Cette section décrit comment gérer la grappe.

Réseau de gestion

Nous vous recommandons de connecter toutes les unités à un seul réseau de gestion. Ce réseau est distinct de la liaison de commande de grappe.

Interface de gestion

Vous devez affecter une interface de type gestion à la grappe. Cette interface est une interface individuelle spéciale, par opposition à une interface étendue. L'interface de gestion vous permet de vous connecter directement à chaque unité.

L'adresse IP de la grappe principale est une adresse fixe pour la grappe qui appartient toujours à l'unité de contrôle actuelle. Vous configurez également une plage d'adresses de sorte que chaque unité, y compris l'unité de contrôle actuelle, puisse utiliser une adresse locale de la plage. L'adresse IP de la grappe principale fournit un accès de gestion cohérent à une adresse; lorsqu'une unité de contrôle change, l'adresse IP de la grappe principale est déplacée vers la nouvelle unité de contrôle, de sorte que la gestion de la grappe se poursuit de façon transparente.

Par exemple, vous pouvez gérer la grappe en vous connectant à l'adresse IP de la grappe principale, qui est toujours associée à l'unité de contrôle actuelle. Pour gérer un membre, vous pouvez vous connecter à l'adresse IP locale.



Remarque

Le trafic vers la boîte doit être acheminé vers l'adresse IP de gestion du nœud; le trafic vers la boîte n'est pas transféré sur la liaison de commande de grappe vers un autre nœud.

Pour le trafic de gestion sortant tel que TFTP ou syslog, chaque unité, y compris l'unité de contrôle, utilise l'adresse IP locale pour se connecter au serveur.

Gestion de l'unité de contrôle vs Gestion de l'unité de données

Toutes la gestion et la surveillance peuvent avoir lieu sur le nœud de contrôle. À partir du nœud de contrôle, vous pouvez vérifier les statistiques d'exécution, l'utilisation des ressources ou d'autres informations de surveillance sur tous les nœuds. Vous pouvez également transmettre une commande à tous les nœuds de la grappe et reproduire les messages de console des nœuds de données vers le nœud de contrôle.

Vous pouvez surveiller directement les nœuds de données si vous le souhaitez. Bien que cela soit également disponible à partir du nœud de contrôle, vous pouvez effectuer la gestion de fichiers sur les nœuds de données (y compris la sauvegarde de la configuration et la mise à jour des images). Les fonctions suivantes ne sont pas disponibles à partir du nœud de commande :

- Surveillance des statistiques propres à la grappe par nœud.
- Surveillance des journaux système par nœud (sauf pour les journaux système envoyés à la console lorsque la duplication de la console est activée).
- SNMP
- NetFlow

duplication de clé de chiffrement

Lorsque vous créez une clé de chiffrement sur le nœud de contrôle, la clé est répliquée sur tous les nœuds de données. Si vous avez une session SSH sur l'adresse IP de la grappe principale, vous serez déconnecté en cas de défaillance du nœud de contrôle. Le nouveau nœud de contrôle utilise la même clé pour les connexions SSH, de sorte que vous n'avez pas besoin de mettre à jour la clé d'hôte SSH mise en cache lorsque vous vous reconnectez au nouveau nœud de contrôle.

Incompatibilité d'adresse IP du certificat de connexion ASDM

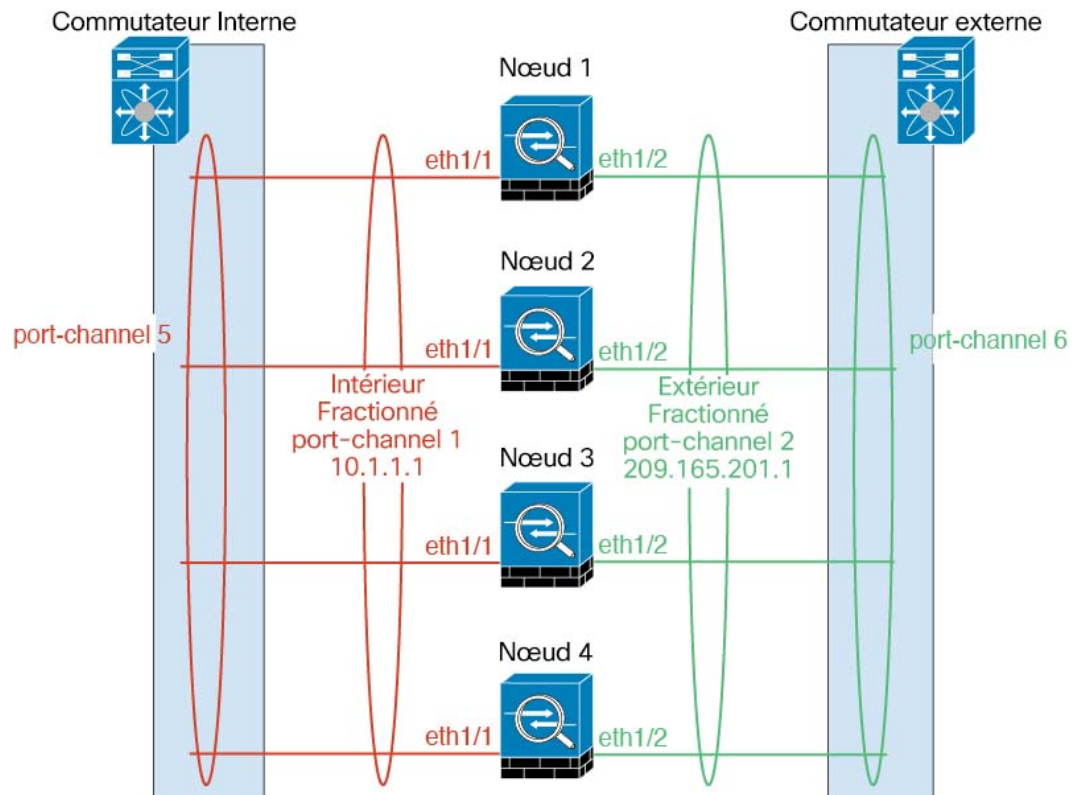
Par défaut, un certificat autosigné est utilisé pour la connexion ASDM en fonction de l'adresse IP locale. Si vous vous connectez à l'adresse IP de la grappe principale à l'aide d'ASDM, un message d'avertissement concernant une adresse IP non concordante peut s'afficher, car le certificat utilise l'adresse IP locale, et non l'adresse IP de la grappe principale. Vous pouvez ignorer le message et établir la connexion ASDM. Cependant, pour éviter ce type d'avertissement, vous pouvez inscrire un certificat qui contient l'adresse IP de la grappe principale et toutes les adresses IP locales de l'ensemble d'adresses IP. Vous pouvez ensuite utiliser ce certificat pour chaque membre de la grappe. Consultez <https://www.cisco.com/c/en/us/td/docs/security/asdm/identity-cert/cert-install.html> pour de plus amples renseignements.

EtherChannels étendus (recommandé)

Vous pouvez regrouper une ou plusieurs interfaces par châssis dans un EtherChannel qui s'étend sur tous les châssis de la grappe. L'EtherChannel agrège le trafic sur toutes les interfaces actives disponibles dans le canal.

Un EtherChannel étendu peut être configuré dans les modes de pare-feu avec routage et transparent. En mode routé, l'EtherChannel est configuré comme une interface routée avec une seule adresse IP. En mode transparent, l'adresse IP est attribuée aux BVI, et non à l'interface du membre du groupe de ponts.

L'EtherChannel assure intrinsèquement l'équilibrage de la charge dans le cadre du fonctionnement de base.



Mise en grappe inter-sites

Pour les installations inter-sites, vous pouvez profiter de la mise en grappe d'ASA tant que vous suivez les lignes directrices recommandées.

Vous pouvez configurer chaque châssis de grappe pour qu'il appartienne à un ID de site distinct.

Les ID de site fonctionnent avec des adresses MAC et des adresses IP spécifiques au site. Les paquets qui sortent de la grappe utilisent une adresse MAC et une adresse IP propres au site, tandis que les paquets reçus par la grappe utilisent une adresse MAC et une adresse IP globales. Cette fonctionnalité empêche les commutateurs d'apprendre la même adresse MAC globale à partir des deux sites sur deux ports différents, ce qui entraîne une oscillation MAC; au lieu de cela, ils connaissent uniquement l'adresse MAC du site. Les adresses MAC et les adresses IP spécifiques au site sont prises en charge pour le mode routé qui utilise uniquement les EtherChannels étendus.

Les ID de site sont également utilisés pour permettre la mobilité de flux à l'aide de l'inspection LISP, la localisation du directeur pour améliorer les performances et réduire la latence aller-retour pour la mise en grappe inter-sites pour les centres de données, et la redondance des sites pour les connexions où un propriétaire de secours d'un flux se trouve toujours sur un autre site que celui du propriétaire.

Consultez les sections suivantes pour en savoir plus sur la mise en grappe inter-sites :

- Dimensionnement de l'interconnexion du centre de données : [Exigences et conditions préalables à la mise en grappe sur Châssis Firepower 4100/9300](#), à la page 8
- Lignes directrices inter-sites : [Lignes directrices et limites de la mise en grappe](#), à la page 11
- Configurer la mobilité des flux de grappes : [Configurer la mobilité des flux de grappes](#), à la page 40

- Activer la localisation du directeur : [Activer la localisation du directeur, à la page 38](#)
- Activer la redondance de site : [Activer la localisation du directeur, à la page 38](#)

Exigences et conditions préalables à la mise en grappe sur Châssis Firepower 4100/9300

Maximum d'unités de mise en grappe par modèle

- Firepower 4100 : 16 châssis
- Firepower 9300 : 16 modules. Par exemple, vous pouvez utiliser 1 module dans 16 châssis, ou 2 modules dans 8 châssis, ou toute combinaison qui fournit un maximum de 16 modules.

Exigences matérielles et logicielles pour la mise en grappe inter-châssis

Tous les châssis d'une grappe :

- : pour Firepower 4100 : tous les châssis doivent être du même modèle. Pour le périphérique Firepower 9300 : tous les modules de sécurité doivent être du même type. Par exemple, si vous utilisez la mise en grappe, tous les modules du périphérique Firepower 9300 doivent être des SM-40. Vous pouvez avoir différentes quantités de modules de sécurité installés dans chaque châssis, bien que tous les modules présents dans le châssis doivent appartenir à la grappe, y compris les logements vides.
- Doit exécuter FXOS et le logiciel d'application identiques, sauf au moment d'une mise à niveau d'image. Des versions logicielles non concordantes peuvent entraîner une dégradation des performances. Assurez-vous donc de mettre à niveau tous les nœuds dans la même fenêtre de maintenance.
- Doit inclure la même configuration d'interface pour les interfaces que vous affectez à la grappe, comme la même interface de gestion, les mêmes EtherChannels, les interfaces actives, la vitesse et le duplex, etc. Vous pouvez utiliser différents types de modules de réseau sur le châssis tant que les capacités correspondent pour les mêmes ID d'interface et que les interfaces peuvent être groupées avec succès dans le même EtherChannel étendu. Notez que toutes les interfaces de données doivent être des EtherChannels dans des grappes à plusieurs châssis. Si vous modifiez les interfaces dans FXOS après avoir activé la mise en grappe (en ajoutant ou en supprimant des modules d'interface, ou en configurant EtherChannels, par exemple), vous effectuez les mêmes modifications sur chaque châssis, en commençant par les nœuds de données jusqu'au nœud de contrôle. Notez que si vous supprimez une interface dans FXOS, la configuration de l'ASA conserve les commandes correspondantes afin que vous puissiez effectuer les ajustements nécessaires; la suppression d'une interface de la configuration peut avoir des effets importants. Vous pouvez supprimer manuellement l'ancienne configuration d'interface.
- Doit utiliser le même serveur NTP. Ne réglez pas l'heure manuellement.
- ASA : chaque châssis FXOS doit être enregistré auprès de l'autorité de licence ou du serveur satellite. Il n'y a pas de frais supplémentaires pour les nœuds de données. Pour la réservation de licences permanentes, vous devez acheter des licences distinctes pour chaque châssis. Pour Firewall Threat Defense, toutes les licences sont gérées par Firewall Management Center.

Exigences du commutateur

- Assurez-vous de terminer la configuration du commutateur et de connecter avec succès tous les canaux EtherChannels du châssis aux commutateurs avant de configurer la mise en grappe sur le Châssis Firepower 4100/9300 .
- Pour les caractéristiques de commutateur prises en charge, consultez [Compatibilité Cisco FXOS](#).

dimensionnement de l'interconnexion du centre de données pour la mise en grappe inter-sites

Vous devez réserver la bande passante sur l'interconnexion des centres de données (DCI) pour un trafic de liaison de commande de grappe équivalent au calcul suivant :

$$\frac{\text{Nombre de membres de grappe par site}}{2} \times \text{liaison de commande de grappe par membre}$$

Si le nombre de membres diffère sur chaque site, utilisez le plus grand nombre pour votre calcul. La bande passante minimale pour l'interface DCI ne doit pas être inférieure à la taille de la liaison de commande de grappe pour un membre.

Par exemple :

- Pour 4 membres sur 2 sites :
 - 4 membres de la grappe au total
 - 2 membres sur chaque site
 - Liaison de commande de grappe de 5 Gbit/s par membre

Bande passante DCI réservée = 5 Gbit/s (2/2 x 5 Gbit/s).

- Pour 6 membres sur 3 sites, la taille augmente :
 - Six membres de la grappe au total
 - 3 membres sur le site 1, 2 membres sur le site 2 et 1 membre sur le site 3
 - Liaison de commande de grappe de 10 Gbit/s par membre

Bande passante DCI réservée = 15 Gbit/s (3/2 x 10 Gbit/s).

- Pour 2 membres sur 2 sites :
 - 2 membres de la grappe au total
 - Un membre sur chaque site
 - Liaison de commande de grappe de 10 Gbit/s par membre

Bande passante DCI réservée = 10 Gbit/s (1/2 x 10 Gbit/s = 5 Gbit/s, mais la bande passante minimale ne doit pas être inférieure à la taille de la liaison de commande de grappe (10 Gbit/s)).

Licences pour la mise en grappe sur le Châssis Firepower 4100/9300

Smart Software Manager standard et sur site

La fonctionnalité de mise en grappe elle-même ne nécessite aucune licence. Pour utiliser le cryptage renforcé et d'autres licences facultatives, chaque Châssis Firepower 4100/9300 doit être enregistré auprès de l'autorité de licence ou du serveur Smart Software Manager standard et sur site. Il n'y a pas de frais supplémentaires pour les unités de données.

La licence de cryptage renforcé est automatiquement activée pour les clients qualifiés lorsque vous appliquez le jeton d'enregistrement. Lors de l'utilisation du jeton, chaque châssis doit avoir la même licence de chiffrement. Pour la licence facultative de cryptage renforcé (3DES/AES) activée dans la configuration ASA, consultez ci-dessous.

Dans la configuration de la licence ASA, vous pouvez uniquement configurer la licence Smart sur l'unité de contrôle. La configuration est répliquée sur les unités de données, mais pour certaines licences, elles n'utilisent pas la configuration; elle reste dans un état mis en mémoire cache et seule l'unité de contrôle demande la licence. Les licences sont agrégées en une licence de grappe unique partagée par les unités de grappe, et cette licence agrégée est également mise en mémoire cache sur les unités de données pour être utilisée si l'une d'elles devient l'unité de contrôle ultérieurement. Chaque type de licence est géré comme suit :

- **Essentials** : seule l'unité de contrôle demande la licence Essentials au serveur et les deux unités peuvent l'utiliser en raison de l'agrégation des licences.
- **Contexte** : seule l'unité de contrôle demande la licence de contexte au serveur. La licence Essentials comprend 10 contextes par défaut et est présente sur tous les membres de la grappe. La valeur de la licence Essentials de chaque unité, plus la valeur de la licence de contexte sur l'unité de contrôle, sont combinées jusqu'à la limite de plateformes dans une licence de grappe agrégée. Par exemple :
 - Vous avez 6 modules Firepower 9300 dans la grappe. La licence Essentials comprend 10 contextes; pour 6 unités, ces licences ajoutent jusqu'à 60 contextes. Vous configurez une licence de 20 contextes supplémentaires sur l'unité de contrôle. Par conséquent, la licence de grappe agrégée comprend 80 contextes. Comme la limite de plateformes pour un module est de 250, la licence combinée autorise un maximum de 250 contextes; les 80 contextes sont dans la limite. Par conséquent, vous pouvez configurer jusqu'à 80 contextes sur l'unité de contrôle; chaque unité de données aura également 80 contextes par la réplication de la configuration.
 - Vous avez 3 unités Firepower 4112 dans la grappe. La licence Essentials comprend 10 contextes; pour 3 unités, ces licences ajoutent jusqu'à 30 contextes. Vous configurez une licence de 250 contextes supplémentaires sur l'unité de contrôle. Par conséquent, la licence de grappe agrégée comprend 280 contextes. Comme la limite de plateformes pour une unité est de 250, la licence combinée autorise un maximum de 250 contextes; les 280 contextes sont au-dessus de la limite. Par conséquent, vous ne pouvez configurer que 250 contextes sur l'unité de contrôle; chaque unité de données aura également 250 contextes par la réplication de la configuration. Dans ce cas, vous devez uniquement configurer la licence de contexte de l'unité de contrôle pour 220 contextes.
- **Opérateur** : requis pour le VPN S2S distribué. Cette licence est un droit par unité, et chaque unité demande sa propre licence au serveur.
- **Cryptage renforcé (3DES)** : pour un déploiement Cisco Smart Software Manager sur site antérieur à la version 2.3.0, si votre compte Smart n'est pas autorisé pour le cryptage renforcé, mais que Cisco a

déterminé que vous êtes autorisé à utiliser le cryptage renforcé, vous pouvez ajouter manuellement une licence de cryptage renforcé à votre compte. Cette licence est un droit par unité, et chaque unité demande sa propre licence au serveur.

Si une nouvelle unité de contrôle est choisie, la nouvelle unité de contrôle continue d'utiliser la licence agrégée. Elle utilise également la configuration de licence mise en mémoire cache pour demander à nouveau la licence de l'unité de contrôle. Lorsque l'ancienne unité de contrôle rejoint la grappe en tant qu'unité de données, elle libère le droit de licence de l'unité de contrôle. Avant que l'unité de données ne libère la licence, la licence de l'unité de contrôle peut être dans un état non conforme s'il n'y a aucune licence disponible dans le compte. La licence conservée est valide pendant 30 jours, mais si elle n'est toujours pas conforme après le délai de grâce, vous ne pourrez pas modifier la configuration des fonctionnalités nécessitant des licences spéciales; l'opération n'est pas affectée par le reste. La nouvelle unité active envoie une demande de renouvellement d'autorisation toutes les 12 heures jusqu'à ce que la licence soit conforme. Vous devez éviter de modifier la configuration jusqu'à ce que les demandes de licences soient complètement traitées. Si une unité quitte la grappe, la configuration de contrôle mise en mémoire cache est supprimée, tandis que les droits par unité sont conservés. En particulier, vous devrez demander à nouveau la licence de contexte sur les unités qui ne sont pas des grappes.

Réservation de licence permanente

Pour la réservation de licences permanentes, vous devez acheter des licences distinctes pour chaque châssis et activer les licences *avant* de configurer la mise en grappe.

Lignes directrices et limites de la mise en grappe

Commutateurs pour la mise en grappe

- Assurez-vous que les commutateurs connectés correspondent aux unités de transfert maximales MTU des interfaces de données et de l'interface de liaison de commande de grappe. Vous devez configurer la MTU de l'interface de la liaison de commande de grappe pour qu'elle soit au moins 100 octets supérieure à la MTU de l'interface de données. Assurez-vous donc de configurer le commutateur de connexion de la liaison de commande de grappe correctement. Étant donné que le trafic de liaison de commande de grappe comprend le transfert de paquets de données, la liaison de commande de grappe doit prendre en charge toute la taille d'un paquet de données plus la surcharge de trafic de grappe. De plus, nous ne recommandons pas de définir la MTU de la liaison de commande de grappe entre 2 561 et 8 362. En raison de la gestion du groupe de blocs, cette taille de MTU n'est pas optimale pour le fonctionnement du système.
- Pour les systèmes Cisco IOS XR, si vous souhaitez définir une MTU autre que celle par défaut, définissez la MTU de l'interface IOS XR sur 14 octets au-dessus de la MTU du périphérique de la grappe. Sinon, les tentatives d'homologation de contiguïté OSPF peuvent échouer, sauf si l'option **mtu-ignore** est utilisée. Notez que la MTU du périphérique de grappe doit correspondre à la MTU *IPv4* d'IOS XR. Cet ajustement n'est pas nécessaire pour les commutateurs Cisco Catalyst et Cisco Nexus.
- Sur le ou les commutateurs pour les interfaces de liaison de commande de grappe, vous pouvez éventuellement activer Spanning Tree PortFast sur les ports de commutateur connectés à l'unité de la grappe pour accélérer le processus de jonction des nouvelles unités.
- Sur le commutateur, nous vous recommandons d'utiliser l'un des algorithmes d'équilibrage de charges EtherChannel suivants : **source-dest-ip** or **src-dst-mixed-ip-port** (reportez-vous à la commande **port-channel load-balance** de Cisco Nexus OS et Cisco IOS-XE). N'utilisez pas de mot clé **vlan** dans l'algorithme d'équilibrage de charge, car cela pourrait entraîner une répartition inégale du trafic vers les

périphériques d'une grappe. *Ne modifiez pas* l'algorithme d'équilibrage de charge par défaut sur le périphérique de la grappe.

- Si vous modifiez l'algorithme d'équilibrage de charge de l'EtherChannel sur le commutateur, l'interface EtherChannel du commutateur arrête temporairement de transférer le trafic et le protocole Spanning Tree redémarre. Il faudra attendre un certain temps avant que le trafic ne redevienne fluide.
- Les commutateurs sur le chemin de la liaison de commande de grappe ne doivent pas vérifier la somme de contrôle L4. Le trafic redirigé sur la liaison de commande de grappe n'a pas une somme de contrôle L4 correcte. Les commutateurs qui vérifient la somme de contrôle L4 pourraient entraîner l'abandon du trafic.
- Le temps d'arrêt du groupage du canal de port ne doit pas dépasser l'intervalle Keepalive configuré.
- Sur les EtherChannels de 2e génération, l'algorithme de distribution de hachage par défaut est adaptatif. Pour éviter le trafic symétrique dans une conception VSS, modifiez l'algorithme de hachage sur le canal de port connecté au périphérique de la grappe à fixe :

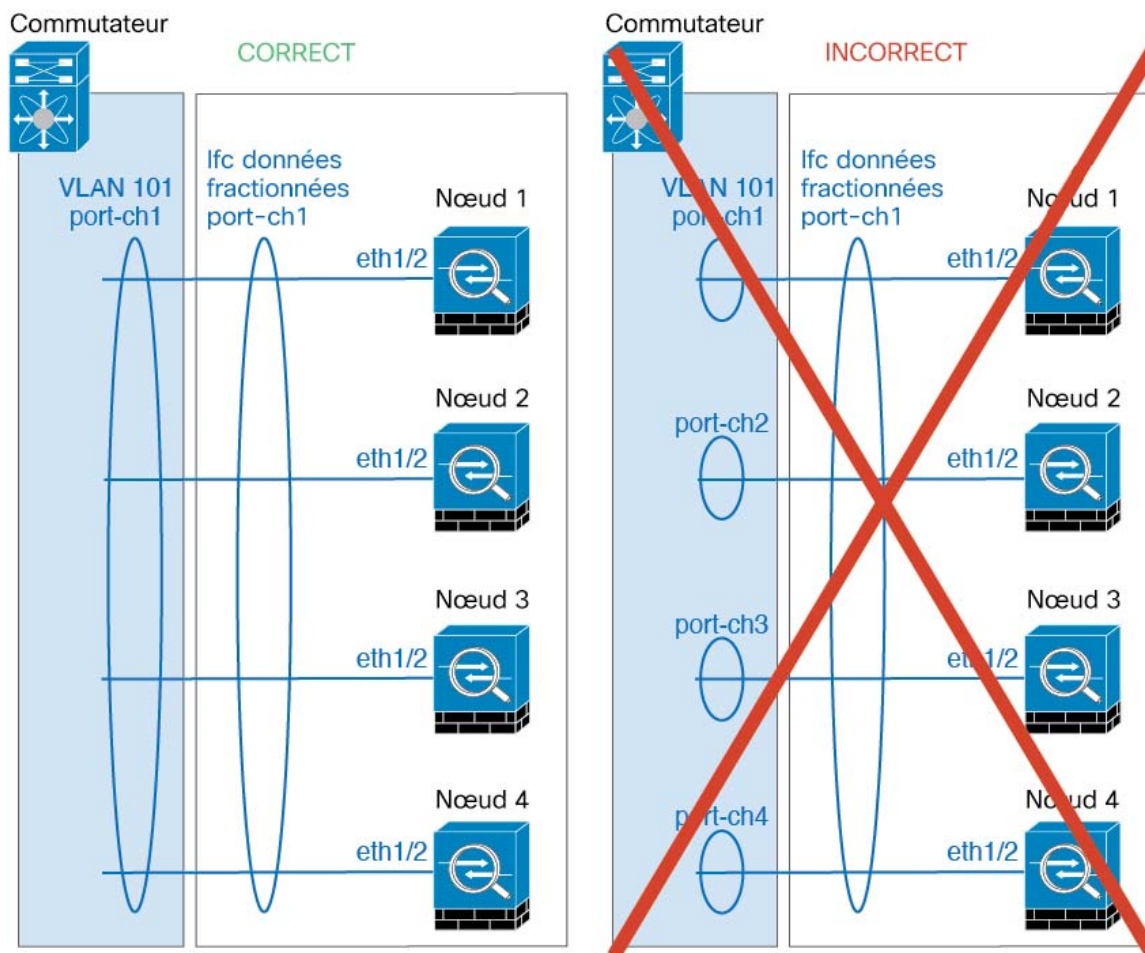
```
router(config)# port-channel id hash-distribution fixed
```

Ne modifiez pas l'algorithme globalement; vous pouvez profiter de l'algorithme adaptatif pour la liaison homologue VSS.

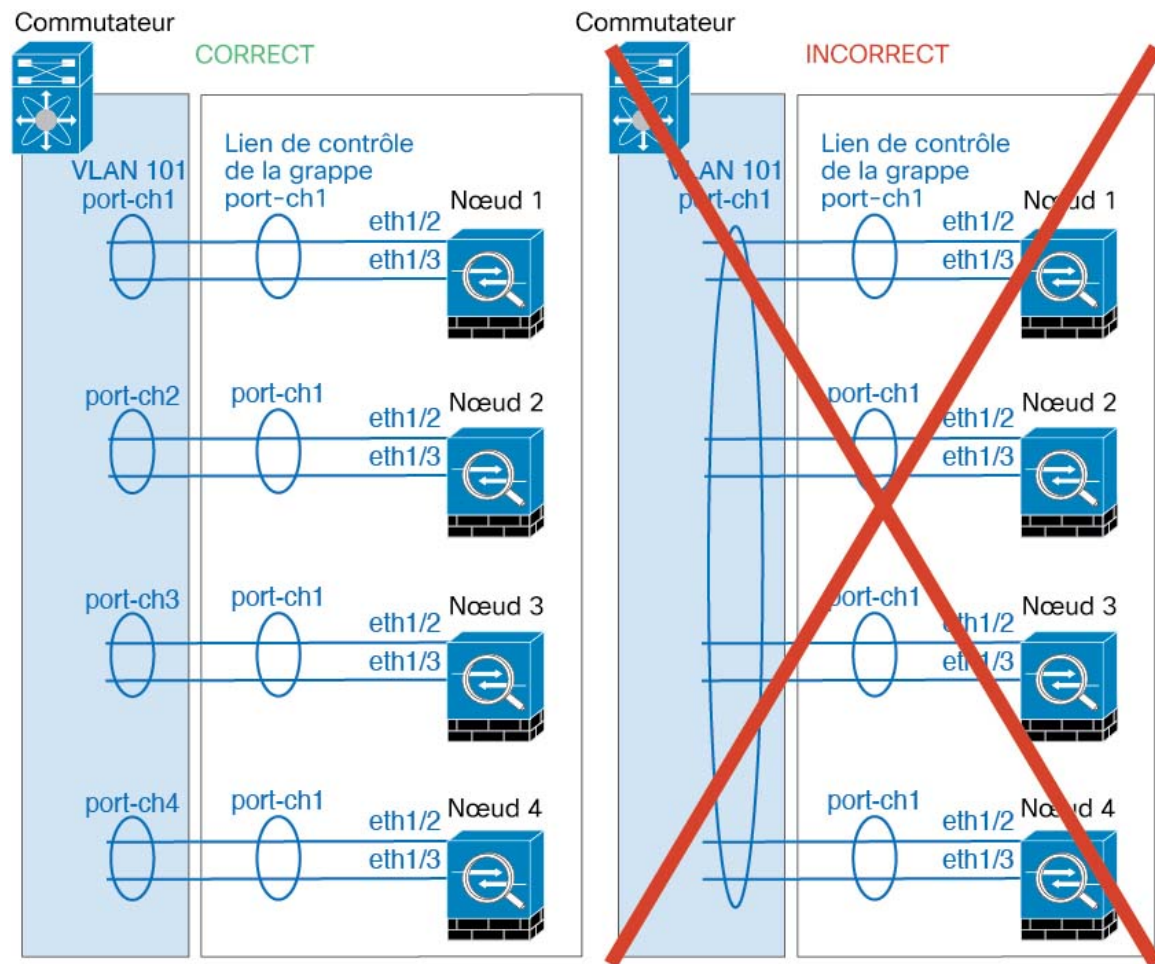
- Contrairement aux grappes de matériel ASA, les grappes Firepower 4100/9300 prennent en charge la convergence progressive LACP. Ainsi, pour la plateforme, vous pouvez laisser la convergence progressive LACP activée sur les commutateurs Cisco Nexus connectés.
- Lorsque vous voyez le regroupement lent d'un EtherChannel étendu sur le commutateur, vous pouvez activer un débit LACP rapide pour une interface individuelle sur le commutateur. Le débit du protocole LACP de FXOS EtherChannels est rapide par défaut. Notez que certains commutateurs, comme la série Nexus, ne prennent pas en charge le débit LACP rapide lors des mises à niveau logicielles en service (ISSU). Nous ne recommandons donc pas l'utilisation des ISSU avec la mise en grappe.

EtherChannels pour la mise en grappe

- Dans les versions du logiciel Cisco IOS Catalyst 3750-X antérieures à la 15.1(1)S2, l'unité de grappe ne prenait pas en charge la connexion d'un EtherChannel à une pile de commutateurs. Avec les paramètres par défaut du commutateur, si l'EtherChannel de l'unité de grappe est connecté de manière croisée et si le commutateur de l'unité de contrôle est hors tension, l'EtherChannel connecté au commutateur restant ne s'activera pas. Pour améliorer la compatibilité, définissez la commande **stack-mac persistent timer** sur une valeur suffisamment grande pour prendre en compte le temps de rechargement; par exemple, 8 minutes ou 0 pour indéfini. Vous pouvez également effectuer une mise à niveau vers une version plus stable du logiciel du commutateur, comme par exemple 15.1(1)S2.
- Configuration des EtherChannels étendus vs. locaux au périphérique : veillez à configurer le commutateur de manière appropriée pour les EtherChannels étendus par rapport aux EtherChannels locaux au périphérique.
 - EtherChannels étendus : pour les EtherChannels *étendus* des unités de grappe, qui s'étendent sur tous les membres de la grappe, les interfaces sont combinées en un seul EtherChannel sur le commutateur. Vérifiez que chaque interface se trouve dans le même groupe de canaux sur le commutateur.



- EtherChannels locaux au périphérique : pour les EtherChannels *locaux au périphérique* de grappe, y compris tous les EtherChannels configurés pour la liaison de commande de la grappe, veillez à configurer des EtherChannels isolés sur le commutateur; ne combinez pas plusieurs EtherChannels d'unités de grappe en un seul EtherChannel sur le commutateur.



Mise en grappe inter-sites

Consultez les consignes suivantes concernant la mise en grappe inter-sites :

- La latence de la liaison de commande de grappe doit être inférieure à 20 ms temps aller-retour (RTT).
- La liaison de commande de grappe doit être fiable, sans paquets dans le désordre ou abandonnés; par exemple, vous devez utiliser un lien dédié.
- Ne configurez pas le rééquilibrage de connexion; il ne faut pas que les connexions soient rééquilibrées entre les membres de la grappe sur un site différent.
- Le ASA ne chiffre pas le trafic de données transféré sur la liaison de commande de grappe, car il s'agit d'un lien dédié, même lorsqu'il est utilisé sur une interconnexion de centre de données (DCI). Si vous utilisez la virtualisation du transport par superposition (OTV), ou étendez la liaison de commande de grappe en dehors du domaine administratif local, vous pouvez configurer le chiffrement sur vos routeurs de frontière, comme 802.1AE MacSec sur OTV.
- La mise en œuvre de la grappe ne fait pas de différence entre les membres de plusieurs sites pour les connexions entrantes; par conséquent, les rôles de connexion pour une connexion donnée peuvent s'étendre sur plusieurs sites. Il s'agit du comportement attendu. Toutefois, si vous activez la localisation de directeur,

le rôle de directeur local est toujours choisi sur le même site que le propriétaire de la connexion (selon l'ID du site). En outre, le directeur local choisit un nouveau propriétaire sur le même site si le propriétaire d'origine tombe en panne (Remarque : si le trafic est dissymétrique entre les sites et qu'il y a un trafic continu en provenance du site distant après la défaillance du propriétaire d'origine, alors un nœud du site distant peut devenir le nouveau propriétaire s'il reçoit un paquet de données pendant la fenêtre de réhbergement.).

- Pour la localisation du directeur, les types de trafic suivants ne prennent pas en charge la localisation : le trafic NAT ou PAT; trafic inspecté par SCTP; requête du propriétaire de la fragmentation.
- Dans un déploiement Nord-Sud avec des flux UDP persistants, des boucles de routage peuvent se produire si les nœuds du site propriétaire initial du flux tombent en panne, puis reviennent en ligne, ce qui entraîne une redirection du flux vers le site d'origine. Si le nouveau propriétaire, sur l'autre site, ne dispose pas d'une route vers la destination, il redirigera le flux vers Internet, créant ainsi une boucle. Dans ce cas, utilisez la commande **clear conn** sur le nouveau propriétaire pour forcer la réinitialisation du flux.
- Pour le mode transparent, si la grappe est placée entre une paire de routeurs internes et externes (AKA insertion nord-sud), vous devez vous assurer que les deux routeurs internes partagent une adresse MAC et que les deux routeurs externes partagent une adresse MAC. Lorsqu'un membre de la grappe sur le site 1 transfère une connexion à un membre du site 2, l'adresse MAC de destination est conservée. Le paquet n'atteindra le routeur du site 2 que si l'adresse MAC est la même que celle du routeur du site 1.
- Pour le mode transparent, si la grappe est placée entre les réseaux de données et le routeur de passerelle de chaque site pour le pare-feu entre les réseaux internes (AKA insertion est-ouest), chaque routeur de passerelle doit utiliser un protocole FHRP (First Hop Redundancy Protocol) comme HSRP pour fournir des destinations d'adresses IP et MAC virtuelles identiques sur chaque site. Les VLAN de données sont étendus à l'ensemble des sites à l'aide de la virtualisation du transport par superposition (OTV), ou d'un outil similaire. Vous devez créer des filtres pour empêcher le trafic destiné au routeur de la passerelle local d'être envoyé à l'autre site sur DCI. Si le routeur de la passerelle devient inaccessible sur un site, vous devez supprimer les filtres pour que le trafic puisse atteindre la passerelle de l'autre site.
- Pour le mode transparent, si la grappe est connectée à un routeur HSRP, vous devez ajouter l'adresse MAC HSRP du routeur en tant qu'entrée de tableau d'adresses MAC statiques sur ASA (voir [Ajouter une adresse MAC statique pour les groupes de ponts](#)). Lorsque des routeurs adjacents utilisent HSRP, le trafic destiné à l'adresse IP HSRP sera envoyé à l'adresse MAC HSRP, mais le trafic retour proviendra de l'adresse MAC de l'interface d'un routeur particulier dans la paire HSRP. Par conséquent, le tableau d'adresses MAC ASA n'est généralement mis à jour que lorsque l'entrée du tableau ARP ASA pour l'adresse IP HSRP expire et que le ASA envoie une demande ARP et reçoit une réponse. Étant donné que les entrées du tableau ARP de ASA expirent après 14 400 secondes par défaut, mais que l'entrée du tableau d'adresses MAC expire après 300 secondes par défaut, une entrée d'adresse MAC statique est requise pour éviter les pertes de trafic liées à l'expiration du tableau.
- Pour le mode routé à l'aide de l'EtherChannel étendu, configurez les adresses MAC propres au site. Étendre les VLAN de données à tous les sites à l'aide d'OTV ou d'un outil similaire. Vous devez créer des filtres pour empêcher le trafic destiné à l'adresse MAC globale d'être envoyé à l'autre site par DCI. Si la grappe devient inaccessible sur un site, vous devez supprimer les filtres pour que le trafic puisse atteindre les nœuds de la grappe de l'autre site. Le routage dynamique n'est pas pris en charge lorsqu'une grappe inter-sites sert de routeur de premier saut pour un segment étendu.

Directives supplémentaires

- Lorsque des modifications de topologie surviennent (par exemple, ajout ou suppression d'une interface de données, activation ou désactivation d'une interface sur l'Châssis Firepower 4100/9300 ou le

commutateur, ajout d'un commutateur supplémentaire pour former un VSS, vPC, StackWise, ou StackWise Virtual), vous devez désactiver le contrôle d'intégrité et désactiver la surveillance d'interface pour les interfaces désactivées. Lorsque la modification de la topologie est terminée et que la modification de la configuration est synchronisée avec toutes les unités, vous pouvez réactiver le contrôle d'intégrité.

- lors de l'ajout d'une unité à une grappe existante ou lors du rechargement d'une unité, il se produira une perte temporaire et limitée de paquets ou de connexion; c'est un comportement attendu. Dans certains cas, les paquets abandonnés peuvent suspendre les connexions; Par exemple, la suppression d'un paquet FIN/ACK pour une connexion FTP entraînera le blocage du client FTP. Dans ce cas, vous devez rétablir la connexion FTP.
- Si vous utilisez un serveur Windows 2003 connecté à une interface EtherChannel étendue, lorsque le port du serveur syslog est en panne et que le serveur ne limite pas les messages d'erreur ICMP, un grand nombre de messages ICMP sont renvoyés à la grappe. Ces messages peuvent faire en sorte que certaines unités de la grappe connaissent un niveau élevé de CPU, ce qui peut affecter les performances. Nous vous recommandons de limiter les messages d'erreur ICMP.
- Nous vous recommandons de connecter les EtherChannels à un VSS, à un vPC, à StackWise ou à StackWise Virtual pour la redondance.
- Dans un châssis, vous ne pouvez pas mettre en grappe certains modules de sécurité et exécuter d'autres modules de sécurité en mode autonome; vous devez inclure tous les modules de sécurité dans la grappe.

Valeurs par défaut

- La fonction de vérification de l'intégrité de la grappe est activée par défaut avec un délai d'attente de 3 secondes. La surveillance de l'intégrité des interfaces est activée sur toutes les interfaces par défaut.
- Le rééquilibrage de la connexion est désactivé par défaut. Si vous activez le rééquilibrage de la connexion, l'intervalle par défaut entre les échanges d'informations de chargement est de 5 secondes.
- La fonction de jonction automatique de la grappe pour une liaison de commande de grappe défaillante est définie pour permettre un nombre illimité de tentatives toutes les 5 minutes.
- La fonction de jonction automatique de la grappe pour une interface de données défaillante est définie à 3 tentatives toutes les 5 minutes, avec un intervalle croissant défini à 2.
- Un délai de duplication de connexion de 5 secondes est activé par défaut pour le trafic HTTP.

Configurer la mise en grappe sur le Châssis Firepower 4100/9300

Vous pouvez facilement déployer la grappe à partir du superviseur Châssis Firepower 4100/9300 . Toute la configuration initiale est générée automatiquement pour chaque unité. Cette section décrit la configuration du démarrage par défaut et la personnalisation facultative que vous pouvez effectuer sur l'ASA. Cette section décrit également comment gérer les membres de la grappe à partir de l'ASA. Vous pouvez également gérer l'appartenance à la grappe à partir du Châssis Firepower 4100/9300 . Pour obtenir de plus amples renseignements, consultez la documentation du Châssis Firepower 4100/9300 .

Procédure

Étape 1	FXOS : Ajouter une grappe ASA, à la page 17
Étape 2	ASA : Modifier le mode de pare-feu et le mode de contexte, à la page 27
Étape 3	ASA : Configurer les interfaces de données, à la page 28
Étape 4	ASA : Personnaliser la configuration de la grappe, à la page 31
Étape 5	ASA : Gérer les membres de la grappe, à la page 53

FXOS : Ajouter une grappe ASA

Vous pouvez ajouter un seul châssis Firepower 9300 en tant que grappe intra-châssis ou ajouter plusieurs châssis pour la mise en grappe inter-châssis. Pour la mise en grappe inter-châssis, vous devez configurer chaque châssis séparément. Ajoutez la grappe sur un châssis; vous pouvez ensuite saisir la plupart des mêmes paramètres sur le châssis suivant.

Créer une grappe ASA

Définissez la portée sur la version de l'image.

Vous pouvez facilement déployer la grappe à partir du superviseur Châssis Firepower 4100/9300. Toute la configuration initiale est générée automatiquement pour chaque unité.

Pour la mise en grappe sur plusieurs châssis, vous devez configurer chaque châssis séparément. Déployer la grappe sur un châssis; vous pouvez ensuite copier la configuration de démarrage du premier châssis au châssis suivant pour faciliter le déploiement.

Dans un châssis Firepower 9300, vous devez activer la mise en grappe pour les 3 logements de module, même si aucun module n'est installé. Si vous ne configurez pas les 3 modules, la grappe ne s'affichera pas.

Pour le mode à contextes multiples, vous devez d'abord déployer le périphérique logique, puis activer le mode à contextes multiples dans l'application ASA.

Lorsque vous déployez une grappe, le superviseur Châssis Firepower 4100/9300 configure chaque application ASA avec la configuration de démarrage suivante. Vous pouvez modifier ultérieurement, si vous le souhaitez, des parties de la configuration de démarrage à partir de l'ASA (en **gras**).

```
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

```

**Remarque**

Le nom **local-unit** ne peut être modifié que si vous désactivez la mise en grappe.

Avant de commencer

- Téléchargez l'image de l'application que vous voulez utiliser pour l'appareil logique à partir de Cisco.com, puis téléchargez cette image sur le serveur de l'application Châssis Firepower 4100/9300 .
- Recueillez les informations suivantes :
 - ID de l'interface de gestion, adresses IP et masque de réseau
 - Adresse IP de la passerelle

Procédure**Étape 1**

Configurer les interfaces.

- Ajoutez au moins une interface de type de données ou un EtherChannel (également appelé canal de port) avant de déployer la grappe. Reportez-vous aux sections [Ajouter un canal EtherChannel \(canal de port\)](#) ou [Configurer une interface physique](#).

Pour la mise en grappe sur plusieurs châssis, toutes les interfaces de données doivent être des EtherChannels étendus avec au moins une interface membre. Ajoutez les mêmes EtherChannels sur chaque châssis. Combinez les interfaces membres de toutes les unités de la grappe en un seul EtherChannel sur le commutateur. Consultez [Lignes directrices et limites de la mise en grappe, à la page 11](#) pour obtenir des renseignements sur les EtherChannels.

- Ajoutez une interface de type de gestion ou un EtherChannel. Reportez-vous aux sections [Ajouter un canal EtherChannel \(canal de port\)](#) ou [Configurer une interface physique](#).

L'interface de gestion est requise. Notez que cette interface de gestion n'est pas la même que l'interface de gestion du châssis qui est utilisée uniquement pour la gestion de ce dernier (dans FXOS, vous pouvez voir l'interface de gestion du châssis affichée comme MGMT, management0, ou d'autres noms similaires).

Pour la mise en grappe sur plusieurs châssis, ajoutez la même interface de gestion sur chaque châssis.

- Pour la mise en grappe sur plusieurs châssis, ajoutez une interface membre à l'EtherChannel de la liaison de commande de grappe (par défaut, le canal de port 48). Consultez [Ajouter un canal EtherChannel \(canal de port\)](#).

N'ajoutez pas d'interface membre pour une grappe isolée aux modules de sécurité dans un châssis Firepower 9300. Si vous ajoutez un membre, le châssis suppose que cette grappe utilisera plusieurs châssis et vous permettra uniquement d'utiliser des EtherChannels étendus, par exemple.

Ajoutez les mêmes interfaces membre sur chaque châssis. La liaison de commande de grappe est un EtherChannel local au périphérique sur chaque châssis. Utilisez des EtherChannels distincts sur le commutateur pour chaque périphérique. Consultez [Lignes directrices et limites de la mise en grappe](#), à la page 11 pour obtenir des renseignements sur les EtherChannels.

Étape 2 Entrez le mode de services de sécurité.

scope ssa

Exemple :

```
Firepower# scope ssa
Firepower /ssa #
```

Étape 3 Définissez les paramètres d'instance de l'application, y compris la version de l'image.

a) Affichez les images disponibles. Notez le numéro de version que vous souhaitez utiliser.

show app

Exemple :

```
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) Définissez la portée sur la version de l'image.

scope app asa application_version

Exemple :

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

c) Définissez cette version comme étant la version par défaut.

set-default

Exemple :

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

d) Quittez le mode ssa.

exit

Exemple :

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

Exemple :

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

Étape 4

Créer la grappe.

enter logical-device device_name asa slots clustered

- *device_name* : utilisé par le superviseur Châssis Firepower 4100/9300 pour configurer les paramètres de mise en grappe et affecter des interfaces; il ne s'agit pas du nom de grappe utilisé dans la configuration du module de sécurité. Vous devez préciser les trois modules de sécurité, même si vous n'avez pas encore installé le matériel.
- *slots (logements)* : attribue les modules de châssis à la grappe. Pour Firepower 4100, spécifiez **1**. Pour Firepower 9300, spécifiez **1,2,3**. Vous devez activer la mise en grappe pour les 3 logements de module dans un châssis Firepower 9300, même si aucun module n'est installé. Si vous ne configurez pas les 3 modules, la grappe ne s'affichera pas.

Exemple :

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

Étape 5

Configurer les paramètres de démarrage de la grappe.

Ces paramètres sont destinés uniquement au déploiement initial ou à la reprise après sinistre. Pour un fonctionnement normal, vous pouvez ultérieurement modifier la plupart des valeurs dans la configuration de l'interface de ligne de commande de l'application.

- a) Créez l'objet de démarrage de grappe.

enter cluster-bootstrap**Exemple :**

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

- b) Définissez l'ID de châssis.

set chassis-id ID

Chaque châssis de la grappe a besoin d'un ID unique.

- c) Pour la mise en grappe inter-sites, définissez l'ID de site entre 1 et 8.

set site-id nombre.

Pour supprimer l'ID de site, définissez la valeur à 0.

Exemple :

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

- d) Configurez une clé d'authentification pour contrôler le trafic dans la liaison de commande de la grappe.

set key

Exemple :

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

Vous êtes invité à saisir le code secret partagé.

Le secret partagé est une chaîne ASCII comptant de 1 à 63 caractères. Le code secret partagé est utilisé pour générer la clé de chiffrement. Cette option n'influe pas sur le trafic datapath, y compris sur la mise à jour de l'état de connexion et les paquets transférés, qui sont toujours envoyés en clair.

- e) Définissez le mode d'interface de grappe.

set mode spanned-etherchannel

Le mode EtherChannel étendu est le seul mode pris en charge.

Exemple :

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

- f) Définissez le nom du groupe de grappes dans la configuration du module de sécurité.

set service-type *nom_groupe*

Le nom doit être une chaîne ASCII comptant de 1 à 38 caractères.

Exemple :

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

- g) (Facultatif) Définissez le réseau IP de la liaison de commande de la grappe.

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

Par défaut, la liaison de commande de grappe utilise le réseau 127.2.0.0/16. Cependant, certains déploiements réseau ne permettent pas le passage du trafic 127.2.0.0/16. Dans ce cas, vous pouvez spécifier une adresse /16 sur un réseau unique pour la grappe.

- ***a.b.0.0*** : spécifiez toute adresse réseau /16, à l'exception des adresses de boucle avec retour (127.0.0.0/8) et de multidiffusion (224.0.0.0/4). Si vous définissez la valeur à 0.0.0.0, le réseau par défaut est utilisé : 127.2.0.0.

Le châssis génère automatiquement l'adresse IP de l'interface de liaison de commande de grappe pour chaque unité en fonction de l'ID du châssis et de l'ID de logement : *a.b.id_châssis.id_logement*.

Exemple :

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

- h) Configurez les informations relatives à l'adresse IP de gestion.

Ces informations sont utilisées pour configurer une interface de gestion dans la configuration du module de sécurité.

1. Configurez un ensemble d'adresses IP locales, dont une sera affectée à chaque unité de grappe pour l'interface.

```
set ipv4 pool start_ip end_ip
```

```
set ipv6 pool start_ip end_ip
```

Incluez au moins autant d'adresses qu'il y a d'unités dans la grappe. Notez que pour le périphérique Firepower 9300, vous devez inclure 3 adresses par châssis, même si tous les logements de module ne sont pas remplis. Si vous prévoyez d'étendre la grappe, incluez des adresses supplémentaires. L'adresse IP virtuelle (appelée adresse IP de la grappe principale) qui appartient à l'unité de contrôle actuelle ne fait *pas* partie de cet ensemble; Assurez-vous de réserver une adresse IP sur le même réseau pour l'adresse IP de la grappe principale. Vous pouvez utiliser des adresses IPv4 et/ou IPv6.

2. Configurez l'adresse IP de l'ensemble d'adresses principal pour l'interface de gestion.

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

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

Cette adresse IP doit être sur le même réseau que les adresses de l'ensemble d'adresses de la grappe, mais ne doit pas faire partie de ce dernier.

3. Saisissez l'adresse de la passerelle du réseau.

```
set ipv4 gateway ip_address
```

```
set ipv6 gateway ip_address
```

Exemple :

```
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) Quittez le mode de démarrage de la grappe.

```
exit
```

Exemple :

```
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/* #

```

Étape 6 Configurer les paramètres de démarrage de gestion.

Ces paramètres sont destinés uniquement au déploiement initial ou à la reprise après sinistre. Pour un fonctionnement normal, vous pouvez ultérieurement modifier la plupart des valeurs dans la configuration de l'interface de ligne de commande de l'application.

- a) Créez l'objet de démarrage de gestion.

enter mgmt-bootstrap asa

Exemple :

```

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

```

- b) Spécifiez le mot de passe d'administrateur et d'activation.

create bootstrap-key-secret PASSWORD

set value

Saisissez une valeur *mot de passe*

Confirmez la valeur : *mot de passe*

exit

Exemple :

L'utilisateur admin ASA et le mot de passe d'activation préconfigurés sont utiles pour la récupération du mot de passe. si vous avez un accès FXOS, vous pouvez réinitialiser le mot de passe de l'utilisateur admin en cas d'oubli.

Exemple :

```

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) Spécifiez le mode de pare-feu, routage ou transparent.

create bootstrap-key FIREWALL_MODE

set value {routed | transparent}

exit

En mode routage, l'appareil est considéré comme un saut de routeur dans le réseau. Chaque interface par laquelle vous souhaitez acheminer le trafic réseau se trouve sur un sous-réseau différent. Un pare-feu transparent, en revanche, est un pare-feu de couche 2 qui agit comme une « présence sur le réseau câblé » ou un « pare-feu furtif », et qui n'est pas considéré comme un saut de routeur vers les appareils connectés.

Le mode pare-feu est uniquement défini lors du déploiement initial. Si vous appliquez à nouveau les paramètres de démarrage, ce paramètre n'est pas utilisé.

Exemple :

```
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) Quittez le mode de démarrage de gestion.

exit

Exemple :

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

Étape 7 Enregistrez la configuration.

commit-buffer

Le châssis déploie le périphérique logique en téléchargeant la version de logiciel spécifiée et en envoyant les paramètres de l'interface de gestion et de configuration du démarrage à l'instance d'application. Vérifiez l'état du déploiement à l'aide de la commande **show app-instance**. L'instance d'application est en cours d'exécution et prête à être utilisée lorsque l'état **Admin State** est **Enabled** (Activé) et que l'état **Oper State** est **Online** (En ligne).

Exemple :

```
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		

Étape 8 Pour ajouter un autre châssis à la grappe, répétez cette procédure, sauf que vous devez configurer un unique **chassis-id** et le bon **site-id**; Sinon, utilisez la même configuration pour les deux châssis.

Vérifiez que la configuration de l'interface est la même sur le nouveau châssis. Vous pouvez exporter et importer la configuration du châssis FXOS pour faciliter ce processus.

Étape 9 Connectez-vous à l'unité de contrôle ASA pour personnaliser votre configuration de mise en grappe.

Exemple

Pour le châssis 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

```

Pour le châssis 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
exit
scope app asa 9.5.2.1
set-default
exit
commit-buffer

```

Ajouter d'autres membres à la grappe

Ajoutez ou remplacez le membre de la grappe ASA.



Remarque

Cette procédure s'applique uniquement à l'ajout ou au remplacement d'un *châssis*; si vous ajoutez ou remplacez un module sur un Firepower 9300 pour lequel la mise en grappe est déjà activée, le module sera ajouté automatiquement.

Avant de commencer

- Assurez-vous que votre grappe existante dispose d'un nombre suffisant d'adresses IP dans l'ensemble d'adresses IP de gestion pour ce nouveau membre. Sinon, vous devez modifier la configuration de démarrage de grappe existante sur chaque châssis avant d'ajouter ce nouveau membre. Cette modification entraîne le redémarrage du périphérique logique.
- La configuration de l'interface doit être la même sur le nouveau châssis. Vous pouvez exporter et importer la configuration du châssis FXOS pour faciliter ce processus.
- Pour le mode de contexte multiple, activez le mode de contexte multiple dans l'application ASA sur le premier membre de grappe; les membres de la grappe héritent automatiquement de la configuration en mode de contexte multiple.

Procédure

Étape 1

cliquez sur **OK**.

Étape 2

Pour ajouter un autre châssis à la grappe, répétez la procédure dans [Créer une grappe ASA, à la page 17](#), sauf que vous devez configurer un **chassis-id** unique et le bon **site-id**; sinon, utilisez la même configuration pour le nouveau châssis.

ASA : Modifier le mode de pare-feu et le mode de contexte

Par défaut, le châssis FXOS déploie une grappe en mode de pare-feu routé et en mode de contexte unique.

- Changer le mode de pare-feu : pour changer le mode après le déploiement, changez le mode sur l'unité de contrôle; le mode est automatiquement modifié sur toutes les unités de données pour qu'elles

correspondent. Consultez [Définir le mode du pare-feu](#). En mode de contexte multiple, vous définissez le mode de pare-feu par contexte. Consultez le

- Passer en mode de contexte multiple : pour passer en mode de contexte multiple après le déploiement, modifiez le mode sur l'unité de contrôle; le mode est automatiquement modifié sur toutes les unités de données pour qu'elles correspondent. Voir [Activer le mode de contexte multiple](#). Consultez le

ASA : Configurer les interfaces de données

Cette procédure configure les paramètres de base pour chaque interface de données que vous avez affectée à la grappe lorsque vous l'avez déployée dans FXOS. Pour la mise en grappe sur plusieurs châssis, les interfaces de données sont toujours des interfaces EtherChannel étendus.



Remarque

L'interface de gestion a été préconfigurée lorsque vous avez déployé la grappe. Vous pouvez également modifier les paramètres de l'interface de gestion dans l'ASA, mais cette procédure se concentre sur les interfaces de données. L'interface de gestion est une interface individuelle spéciale, par opposition à une interface étendue. Consultez [Interface de gestion, à la page 5](#) pour de plus amples renseignements.

Avant de commencer

- Pour le mode de contexte multiple, commencez cette procédure dans l'espace d'exécution du système. Si vous n'êtes pas déjà en mode de configuration système, saisissez la commande **changeto system**.
- Pour le mode transparent, configurez le groupe de ponts. Consultez [Configurer la BVI \(Bridge Virtual Interface\)](#).
- Lorsque vous utilisez des EtherChannels étendus pour une grappe avec plusieurs châssis, l'interface du canal de port ne s'affiche pas tant que la mise en grappe n'est pas complètement activée. Cette exigence empêche le trafic d'être transféré vers un nœud qui n'est pas un nœud actif dans la grappe.

Procédure

Étape 1

Spécifiez l'ID de l'interface.

interface ID

Reportez-vous au châssis FXOS pour connaître les interfaces attribuées à cette grappe. L'ID d'interface peut être :

- **canal de port** *entier*
- **ethernet** *logement/port*

Exemple :

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

Étape 2

Activez l'interface :

no shutdown

Étape 3 (Facultatif) Si vous créez des sous-interfaces VLAN sur cette interface, faites-le maintenant.

Exemple :

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

Le reste de cette procédure s'applique aux sous-interfaces.

Étape 4 (Mode de contexte multiple) Affectez l'interface à un contexte, puis changez de contexte et passez en mode d'interface.

Exemple :

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

Pour le mode de contexte multiple, le reste de la configuration d'interface se produit dans chaque contexte.

Étape 5 Nommez l'interface :

nameif name**Exemple :**

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

Le *nom* est une chaîne de texte de 48 caractères maximum et n'est pas sensible à la casse. Vous pouvez modifier le nom en saisissant à nouveau cette commande avec une nouvelle valeur.

Étape 6 Effectuez l'une des opérations suivantes, selon le mode de pare-feu.

- Mode routé : définissez l'adresse IPv4 et/ou IPv6 :

(IPv4)

ip address ip_address [mask]

(IPv6)

ipv6 address ipv6-prefix/prefix-length

Exemple :

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

Les configurations automatiques DHCP, PPPoE et IPv6 ne sont pas prises en charge. Pour les connexions point à point, vous pouvez spécifier un filtre d'adresse locale de 31 bits (255.255.255.254). Dans ce cas, aucune adresse IP n'est réservée pour les adresses de réseau ou de diffusion. La configuration manuelle de l'adresse Link-Local n'est pas non plus prise en charge.

- Mode transparent : affectez l'interface à un groupe de ponts :

bridge-group *number*

Exemple :

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

Où *le nombre* est un entier compris entre 1 et 100. Vous pouvez attribuer jusqu'à 64 interfaces à un groupe de ponts. Vous ne pouvez pas attribuer la même interface à plusieurs groupes de ponts. Notez que la configuration des BVI comprend l'adresse IP.

Étape 7 Définissez le niveau de sécurité :

security-level *number*

Exemple :

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

Où *number* est un entier compris entre 0 (le plus bas) et 100 (le plus élevé).

Étape 8 (Mise en grappe sur plusieurs châssis) Configurez une adresse MAC globale unique et manuelle pour un EtherChannel étendu afin d'éviter d'éventuels problèmes de connectivité au réseau.

mac-address *mac_address*

- *mac_address* : l'adresse MAC est au format H.H.H., où H est une valeur hexadécimale de 16 bits. Par exemple, l'adresse MAC 00-0C-F1-42-4C-DE est saisie comme suit : 000C.F142.4CDE. Les deux premiers octets d'une adresse MAC manuelle ne peuvent pas être A2 si vous souhaitez également utiliser des adresses MAC générées automatiquement.

Vous devez configurer une adresse MAC unique qui n'est pas utilisée actuellement sur votre réseau. Dans le cas d'une adresse MAC configurée manuellement, l'adresse MAC reste celle de l'unité de contrôle actuelle. Si vous ne configurez pas d'adresse MAC, si l'unité de contrôle change, la nouvelle unité de contrôle utilisera une nouvelle adresse MAC pour l'interface, ce qui peut provoquer une panne temporaire du réseau.

En mode de contexte multiple, si vous partagez une interface entre les contextes, vous devez plutôt activer la génération automatique des adresses MAC pour ne pas avoir à définir l'adresse MAC manuellement. Notez que vous devez configurer manuellement l'adresse MAC à l'aide de cette commande pour les interfaces *non partagées*.

Exemple :

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

Étape 9 (Mise en grappe inter-sites) Configurez une adresse MAC spécifique au site et, pour le mode routé, une adresse IP pour chaque site :

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

Exemple :

```
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.ddd site-id 4 site-ip 10.9.9.4
```

Les adresses IP spécifiques au site doivent se trouver sur le même sous-réseau que l'adresse IP globale. L'adresse MAC spécifique au site et l'adresse IP utilisées par une unité dépendent de l'ID de site que vous spécifiez dans la configuration de démarrage de chaque unité.

ASA : Personnaliser la configuration de la grappe

Si vous souhaitez modifier les paramètres de démarrage après avoir déployé la grappe ou configuré des options supplémentaires, comme la surveillance de l'intégrité de la grappe, le délai de réplication de la connexion TCP, la mobilité des flux et d'autres optimisations, vous pouvez le faire sur l'unité de contrôle.

Configurer les paramètres de grappe ASA de base

Vous pouvez personnaliser les paramètres de grappe sur le nœud de contrôle.

Avant de commencer

- Pour le mode de contexte multiple, effectuez cette procédure dans l'espace d'exécution du système sur l'unité de contrôle. Pour passer du contexte à l'espace d'exécution du système, saisissez la commande **changeto system**.
- Le nom de l'unité locale et plusieurs autres options ne peuvent être définis que sur le châssis FXOS, ou ils ne peuvent être modifiés que sur l'ASA si vous désactivez la mise en grappe. Ils ne sont donc pas inclus dans la procédure suivante.

Procédure

Étape 1 Confirmez qu'il s'agit de l'unité de contrôle :

show cluster info

Exemple :

```
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

```

Si une autre unité est l'unité de contrôle, quittez la connexion et connectez-vous à l'unité appropriée.

Étape 2 Spécifiez l'unité de transfert maximale de l'interface de liaison de commande de grappe, qui doit être supérieure d'au moins 100 octets à la MTU la plus élevée des interfaces de données.

mtu cluster bytes

Exemple :

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

Nous vous suggérons de définir la MTU de la liaison de commande de grappe au maximum; la valeur minimale est de 1 400 octets. De plus, nous ne conseillons pas de définir la MTU de la liaison de commande de grappe entre 2 561 et 8 362. En raison de la gestion du groupe de blocs, cette taille de MTU n'est pas optimale pour le fonctionnement du système. Étant donné que le trafic de liaison de commande de grappe comprend le transfert de paquets de données, la liaison de commande de grappe doit prendre en charge toute la taille d'un paquet de données plus la surcharge de trafic de grappe. Par exemple, comme la MTU maximale est de 9 184 octets, la MTU de l'interface de données la plus élevée peut s'établir à 9 084 octets, tandis que la liaison de commande de grappe peut être définie sur 9 184 octets.

Étape 3 Entrez en mode de configuration de la grappe :

cluster group name

Étape 4 (Facultatif) Activez la réplication de la console des unités de données vers l'unité de contrôle :

console-replicate

Par défaut, cette fonction est désactivée. L'ASA imprime certains messages directement sur la console pour certains événements critiques. Si vous activez la réplication de console, les unités de données envoient les messages de console à l'unité de contrôle de sorte qu'il suffit de surveiller un seul port de console pour la grappe.

Étape 5 Définissez le niveau de trace minimal pour les événements de mise en grappe :

trace-level level

Définissez le niveau minimum comme souhaité :

- **critical**—Événements critiques (gravité=1)

- **warning**—Avertissements (gravité=2)
- **informational**—Événements informationnels (gravité=3)
- **debug**—Événements de débogage (gravité= 4)

Étape 6

(Facultatif) (Firepower 9300 uniquement) Assurez-vous que les modules de sécurité d'un châssis rejoignent la grappe simultanément afin que le trafic soit réparti uniformément entre les modules. Si un module se joint très avant les autres modules, il peut recevoir plus de trafic que souhaité, car les autres modules ne peuvent pas encore partager la charge.

unit parallel-join num_of_units max-bundle-delay max_delay_time

- **num_of_units** : spécifie le nombre minimum de modules dans le même châssis qui doivent être prêts avant qu'un module puisse rejoindre la grappe (entre 1 et 3). La valeur par défaut est 1, ce qui signifie qu'un module n'attendra pas que les autres modules soient prêts avant de rejoindre la grappe. Par exemple, si vous définissez la valeur à 3, chaque module attendra *max_delay_time* ou jusqu'à ce que les trois modules soient prêts avant de rejoindre la grappe. Les trois modules demanderont de rejoindre la grappe presque simultanément et commenceront tous à recevoir du trafic à peu près au même moment.
- **max_delay_time** : spécifie le délai maximum en minutes avant qu'un module arrête d'attendre que d'autres modules soient prêts avant de rejoindre la grappe, entre 0 et 30 minutes. La valeur par défaut est 0, ce qui signifie que le module n'attendra pas que les autres modules soient prêts avant de rejoindre la grappe. Si vous définissez le *num_of_units* sur 1, cette valeur doit être de 0. Si vous définissez le *num_of_units* sur 2 ou 3, cette valeur doit être de 1 ou plus. Cette minuterie est définie par module, mais lorsque le premier module rejoint la grappe, toutes les autres minuteries du module se terminent et les modules restants rejoignent la grappe.

Par exemple, vous définissez le *num_of_units* sur 3 et le *max_delay_time* sur 5 minutes. Lorsque le module 1 apparaît, il démarre sa minuterie de 5 minutes. Le module 2 apparaît 2 minutes plus tard et démarre sa minuterie de 5 minutes. Le module 3 apparaît 1 minute plus tard. Par conséquent, tous les modules rejoindront désormais la grappe au bout de 4 minutes; ils n'attendront pas la fin de la minuterie. Si le module 3 n'apparaît jamais, le module 1 rejoindra la grappe à la fin de sa minuterie de 5 minutes et le module 2 rejoindra également la grappe, même s'il reste 2 minutes à s'écouler; il n'attendra pas la fin de sa minuterie.

Étape 7

Configurez le nombre maximum de membres de la grappe.

cluster-member-limit number

- **number** : de 2 à 16. Par défaut, c'est 16.

Si vous savez que votre grappe sera inférieure au maximum de 16 unités, nous vous conseillons de définir le nombre d'unités réel planifié. La définition du nombre maximum d'unités permet à la grappe de mieux gérer les ressources. Par exemple, si vous utilisez la traduction d'adresse de port (PAT), l'unité de contrôle peut ensuite allouer des blocs de ports au nombre de membres planifié sans avoir à réserver des ports pour des unités supplémentaires que vous ne comptez pas utiliser.

Configurer les paramètres de surveillance de l'intégrité et de jonction automatique

Cette procédure permet de configurer la surveillance de l'intégrité des unités et des interfaces.

Vous pouvez désactiver la surveillance de l'intégrité des interfaces non essentielles, par exemple, l'interface de gestion. Vous pouvez superviser n'importe quel ID de canal de port ou ID d'interface physique unique.

La surveillance de l'intégrité n'est pas effectuée sur les sous-interfaces VLAN ou les interfaces virtuelles telles que les VNI ou les BVI. Vous ne pouvez pas configurer la surveillance pour la liaison de commande de grappe; elle est toujours surveillée.

Procédure

Étape 1 Entrez en mode de configuration de la grappe :

cluster group *name*

Étape 2 Personnalisez la fonctionnalité de contrôle de l'intégrité des unités de la grappe :

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

La commande **holdtime** détermine l'intervalle de temps entre les messages d'état heartbeat des unités, entre 0,3 et 45 secondes; la valeur par défaut est de 3 secondes.

Pour déterminer l'intégrité des unités, les unités de grappe ASA envoient à d'autres unités des messages de signal de présence heartbeat sur la liaison de commande de grappe. Si une unité ne reçoit aucun message heartbeat d'une unité homologue pendant la période d'attente, l'unité homologue est considérée comme ne répondant pas ou morte.

Lorsque des modifications de topologie surviennent (par exemple, ajout ou suppression d'une interface de données, activation ou désactivation d'une interface sur l'ASA, le Châssis Firepower 4100/9300 ou le commutateur, ajout d'un commutateur supplémentaire pour former un VSS, vPC, StackWise, ou StackWise Virtual), vous devez désactiver le contrôle d'intégrité et désactiver la surveillance d'interface pour les interfaces désactivées (**no health-check monitor-interface**). Lorsque la modification de la topologie est terminée et que la modification de la configuration est synchronisée avec toutes les unités, vous pouvez réactiver le contrôle d'intégrité.

Exemple :

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

Étape 3 Désactivez le contrôle d'intégrité de l'interface sur une interface :

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

La vérification de l'intégrité de l'interface surveille les défaillances de liaison. Si tous les ports physiques d'une interface logique donnée tombent en panne sur une unité particulière, mais qu'il y a des ports actifs sous la même interface logique sur d'autres unités, l'unité est supprimée de la grappe. Le délai avant que l'ASA ne supprime un membre de la grappe dépend du type d'interface et du fait que l'unité soit un membre établi ou en voie de se joindre à la grappe.

Le contrôle de l'intégrité est activé par défaut pour toutes les interfaces. Vous pouvez le désactiver pour chaque interface en utilisant la forme **no** de cette commande. Vous pouvez désactiver la surveillance de l'intégrité des interfaces non essentielles, par exemple, l'interface de gestion. La surveillance de l'intégrité n'est pas effectuée sur les sous-interfaces VLAN ou les interfaces virtuelles telles que les VNI ou les BVI. Vous ne pouvez pas configurer la surveillance pour la liaison de commande de grappe; elle est toujours surveillée. Spécifiez le **service-application** pour désactiver la surveillance d'une application décorative.

Lorsque des modifications de topologie surviennent (par exemple, ajout ou suppression d'une interface de données, activation ou désactivation d'une interface sur l'ASA, le Châssis Firepower 4100/9300 ou le commutateur, ajout d'un commutateur supplémentaire pour former un VSS, vPC, StackWise, ou StackWise

Virtual), vous devez désactiver le contrôle d'intégrité (**no health-check**) et désactiver la surveillance d'interface pour les interfaces désactivées. Lorsque la modification de la topologie est terminée et que la modification de la configuration est synchronisée avec toutes les unités, vous pouvez réactiver le contrôle d'intégrité.

Exemple :

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

Étape 4

Personnalisez les paramètres de grappe de la jonction automatique après l'échec de la vérification de l'intégrité :

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

- **system** : spécifie les paramètres de jonction automatique pour les erreurs internes. Les défaillances internes comprennent : des états d'applications non uniformes; et ainsi de suite.
- **unlimited** : (Par défaut pour le **cluster-interface**) Ne limite pas le nombre de tentatives de jonction.
- **auto-rejoin-max** : définit le nombre de tentatives de jonction, entre 0 et 65535. **0** désactive la jonction automatique. La valeur par défaut pour les commandes **data-interface** et **system** est de 3.
- **auto_rejoin_interval** : permet de définir la durée de l'intervalle en minutes entre les tentatives de jonction en sélectionnant un intervalle entre 2 et 60. La valeur par défaut est 5 minutes. La durée totale maximale pendant laquelle l'unité tente de rejoindre la grappe est limitée à 14400 minutes (10 jours) à compter du dernier échec.
- **auto_rejoin_interval_variation** : définit si la durée de l'intervalle augmente. Définissez la valeur entre 1 et 3 : **1** (pas de changement); **2** (2 x la durée précédente) ou **3** (3 x la durée précédente). Par exemple, si vous définissez la durée de l'intervalle à 5 minutes et la variation à 2, la première tentative survient après 5 minutes; la deuxième tentative, après 10 minutes (2 x 5); la troisième tentative, après 20 minutes (2 x 10), etc. La valeur par défaut est **1** pour l'interface de grappe et **2** pour l'interface de données et le système

Exemple :

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

Étape 5

Configurez le délai antirebond avant que l'ASA considère une interface comme défaillante et que l'unité ne soit supprimée de la grappe.

health-check monitor-interface debounce-time ms

Définissez le délai antirebond entre 300 et 9 000 ms. La valeur par défaut est 500ms. Des valeurs inférieures permettent une détection plus rapide des défaillances d'interface. Notez que la configuration d'un délai antirebond inférieur augmente les risques de faux positifs. Lorsqu'une mise à jour d'état d'interface se produit, l'ASA attend le nombre de millisecondes spécifié avant de marquer l'interface comme défaillante, et l'unité est supprimée de la grappe. Par exemple, dans le cas d'un EtherChannel qui passe d'un état inactif à un état opérationnel (par exemple, le commutateur a rechargé ou le commutateur a activé un EtherChannel), un délai de réponse plus long peut empêcher l'interface de sembler être défaillante sur une unité de la grappe juste parce qu'une autre unité de la grappe a été plus rapide à regrouper les ports.

Exemple :

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

Étape 6 Configurez l'intervalle de vérification de l'intégrité du châssis :**app-agent heartbeat** [*interval ms*] [*retry-count number*]

- **interval ms** : définissez la durée entre les pulsations, entre 100 et 6 000 ms, en multiples de 100. La valeur par défaut est 1000ms.
- **retry-count number** : définissez le nombre de tentatives, entre 1 et 30. La valeur par défaut est de 3 tentatives.

L'ASA vérifie s'il peut communiquer via le fond de panier avec le châssis hôte.

La durée minimum combinée (*intervalle x nombre de tentatives*) ne peut pas être inférieure à 600 ms. Par exemple, si vous définissez l'intervalle sur 100 et le nombre de nouvelles tentatives sur 3, la durée totale combinée est de 300 ms, ce qui n'est pas pris en charge. Par exemple, vous pouvez définir l'intervalle sur 100 et le nombre de nouvelles tentatives sur 6 pour respecter la durée minimum (600 ms).

Exemple :

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

Étape 7 (Facultatif) Configurez la surveillance de la charge de trafic.**load-monitor** [*frequency seconds*] [*intervals intervals*]

- **frequency seconds** : définit le délai en secondes entre les messages de surveillance, entre 10 et 360 secondes. La valeur par défaut est de 20 secondes.
- **intervals intervals** : définit le nombre d'intervalles pour lesquels l'ASA conserve les données, entre 1 et 60. La valeur par défaut est 30.

Vous pouvez superviser la charge de trafic pour les membres de la grappe, y compris le nombre total de connexions, l'utilisation du CPU et de la mémoire, ainsi que les abandons de tampon. Si la charge est trop élevée, vous pouvez choisir de désactiver manuellement la mise en grappe sur l'unité si les unités restantes peuvent gérer la charge ou de régler l'équilibrage de charges sur le commutateur externe. Cette fonction est activée par défaut. Par exemple, pour la mise en grappe inter-châssis sur le Firepower 9300 avec 3 modules de sécurité dans chaque châssis, si deux modules de sécurité d'un châssis quittent la grappe, la même quantité de trafic vers le châssis sera envoyée au module restant et risque de le submerger. Vous pouvez superviser périodiquement la charge de trafic. Si la charge est trop élevée, vous pouvez choisir de désactiver manuellement la mise en grappe sur l'unité.

Utilisez la commande **show cluster info load-monitor** pour afficher la charge de trafic.

Exemple :

```
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

Configurer la rééquilibrage de la connexion et le délai de réplication TCP de la grappe

Vous pouvez configurer le rééquilibrage des connexions. Vous pouvez activer le délai de réplication de grappe pour les connexions TCP pour aider à éliminer le « travail inutile » lié aux flux de courte durée en retardant la création du flux directeur/de sauvegarde. Notez que si une unité tombe en panne avant la création du flux directeur/de sauvegarde, ces flux ne peuvent pas être récupérés. De même, si le trafic est rééquilibré vers une autre unité avant la création du flux, le flux ne peut pas être récupéré. Vous ne devez pas activer la réplication TCP pour le trafic sur lequel vous désactivez la répartition aléatoire TCP.

Procédure

Étape 1 Entrez en mode de configuration de la grappe :

cluster group *name*

Étape 2 (Facultatif) Activez le rééquilibrage des connexions pour le trafic TCP :

conn-rebalance [**frequency** *seconds*]

Exemple :

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

Cette commande est désactivée par défaut. Si cette option est activée, les ASA échangent périodiquement des renseignements sur les connexions par seconde et déchargent les nouvelles connexions des appareils ayant plus de connexions par seconde vers les appareils moins chargés. Les connexions existantes ne sont jamais déplacées. De plus, comme cette commande rééquilibre uniquement en fonction des connexions par seconde, le nombre total de connexions établies sur chaque nœud n'est pas pris en compte, et le nombre total de connexions peut ne pas être égal. La fréquence, comprise entre 1 et 360 secondes, spécifie la fréquence à laquelle les informations de chargement sont échangées. La valeur par défaut est de 5 secondes.

Une fois qu'une connexion est déchargée sur un autre nœud, elle devient une connexion asymétrique.

Ne configurez pas le rééquilibrage de connexion pour les topologies inter-sites; il ne faut pas que les nouvelles connexions soient rééquilibrées entre les membres de la grappe sur un site différent.

Étape 3 Activez le délai de réplication de la grappe pour les connexions TCP :

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*]}

Exemple :

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

Définissez les *secondes* entre 1 et 15. Le délai **http** est activé par défaut pendant 5 secondes.

Configurer les fonctionnalités inter-sites

Pour la mise en grappe inter-sites, vous pouvez personnaliser votre configuration pour améliorer la redondance et la stabilité.

Activer la localisation du directeur

Pour améliorer les performances et réduire la latence d'aller-retour dans le cas de la mise en grappe inter-sites pour les centres de données, vous pouvez activer la localisation du directeur. Les nouvelles connexions sont généralement équilibrées et appartiennent aux membres de la grappe d'un site donné. Cependant, l'ASA attribue le rôle de directeur à un membre sur *n'importe quel* site. La localisation du directeur permet des rôles de directeur supplémentaires : un directeur local sur le même site que le propriétaire et un directeur global qui peut se trouver sur n'importe quel site. Le fait de maintenir le propriétaire et le directeur sur le même site améliore les performances. En cas de défaillance du propriétaire d'origine, le directeur local choisit un nouveau propriétaire de connexion sur le même site. Le directeur global est utilisé si un membre d'une grappe reçoit des paquets pour une connexion appartenant à un autre site.

Avant de commencer

- Définissez l'ID de site du châssis sur le superviseur Châssis Firepower 4100/9300 .
- Les types de trafic suivants ne prennent pas en charge la localisation : trafic NAT ou PAT; trafic inspecté par SCTP; requête du propriétaire de la fragmentation.

Procédure

Étape 1 Entrez en mode de configuration de la grappe :

cluster group name

Exemple :

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

Étape 2 Activez la localisation du directeur :

director-localization

Activer la redondance de site

Pour protéger les flux contre une défaillance de site, vous pouvez activer la redondance de site. Si le propriétaire de connexion de secours se trouve sur le même site que le propriétaire, un propriétaire de secours supplémentaire sera choisi dans un autre site pour protéger les flux d'une défaillance du site.

Avant de commencer

- Définissez l'ID de site du châssis sur le superviseur Châssis Firepower 4100/9300 .

Procédure

Étape 1 Entrez en mode de configuration de la grappe :

cluster group *name*

Exemple :

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

Étape 2 Activez la redondance de site.

site-redundancy

Configurer l'ARP gratuit par site

L'ASA génère des paquets ARP gratuits (GARP) pour maintenir l'infrastructure de commutation à jour : le membre ayant la priorité la plus élevée sur chaque site génère périodiquement du trafic GARP pour les adresses MAC et IP globales.

Lorsque vous utilisez des adresses MAC et IP par site, les paquets provenant de la grappe utilisent une adresse MAC et une adresse IP propres au site, tandis que les paquets reçus par la grappe utilisent une adresse MAC et une adresse IP globales. Si le trafic n'est pas généré périodiquement à partir de l'adresse MAC globale, vous pourriez rencontrer un délai d'expiration de l'adresse MAC sur vos commutateurs pour l'adresse MAC globale. Après un délai d'expiration, le trafic destiné à l'adresse MAC globale sera inondé dans l'ensemble de l'infrastructure de commutation, ce qui peut entraîner des problèmes de performance et de sécurité.

Le GARP est activé par défaut lorsque vous définissez l'ID de site pour chaque unité et l'adresse MAC du site pour chaque EtherChannel étendu. Vous pouvez personnaliser l'intervalle GARP ou vous pouvez désactiver GARP.

Avant de commencer

- Définissez l'ID de site pour le membre de la grappe dans la configuration de démarrage.
- Définissez l'adresse MAC par site pour l'EtherChannel étendu dans la configuration de l'unité de contrôle.

Procédure

Étape 1 Entrez en mode de configuration de la grappe.

cluster group *name*

Exemple :

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

```
ciscoasa(cfg-cluster)#
```

Étape 2 Personnalisez l'intervalle GARP.

site-periodic-garp interval seconds

- *seconds* : Définissez le temps en secondes entre la génération du GARP, entre 1 et 1000000 secondes. La valeur par défaut est de 290 secondes.

Pour désactiver GARP, saisissez **no site-periodic-garp interval**.

Exemple :

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

Configurer la mobilité des flux de grappes

Vous pouvez inspecter le trafic LISP pour activer la mobilité des flux lorsqu'un serveur se déplace entre les sites.

À propos de l'inspection LISP

Vous pouvez inspecter le trafic LISP pour activer la mobilité des flux entre les sites.

À propos de LISP

La mobilité des machines virtuelles de centres de données, comme VMware VMotion, permet aux serveurs de migrer entre les centres de données tout en maintenant les connexions avec les clients. Pour prendre en charge une telle mobilité de serveur de centre de données, les routeurs doivent être en mesure de mettre à jour la route d'entrée vers le serveur lorsqu'il se déplace. L'architecture Cisco Locator/Protocole de séparation d'ID (LISP) sépare l'identité du périphérique, ou l'identifiant de terminal (EID), de son emplacement, ou localisateur de routage, dans deux espaces de numérotation différents, ce qui rend la migration du serveur transparente pour les clients. Par exemple, lorsqu'un serveur est déplacé vers un nouveau site et qu'un client envoie du trafic vers le serveur, le routeur redirige le trafic vers le nouvel emplacement.

LISP nécessite des routeurs et des serveurs dans certains rôles, tels que le routeur de tunnel de sortie (ETR) LISP, le routeur de tunnel d'entrée (ITR), les routeurs de premier saut, le résolveur de carte (MR) et le serveur de carte (MS). Lorsque le routeur du premier saut du serveur détecte que le serveur est connecté à un autre routeur, il met à jour tous les autres routeurs et les bases de données afin que l'ITR connectée au client puisse intercepter, encapsuler et envoyer le trafic vers le nouvel emplacement du serveur.

Prise en charge du LISP Cisco Secure Firewall ASA

L'ASA n'exécute pas le LISP lui-même; il peut, cependant, inspecter le trafic LISP pour détecter les changements d'emplacement, puis utiliser ces informations pour un fonctionnement de mise en grappe transparent. Sans l'intégration du LISP, lorsqu'un serveur est déplacé vers un nouveau site, le trafic est acheminé vers un membre de la grappe ASA sur le nouveau site plutôt que vers le propriétaire d'origine du flux. Le nouvel ASA transfère le trafic vers l'ASA au niveau de l'ancien site, puis l'ancien ASA doit renvoyer le trafic vers le nouveau site pour qu'il atteigne le serveur. Ce flux de trafic est sous-optimal et est connu comme « tromboning » ou « épingle à cheveux ».

Grâce à l'intégration du LISP, les membres de la grappe ASA peuvent inspecter le trafic LISP passant entre le routeur du premier saut et l'ETR ou l'ITR, puis peuvent changer le propriétaire du flux pour qu'il se trouve sur le nouveau site.

Lignes directrices LISP

- Les membres de la grappe ASA doivent se trouver entre le routeur du premier saut et l'ITR ou l'ETR du site. La grappe ASA elle-même ne peut pas être le premier routeur de saut pour un segment étendu.
- Seuls les flux entièrement distribués sont pris en charge; les flux centralisés, les flux semi-distribués ou les flux appartenant à des nœuds individuels ne sont pas transférés vers de nouveaux propriétaires. Les flux semi-distribués comprennent des applications, telles que SIP, où tous les flux enfants appartiennent au même ASA qui possède le flux parent.
- La grappe déplace uniquement les états de flux de couche 3 et 4; certaines données d'application pourraient être perdues.
- Pour les flux de courte durée ou non essentiels pour l'entreprise, le déplacement du propriétaire peut ne pas valoir la peine. Vous pouvez contrôler les types de trafic pris en charge avec cette fonctionnalité lorsque vous configurez la politique d'inspection. Vous devez limiter la mobilité de flux au trafic essentiel.

Mise en œuvre du LISP ASA

Cette fonctionnalité comprend plusieurs configurations interdépendantes (qui sont toutes décrites dans ce chapitre) :

1. (Facultatif) Limiter les EID inspectés en fonction de l'adresse IP de l'hôte ou du serveur : le routeur de premier saut pourrait envoyer des messages de notification d'EID pour les hôtes ou les réseaux dans lesquels la grappe ASA n'est pas impliquée, vous pouvez donc limiter les EID à ces serveurs ou réseaux pertinents pour votre grappe. Par exemple, si la grappe ne concerne que 2 sites, mais que le LISP est exécuté sur 3 sites, vous ne devez inclure les EID que pour les 2 sites de la grappe.
2. Inspection du trafic LISP : l'ASA inspecte le trafic LISP sur le port UDP 4342 pour détecter le message EID-notify envoyé entre le premier saut du routeur et l'ITR ou l'ETR. L'ASA gère un tableau EID qui met en corrélation l'EID et l'ID de site. Par exemple, vous devez inspecter le trafic LISP avec une adresse IP source du routeur du premier saut et une adresse de destination de l'ITR ou de l'ETR. Remarque : le trafic LISP n'a pas de directeur affecté et le trafic LISP en lui-même ne participe pas au partage de l'état de la grappe.
3. Politique de service pour activer la mobilité des flux sur le trafic spécifié : vous devez activer la mobilité des flux sur le trafic critique. Par exemple, vous pouvez limiter la mobilité de flux au trafic HTTPS uniquement et/ou au trafic vers des serveurs spécifiques.
4. ID de sites : l'ASA utilise l'ID de site pour chaque nœud de la grappe afin de déterminer le nouveau propriétaire.
5. Configuration au niveau de la grappe pour activer la mobilité de flux : vous devez également activer la mobilité de flux au niveau de la grappe. Ce bouton bascule marche/arrêt vous permet d'activer ou de désactiver facilement la mobilité de flux pour une classe particulière de trafic ou d'applications.

Configurer l'inspection LISP

Vous pouvez inspecter le trafic LISP pour activer la mobilité des flux lorsqu'un serveur se déplace entre les sites.

Avant de commencer

- Définissez l'ID de site du châssis sur le superviseur Châssis Firepower 4100/9300 .

- Le trafic LISP n'est pas inclus dans la classe de trafic d'inspection par défaut. Vous devez donc configurer une classe distincte pour le trafic LISP dans le cadre de cette procédure.

Procédure

Étape 1

(Facultatif) Configurez une carte d'inspection LISP pour limiter les EID inspectés en fonction de l'adresse IP et pour configurer la clé partagée LISP :

- Créez une ACL étendue; seule l'adresse IP de destination correspond à l'adresse intégrée de l'EID :

access list *eid_acl_name* **extended permit ip** *source_address mask destination_address mask*

Les ACL IPv4 et IPv6 sont acceptées. Consultez la référence de commande pour connaître la syntaxe **access-list extended** exacte.

- Créez la carte d'inspection LISP et entrez le mode des paramètres :

policy-map type inspect lisp *inspect_map_name*

parameters

- Définissez les EID autorisés en identifiant l'ACL que vous avez créée :

allowed-eid access-list *eid_acl_name*

Le routeur de premier saut ou l'ITR/ETR pourrait envoyer des messages de notification d'EID pour les hôtes ou les réseaux dans lesquels la grappe ASA n'est pas impliquée; vous pouvez donc limiter les EID à ces serveurs ou réseaux pertinents pour votre grappe. Par exemple, si la grappe ne concerne que 2 sites, mais que le LISP est exécuté sur 3 sites, vous ne devez inclure les EID que pour les 2 sites de la grappe.

- Si nécessaire, saisissez la clé partagée :

validate-key *key*

Exemple :

```
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
```

Étape 2

Configurez l'inspection LISP pour le trafic UDP entre le routeur de premier saut et l'ITR ou l'ETR sur le port 4342 :

- Configurez l'ACL étendue pour identifier le trafic LISP :

access list *inspect_acl_name* **extended permit udp** *source_address mask destination_address mask eq 4342*

Vous devez spécifier le port UDP 4342. Les ACL IPv4 et IPv6 sont acceptées. Consultez la référence de commande pour connaître la syntaxe **access-list extended** exacte.

- Créez une carte de trafic pour l'ACL :

class-map *inspect_class_name*

match access-list *inspect_acl_name*

- c) Spécifiez la liste des politiques, la carte de trafic, activez l'inspection à l'aide de la carte d'inspection LISP facultative et appliquez la politique de service à une interface (si nouvelle) :

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

Si vous avez une politique de service existante, spécifiez le nom de la liste des politiques existante. Par défaut, l'ASA inclut une politique globale appelée **global_policy**, donc pour une politique globale, spécifiez ce nom. Vous pouvez également créer une politique de service par interface si vous ne souhaitez pas appliquer la politique à l'échelle mondiale. L'inspection LISP est appliquée au trafic de manière bidirectionnelle, vous n'avez donc pas besoin d'appliquer la politique de service sur les interfaces source et de destination; tout le trafic qui entre ou sort de l'interface à laquelle vous appliquez la liste des politiques est affecté si le trafic correspond à la carte de trafic dans les deux sens.

Exemple :

```
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
```

L'ASA inspecte le trafic LISP pour détecter le message de notification d'EID envoyé entre le premier saut du routeur et l'ITR ou l'ETR. L'ASA gère un tableau EID qui met en corrélation l'EID et l'ID de site.

Étape 3

Activez la mobilité des flux pour une classe de trafic :

- a) Configurez l'ACL étendue pour identifier le trafic opérationnel critique que vous souhaitez réaffecter au site le plus optimal lorsque les serveurs changent de site :

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

Les ACL IPv4 et IPv6 sont acceptées. Consultez la référence de commande pour connaître la syntaxe **access-list extended** exacte. Vous devez activer la mobilité des flux sur le trafic critique. Par exemple, vous pouvez limiter la mobilité de flux au trafic HTTPS uniquement et/ou au trafic vers des serveurs spécifiques.

- b) Créez une carte de trafic pour l'ACL :

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

- c) Spécifiez la même liste des politiques sur laquelle vous avez activé l'inspection LISP, la carte de trafic de flux, et activez la mobilité des flux :

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

Exemple :

```
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
```

Étape 4

Entrez dans le mode de configuration du groupe de grappes et activez la mobilité des flux pour la grappe :

cluster group name

flow-mobility lisp

Ce bouton bascule marche/arrêt vous permet d'activer ou de désactiver facilement la mobilité des flux.

Exemples

L'exemple suivant :

- Limite les EID à ceux sur le réseau 10.10.10.0/24
- Inspecte le trafic LISP (UDP 4342) entre un routeur LISP à l'adresse 192.168.50.89 (à l'intérieur) et un routeur ITR ou ETTR (sur une autre interface ASA) à l'adresse 192.168.10.8
- Active la mobilité des flux pour tout le trafic interne vers un serveur sur 10.10.10.0/24 à l'aide de HTTPS.
- Active la mobilité des flux pour la grappe.

```
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
```

Configurer le VPN de site à site distribué

Par défaut, la grappe utilise le mode VPN de site à site centralisé. Pour profiter de l'évolutivité de la mise en grappe, vous pouvez activer le mode VPN de site à site distribué.

À propos du VPN de site à site distribué

En mode distribué, les connexions VPN IPsec IKEv2 de site à site sont réparties sur les nœuds d'une grappe. La distribution des connexions VPN sur les nœuds d'une grappe permet d'utiliser pleinement la capacité et le débit de la grappe, ce qui augmente considérablement la prise en charge des VPN au-delà des capacités VPN centralisées.

Rôles de connexion du VPN distribué

Les connexions peuvent être équilibrées vers plusieurs nœuds de la grappe. Les rôles de connexion déterminent la façon dont les connexions sont gérées, à la fois en fonctionnement normal et en situation de disponibilité élevée. Lors de l'exécution en mode VPN distribué, les rôles suivants sont affectés aux nœuds de la grappe :

- **Propriétaire de session active** : le nœud qui reçoit initialement la connexion ou qui a effectué la transition d'une session de sauvegarde en une session active. Le propriétaire conserve l'état et traite les paquets pour toute la session, y compris les tunnels IKE et IPsec et tout le trafic qui leur est associé.
- **Propriétaire de la session de sauvegarde** : le nœud qui gère la session de sauvegarde pour une session active existante. En cas d'échec du propriétaire de la session active, le propriétaire de la session de secours devient le propriétaire de la session active et une nouvelle session de sauvegarde est établie sur un autre nœud.
- **Transmetteur** : si le trafic associé à une session VPN est envoyé à un nœud qui n'est pas propriétaire de la session VPN, ce nœud utilisera la liaison de commande de grappe pour transmettre le trafic au nœud propriétaire de la session VPN.
- **Orchestrateur** : l'orchestrateur (toujours le nœud de contrôle de la grappe) est responsable du calcul des sessions qui se déplaceront et de la destination lors de l'exécution d'une redistribution de sessions actives (ASR). Il envoie une demande au nœud propriétaire X de déplacer N sessions vers le nœud Y. Le nœud X répondra à l'orchestrateur une fois terminé, en précisant le nombre de sessions qu'il a pu déplacer.

Caractéristiques de session du VPN distribué

Les sessions VPN de site à site distribuées présentent les caractéristiques suivantes. Sinon, les connexions VPN se comportent normalement, comme si elles n'étaient pas sur une grappe.

- Les sessions VPN sont réparties dans la grappe au niveau de la session. Cela signifie que le même nœud de grappe gère les tunnels IKE et IPsec et tout leur trafic pour une connexion VPN. Si le trafic de session VPN est envoyé à un nœud de grappe qui ne possède pas cette session VPN, le trafic est transféré au nœud de grappe qui est propriétaire de la session VPN.
- Les sessions VPN ont un ID de session unique dans la grappe. À l'aide de l'ID de session, le trafic est validé, des décisions de transfert sont prises et la négociation IKE est terminée.
- Dans une configuration en étoile avec un VPN de site à site, lorsque les clients se connectent au moyen de la grappe (ce que l'on appelle le hair-pinning), le trafic de session entrant et le trafic de session sortant peuvent se trouver sur des nœuds de grappe différents.
- Vous pouvez exiger que la session de sauvegarde soit allouée sur un module de sécurité dans un autre châssis; cela offre une protection contre les défaillances de châssis. Vous pouvez également choisir

d'allouer des sessions de sauvegarde sur n'importe quel nœud de la grappe; cela offre une protection uniquement contre les défaillances de nœuds. Lorsque la grappe contient deux châssis, la sauvegarde sur le châssis distant est fortement conseillée.

Gestion des événements de grappe par le VPN distribué

Événement	VPN distribué
Défaillance de nœud	Pour toutes les sessions actives sur ce nœud défaillant, les sessions de sauvegarde (sur un autre nœud) deviennent actives et les sessions de sauvegarde sont réaffectées sur un autre nœud en fonction de la stratégie de sauvegarde.
Défaillance du châssis	Lorsqu'une stratégie de sauvegarde de châssis distant est utilisée, pour toutes les sessions actives sur le châssis défaillant, les sessions de sauvegarde (sur un nœud de l'autre châssis) deviennent actives. Lorsque les nœuds seront remplacés, les sessions de sauvegarde pour ces sessions désormais actives seront réaffectées sur les nœuds du châssis remplacé. Lorsqu'une stratégie de sauvegarde plate est utilisée, si les sessions actives et de sauvegarde se trouvent sur le châssis défaillant, la connexion sera abandonnée. Toutes les sessions actives avec des sessions de sauvegarde sur un nœud de l'autre châssis reviennent à ces sessions. De nouvelles sessions de sauvegarde seront attribuées sur un autre nœud du châssis restant.
Désactiver un nœud de grappe	Pour toutes les sessions actives sur le nœud de grappe qui sont désactivées, les sessions de sauvegarde (sur un autre nœud) deviennent actives et réaffectent les sessions de sauvegarde sur un autre nœud en fonction de la stratégie de sauvegarde.
Joindre un nœud de grappe	Si le mode de grappe VPN sur le nouveau nœud n'est pas distribué, le nœud de contrôle demandera un changement de mode. Une fois que le mode VPN est compatible, le nœud de la grappe se verra attribuer des sessions actives et de sauvegarde dans le flux des opérations normales.

Modifications IPsec IKEv2

IKEv2 est modifié en mode VPN de site à site distribué de la manière suivante :

- Une identité est utilisée à la place des tuples IP/port. Cela permet de prendre des décisions de transfert appropriées sur les paquets et de nettoyer les connexions précédentes qui peuvent se trouver sur d'autres membres de la grappe.
- Les identifiants (SPI) qui identifient une seule session IKEv2 sont des valeurs aléatoires de 8 octets générées localement et uniques dans toute la grappe. Un SPI intègre un horodatage et un ID de nœud de grappe. À la réception d'un paquet de négociation IKE, si la vérification de l'horodatage ou de l'ID du nœud de la grappe échoue, le paquet est abandonné et un message est enregistré pour indiquer la raison.
- Le traitement IKEv2 a été modifié pour empêcher les négociations NAT-T d'échouer en étant réparties entre les membres de la grappe. Un nouveau domaine de classification ASP, *cluster_isakmp_redirect*, et des règles sont ajoutés lorsque IKEv2 est activé sur une interface. Utilisez la commande **show asp table classify domain cluster_isakmp_redirect** pour afficher les règles.

Haute disponibilité pour le VPN de site à site distribué dans la grappe

Les fonctionnalités suivantes offrent une résilience contre la défaillance unique d'un module ou d'un châssis de sécurité :

- Les sessions VPN qui sont sauvegardées sur un autre module de sécurité de la grappe, sur n'importe quel châssis, supportent les défaillances de module de sécurité.
- Les sessions VPN sauvegardées sur un autre châssis supportent les défaillances de châssis.
- Le nœud de contrôle peut changer sans perdre les sessions VPN de site à site.

Si une défaillance supplémentaire se produit avant que la grappe ne soit stable, les connexions peuvent être perdues si les sessions actives et de secours se trouvent sur les nœuds défaillants.

Toutes les tentatives sont faites pour assurer qu'aucune session ne soit perdue lorsqu'un nœud quitte la grappe de manière progressive, comme la désactivation du mode grappe VPN, le rechargement d'un nœud de grappe et d'autres modifications de châssis anticipées. Pendant ce type d'opérations, les sessions ne seront pas perdues tant que la grappe aura le temps de rétablir les sauvegardes de sessions entre les opérations. Si une sortie progressive est déclenchée sur le dernier nœud de la grappe, elle supprimera rapidement les sessions existantes.

CMPv2

Le certificat d'identification et les paires de clés CMPv2 sont synchronisés sur les nœuds de la grappe. Cependant, seul le nœud de contrôle de la grappe renouvelle automatiquement et resaisit le certificat CMPv2. Le nœud de contrôle synchronise ces nouveaux certificats d'identification et clés avec tous les nœuds de la grappe lors d'un renouvellement. De cette façon, tous les nœuds de la grappe utilisent les certificats CMPv2 pour l'authentification, et tous les nœuds sont capables de prendre le relais en tant que nœud de contrôle.

Licences pour le VPN de site à site distribué

Une licence d'opérateur est requise pour le VPN de site à site distribué, sur chaque membre de la grappe.

Chaque connexion VPN nécessite deux sessions sous licence *Other VPN* (Autre VPN) (la licence *Other VPN* (Autre VPN) fait partie de la licence *Essentials*), une pour la session active et une pour la session de secours. La capacité maximale de sessions VPN de la grappe ne peut pas dépasser la moitié de la capacité sous licence en raison de l'utilisation de deux licences pour chaque session.

Conditions préalables au VPN de site à site distribué

Prise en charge des modèles

- Firepower 9300
- Maximum de 6 modules sur 2 châssis. Vous pouvez installer un nombre différent de modules de sécurité dans chaque châssis, mais nous conseillons une répartition égale.

Maximum de sessions VPN

Chaque module de sécurité prend en charge jusqu'à 6 000 sessions VPN pour un maximum d'environ 36 000 sessions sur 6 nœuds.

Le nombre réel de sessions prises en charge sur un nœud de grappe est déterminé par la capacité de la plateforme, les licences allouées et l'allocation des ressources par contexte. Lorsque l'utilisation est proche de la limite, il peut arriver que la création de sessions échoue, même si la capacité maximum n'a pas été atteinte sur chaque nœud de la grappe. En effet, l'allocation de sessions actives est déterminée par une

commutation externe et l'allocation de sessions de secours est déterminée par un algorithme de grappe interne. Nous invitons les clients à dimensionner leur utilisation en conséquence et à prévoir une marge pour une répartition inégale.

Lignes directrices relatives au VPN de site à site distribué

Mode pare-feu

Le VPN de site à site distribué est pris en charge en mode routé uniquement.

Mode contextuel

Le VPN de site à site distribué fonctionne à la fois en mode de contexte unique et en mode de contexte multiple. Cependant, en mode de contexte multiple, la redistribution des sessions actives est effectuée au niveau du système, et non au niveau du contexte. Cela empêche une session active associée à un contexte de se déplacer vers un membre de grappe qui contient des sessions actives associées à un contexte différent, créant ainsi une charge non pris en charge.

Inspections non prises en charge

Les types d'inspections suivants ne sont pas pris en charge ou sont désactivés en mode VPN de site à site distribué :

- CTIQBE
- DCERPC
- H323, H225 et RAS
- Intercommunication IPsec
- MGCP
- MMP
- NetBIOS
- PPTP
- RADIUS
- RSH
- RTSP
- SCCP (Skinny)
- SunRPC
- TFTP
- WAAS
- WCCP
- XDMCP

Directives supplémentaires

- Seul le VPN de site à site IKEv2 IPsec est pris en charge en mode VPN de site à site distribué. IKEv1 n'est pas pris en charge. Le VPN de site à site IKEv1 est pris en charge en mode VPN centralisé.
- La mise en grappe inter-sites n'est pas prise en charge.
- La PAT dynamique n'est pas disponible en mode VPN de site à site distribué.

Activer le VPN de site à site distribué

Activez le VPN de site à site distribué pour profiter de l'évolutivité de la mise en grappe pour les sessions VPN.

**Remarque**

Le passage du mode VPN entre centralisé et distribué nécessite le rechargement de tous les nœuds de la grappe. La modification du mode de sauvegarde est dynamique et ne mettra pas fin aux sessions.

Avant de commencer

Configurez le VPN de site à site en fonction du guide de configuration de VPN.

Procédure**Étape 1**

Sur le nœud de contrôle de la grappe, passez en mode de configuration de la grappe.

cluster group *name*

Exemple :

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

Étape 2

Activez le VPN de site à site distribué.

vpn-mode distributed backup flat

ou

vpn-mode distributed backup remote-chassis

En mode de sauvegarde **plate**, des sessions de veille sont établies sur n'importe quel autre nœud de la grappe. Cela protégera les utilisateurs contre les défaillances de module; cependant, la protection contre les défaillances de châssis n'est pas garantie.

En mode de sauvegarde du **châssis distant**, des sessions de veille sont établies sur un nœud d'un autre châssis de la grappe. Cela protégera les utilisateurs contre les défaillances de module et de châssis.

Si le **châssis distant** est configuré dans un environnement de châssis unique (configuré intentionnellement ou à la suite d'une défaillance), aucune sauvegarde ne sera créée avant qu'un autre châssis rejoigne la grappe.

Exemple :

```
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]
```

Redistribuer les sessions du VPN S2S distribué

La redistribution des sessions actives redistribue la charge des sessions VPN actives sur les nœuds de la grappe. En raison de la nature dynamique des sessions de début et de fin, la redistribution des sessions actives permet d'équilibrer au mieux les sessions sur tous les nœuds de la grappe. Des actions de redistribution répétées optimiseront l'équilibre.

La redistribution peut être exécutée à tout moment, mais il est conseillé de la réaliser après toute modification de topologie dans la grappe et après qu'un nouveau nœud a rejoint la grappe. L'objectif de la redistribution est de créer une grappe VPN stable. Une grappe VPN stable a un nombre presque égal de sessions actives et de sauvegarde sur les nœuds.

Pour déplacer une session, la session de sauvegarde devient active et un autre nœud est sélectionné pour héberger une nouvelle session de sauvegarde. Le déplacement des sessions dépend de l'emplacement de la sauvegarde de la session active et du nombre de sessions actives déjà sur ce nœud de sauvegarde en particulier. Si le nœud de session de sauvegarde ne peut pas héberger la session active pour une raison quelconque, le nœud d'origine reste propriétaire de la session.

En mode de contexte multiple, la redistribution des sessions actives est effectuée au niveau du système, et non au niveau du contexte individuel. Elle ne se fait pas au niveau du contexte, car une session active dans un contexte peut déplacer un nœud qui contient de nombreuses sessions actives dans un contexte différent, créant ainsi une charge plus importante sur ce nœud de grappe.

Avant de commencer

- Activez les journaux système si vous souhaitez superviser l'activité de redistribution.
- Cette procédure doit être effectuée sur l'unité de contrôle de la grappe.

Procédure

Étape 1 Sur le nœud de contrôle, affichez la répartition des sessions actives et de sauvegarde dans la grappe.

show cluster vpn-sessiondb distribution

Exemple :

Les renseignements sur la distribution s'affichent comme suit :

```
ciscoasa# show cluster vpn-sessiondb distribution
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
```

Chaque ligne contient l'ID du membre, le nom du membre, le nombre de sessions actives et les membres sur lesquels se trouvent les sessions de sauvegarde. Dans l'exemple ci-dessus, on lit les renseignements comme suit :

- Le membre 0 a 209 sessions actives, 111 sessions sont sauvegardées sur le membre 1, 98 sessions sont sauvegardées sur le membre 2
- Le membre 1 a 204 sessions actives, 108 sessions sont sauvegardées sur le membre 0, 96 sessions sont sauvegardées sur le membre 2
- Le membre 2 n'a AUCUNE session active; par conséquent, aucun membre de la grappe ne sauvegarde de sessions pour ce nœud. Ce membre a récemment rejoint la grappe.

Étape 2

Redistribuez les sessions.

cluster redistribute vpn-sessiondb

Cette commande se termine immédiatement (sans message) pendant qu'elle s'exécute en arrière-plan.

Selon le nombre de sessions à redistribuer et la charge sur la grappe, cela peut prendre un certain temps. Les journaux système contenant les expressions suivantes (et d'autres détails sur le système non affichés ici) sont fournis au fur et à mesure que l'activité de redistribution se produit :

Expression du journal système	Notes
Redistribution de la session VPN démarrée	Nœud de contrôle uniquement
Demande envoyée pour déplacer <i>nombre</i> sessions de <i>orig-member-name</i> à <i>dest-member-name</i>	Nœud de contrôle uniquement
Échec de l'envoi du message de redistribution de session à <i>member-name</i>	Nœud de contrôle uniquement
Demande reçue pour déplacer <i>nombre</i> sessions de <i>orig-member-name</i> à <i>dest-member-name</i>	Nœud de données uniquement
Déplacement de <i>nombre</i> sessions à <i>member-name</i>	Le nombre de sessions actives déplacées vers la grappe nommée.
Échec de la réception de la réponse de déplacement de session depuis <i>dest-member-name</i>	Nœud de contrôle uniquement
Session VPN terminée	Nœud de contrôle uniquement
Modification de la topologie de grappe détectée. Redistribution de la session VPN abandonnée.	

Étape 3

Saisissez à nouveau la commande **show cluster vpn-sessiondb distribution** pour afficher les résultats.

FXOS : supprimer un nœud de la grappe

Les sections suivantes décrivent comment supprimer des nœuds de façon temporaire ou permanente de la grappe.

Suppression temporaire

Un nœud de grappe sera automatiquement supprimé de la grappe en raison d'une défaillance matérielle ou réseau, par exemple. Cette suppression est temporaire jusqu'à ce que les conditions soient rectifiées, et qu'il puisse rejoindre la grappe. Vous pouvez également désactiver manuellement la mise en grappe.

Pour vérifier si un appareil se trouve actuellement dans la grappe, vérifiez l'état de la grappe de l'application à l'aide de la commande **show cluster info** :

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

- **Disable clustering in the application** (désactivez la mise en grappe dans l'application) : Vous pouvez désactiver la mise en grappe à l'aide de l'interface de ligne de commande de l'application. Saisissez la commande **cluster remove unit name** pour supprimer tout nœud autre que celui auquel vous êtes connecté. La configuration de démarrage reste inchangée, ainsi que la dernière configuration synchronisée à partir du nœud de contrôle, afin que vous puissiez rajouter le nœud ultérieurement sans perdre votre configuration. Si vous saisissez cette commande sur un nœud de données pour supprimer le nœud de contrôle, un nouveau nœud de contrôle est élu.

Lorsqu'un périphérique devient inactif, toutes les interfaces de données sont fermées; seule l'interface de gestion peut envoyer et recevoir du trafic. Pour reprendre le flux de trafic, réactivez la mise en grappe. L'interface de gestion reste active en utilisant l'adresse IP que le nœud a reçue de la configuration de démarrage. Cependant, si vous rechargez et que le nœud est toujours inactif dans la grappe (par exemple, vous avez enregistré la configuration avec la mise en grappe désactivée), l'interface de gestion est désactivée.

Pour réactiver la mise en grappe, saisissez le **cluster group name** sur l'ASA, puis **enable**.

- **Désactivez l'instance d'application** : Au niveau de l'interface de ligne de commande de FXOS, voyez l'exemple suivant :

```
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 #
```

Pour réactiver :

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

- **Arrêtez le security module/engine** : Au niveau de l'interface de ligne de commande de FXOS, voyez l'exemple suivant :

```
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 #
```

Pour mettre sous tension :

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

- Arrêtez le châssis Au niveau de l'interface de ligne de commande de FXOS, voyez l'exemple suivant :

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

Suppression permanente

Vous pouvez supprimer définitivement un nœud de grappe en utilisant les méthodes suivantes.

- Supprimez le périphérique logique :Au niveau de l'interface de ligne de commande de FXOS, voyez l'exemple suivant :

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

- Supprimez le châssis ou le module de sécurité du service : si vous mettez un périphérique hors du service, vous pouvez ajouter du matériel de remplacement en tant que nouveau nœud de la grappe.

ASA : Gérer les membres de la grappe

Après avoir déployé la grappe, vous pouvez modifier la configuration et gérer les membres de celle-ci.

Devenir un membre inactif

Pour devenir un membre inactif de la grappe, désactivez la mise en grappe sur le nœud tout en laissant intacte la configuration de mise en grappe.



Remarque

Lorsqu'un ASA devient inactif (soit manuellement, soit à la suite d'un échec du contrôle d'intégrité), toutes les interfaces de données sont fermées; seule l'interface de gestion uniquement peut envoyer et recevoir du trafic. Pour reprendre le flux de trafic, réactivez la mise en grappe; ou vous pouvez supprimer complètement le nœud de la grappe. L'interface de gestion reste active en utilisant l'adresse IP que le nœud a reçue de l'ensemble d'adresses IP de la grappe. Cependant, si vous rechargez et que le nœud est toujours inactif dans la grappe (par exemple, vous avez enregistré la configuration avec la mise en grappe désactivée), alors l'interface de gestion est désactivée. Vous devez utiliser le port de console pour toute autre configuration.

Avant de commencer

- Vous devez utiliser le port de console; vous ne pouvez pas activer ni désactiver la mise en grappe à partir d'une connexion distante d'interface de ligne de commande.
- Pour le mode de contexte multiple, réalisez cette procédure dans l'espace d'exécution du système. Si vous n'êtes pas déjà en mode de configuration système, saisissez la commande **changeto system**.

Procédure

Étape 1 Entrez en mode de configuration de la grappe :

cluster group *name*

Exemple :

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

Étape 2 Désactivez la mise en grappe :

no enable

Si ce nœud était le nœud de contrôle, une nouvelle sélection de contrôle a lieu et un autre membre devient le nœud de contrôle.

La configuration de la grappe est conservée, vous pouvez donc réactiver la mise en grappe ultérieurement.

Désactiver une données

Pour désactiver un membre autre que le nœud auquel vous êtes connecté, effectuez les étapes suivantes.



Remarque

Lorsqu'un ASA devient inactif, toutes les interfaces de données sont fermées; seule l'interface de gestion uniquement peut envoyer et recevoir du trafic. Pour reprendre le flux de trafic, réactivez la mise en grappe. L'interface de gestion reste active en utilisant l'adresse IP que le nœud a reçue de l'ensemble d'adresses IP de la grappe. Cependant, si vous rechargez et que le nœud est toujours inactif dans la grappe (par exemple, vous avez enregistré la configuration avec la mise en grappe désactivée), l'interface de gestion est désactivée. Vous devez utiliser le port de console pour toute autre configuration.

Avant de commencer

Pour le mode de contexte multiple, réalisez cette procédure dans l'espace d'exécution du système. Si vous n'êtes pas déjà en mode de configuration système, saisissez la commande **changeto system**.

Procédure

Supprimez le nœud de la grappe.

cluster remove unit *node_name*

La configuration de démarrage reste inchangée, ainsi que la dernière configuration synchronisée à partir du nœud de contrôle, afin que vous puissiez rajouter le nœud ultérieurement sans perdre votre configuration. Si vous saisissez cette commande sur un nœud de données pour supprimer le nœud de contrôle, un nouveau nœud de contrôle est élu.

Pour afficher les noms des membres, saisissez **cluster remove unit ?**, ou saisissez la commande **show cluster info**.

Exemple :

```
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
```

Rejoindre la grappe

Si un nœud a été supprimé de la grappe, par exemple pour une interface défaillante ou si vous avez désactivé manuellement un membre, vous devez rejoindre manuellement la grappe.

Avant de commencer

- Vous devez utiliser le port de console pour réactiver la mise en grappe. Les autres interfaces sont fermées.
- Pour le mode de contexte multiple, réalisez cette procédure dans l'espace d'exécution du système. Si vous n'êtes pas déjà en mode de configuration système, saisissez la commande **changeto system**.
- Assurez-vous que le problème est résolu avant d'essayer de rejoindre la grappe.

Procédure

Étape 1 Sur la console, passez en mode de configuration de grappe :

cluster group name

Exemple :

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

Étape 2 Activez la mise en grappe.

enable

Changer l'unité de contrôle



Mise en garde

La méthode recommandée pour changer le nœud de contrôle est de désactiver la mise en grappe sur celui-ci en attendant un nouveau choix de contrôle, puis de réactiver la mise en grappe. Si vous devez spécifier le nœud exact que vous souhaitez voir devenir le nœud de contrôle, utilisez la procédure décrite dans cette section. Notez toutefois que pour les fonctionnalités centralisées, si vous forcez un changement de nœud de contrôle à l'aide de cette procédure, toutes les connexions sont abandonnées et vous devez rétablir les connexions sur le nouveau nœud de contrôle.

Pour modifier le nœud de contrôle, procédez comme suit.

Avant de commencer

Pour le mode de contexte multiple, réalisez cette procédure dans l'espace d'exécution du système. Si vous n'êtes pas déjà en mode de configuration système, saisissez la commande **changeto system**.

Procédure

Définissez un nouveau nœud comme nœud de contrôle :

cluster control-node unitnode_name

Exemple :

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

Vous devrez vous reconnecter à l'adresse IP de la grappe principale.

Pour afficher les noms des membres, saisissez **cluster control-node unit ?** (pour voir tous les noms à l'exception du nœud actuel), ou saisissez la commande **show cluster info**.

Exécuter une commande à l'échelle de la grappe

Pour envoyer une commande à tous les membres de la grappe ou à un membre en particulier, procédez comme suit. L'envoi d'une commande **show** à tous les membres collecte toutes les sorties et les affiche sur la console de l'unité actuelle. (Notez qu'il existe également des commandes d'affichage (show) que vous pouvez saisir sur l'unité de contrôle pour afficher les statistiques à l'échelle de la grappe.) D'autres commandes, telles que **capture** et **copy**, peuvent également profiter d'une exécution à l'échelle de la grappe.

Procédure

Envoyez une commande à tous les membres ou, si vous spécifiez le nom de l'unité, à un membre en particulier :

cluster exec [unit unit_name] command

Exemple :

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

Pour afficher les noms des membres, saisissez **cluster exec unit ?** (pour afficher tous les noms à l'exception de l'unité actuelle) ou saisissez la commande **show cluster info**.

Exemples

Pour copier simultanément le même fichier de capture de toutes les unités de la grappe vers un serveur TFTP, saisissez la commande suivante sur l'unité de contrôle :

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

Plusieurs fichiers PCAP, un pour chaque unité, sont copiés sur le serveur TFTP. Le nom du fichier de capture de destination est automatiquement associé au nom de l'unité, par exemple capture1_asa1.pcap, capture1_asa2.pcap, etc. Dans cet exemple, asa1 et asa2 sont des noms d'unités de grappe.

L'exemple de sortie suivant pour la commande **cluster exec show memory** affiche les renseignements sur la mémoire pour chaque membre de la grappe :

```
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 : supervision de la grappe ASA sur le Châssis Firepower 4100/9300

Vous pouvez superviser et dépanner l'état de la grappe et les connexions.

Supervision de l'état de la grappe

Consultez les commandes suivantes pour superviser l'état de la grappe :

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

En l'absence de mot clé, la commande **show cluster info** affiche l'état de tous les membres de la grappe.

La commande **show cluster info health** affiche l'intégrité actuelle des interfaces, des unités et de la grappe dans son ensemble.

Consultez la sortie suivante pour la commande **show cluster info** :

```
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**

Indique si l'unité de la grappe rejoindra automatiquement la grappe après un délai et si les conditions de défaillance (comme l'attente de la licence, l'échec de la vérification de l'intégrité du châssis, etc.) sont supprimées. Si l'unité est désactivée de façon permanente ou si l'unité est déjà dans la grappe, cette commande n'affichera aucune sortie.

Consultez les sorties suivantes pour la commande **show cluster info auto-join** :

```
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]}**

Affiche les statistiques relatives au transport pour les éléments suivants :

- **asp** – Statistiques de transport du plan de données.
- **cp** – Statistiques de transport du plan de commande.

Si vous saisissez le mot clé **detail**, vous pouvez afficher l'utilisation du protocole de transport fiable de la grappe afin de cerner les problèmes d'abandon de paquet lorsque le tampon est plein dans le plan de commande. Consultez la sortie suivante pour la commande **show cluster info transport cp detail** :

```

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

Affiche l'historique de la grappe, ainsi que les messages d'erreur indiquant pourquoi une unité de grappe n'a pas pu se joindre ou pourquoi une unité a quitté la grappe.

Capture des paquets à l'échelle de la grappe

Consultez la commande suivante pour capturer des paquets dans une grappe :

cluster exec capture

Pour prendre en charge le dépannage à l'échelle de la grappe, vous pouvez activer la capture du trafic spécifique à la grappe sur le nœud de contrôle à l'aide de la commande **cluster exec capture**, qui est ensuite automatiquement activée sur tous les nœuds de données de la grappe.

Supervision des ressources de la grappe

Consultez la commande suivante pour surveiller les ressources de la grappe :

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

Affiche les données agrégées pour l'ensemble de la grappe. Les options disponibles dépendent du type de données.

Supervision du trafic de la grappe

Consultez la commande suivante pour superviser le trafic de la grappe :

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

La commande **show conn** indique si un flux est un flux de directeur, de sauvegarde ou de transfert. Utilisez la commande **cluster exec show conn** sur n'importe quelle unité pour afficher toutes les connexions. Cette commande peut montrer comment le trafic d'un flux unique arrive à différents ASA dans la grappe. Le débit de la grappe dépend de l'efficacité et de la configuration de l'équilibrage de charges. Cette commande offre un moyen simple d'afficher le trafic d'une connexion dans la grappe et peut vous aider à comprendre comment un équilibreur de charges peut affecter les performances d'un flux.

Voici un exemple de sortie de la commande **show conn detail** :

```
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]**

Les commandes **show cluster info conn-distribution** et **show cluster info packet-distribution** affichent la répartition du trafic sur toutes les unités de grappe. Ces commandes peuvent vous aider à évaluer et à ajuster l'équilibre de charges externe.

La commande **show cluster info loadbalance** affiche les statistiques de rééquilibrage des connexions.

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

La commande **show cluster info load-monitor** affiche la charge de trafic pour les membres de la grappe pour le dernier intervalle ainsi que la moyenne du nombre total d'intervalles configurés (30 par défaut). Utilisez le mot clé **details** pour afficher la valeur de chaque mesure à chaque intervalle.

```

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}**

Affiche les données agrégées pour l'ensemble de la grappe. Les options disponibles dépendent du type de données.

Consultez la sortie suivante pour la commande **show cluster access-list** :

```
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) 0xclce5c49
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) 0xc7dba0d

```

Pour afficher le nombre agrégé des connexions utilisées pour toutes les unités, saisissez :

```

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**

Cette commande est utile pour le dépannage des chemins de données.

Supervision du routage de la grappe

Consultez les commandes suivantes pour le routage de la grappe :

- **show route cluster**
- **debug route cluster**

Affiche les renseignements sur la grappe pour le routage.

- **show lisp eid**

Affiche le tableau EID de l'ASA qui indique les EID et les ID de site.

Consultez la sortie suivante de la commande **cluster exec show lisp eid**.

```
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**

Cette commande est utile pour la résolution de problèmes.

Supervision du VPN S2S distribué

Utilisez les commandes suivantes pour surveiller l'état et la répartition des sessions VPN :

- La répartition globale des sessions est fournie à l'aide de **show cluster vpn-sessiondb distribution**. Si elle est exécutée dans un environnement de contexte multiple, cette commande doit être exécutée dans l'espace d'exécution du système.

Cette commande **show** offre un aperçu rapide des sessions, plutôt que d'avoir à exécuter **show vpn-sessiondb summary** sur chaque nœud.

- Une vue unifiée des connexions VPN sur la grappe à l'aide de la commande **show cluster vpn-sessiondb summary** est également disponible.
- La supervision des périphériques individuels à l'aide de la commande **show vpn-sessiondb** affiche le nombre de sessions actives et de sauvegarde sur un périphérique en plus des renseignements VPN habituels.

Configuration de la journalisation pour la mise en grappe

Consultez la commande suivante pour configurer la journalisation pour la mise en grappe :

logging device-id

Chaque nœud de la grappe génère des messages de journal système de manière indépendante. Vous pouvez utiliser la commande **logging device-id** pour générer des messages de journal système avec des ID de périphérique identiques ou différents afin de faire apparaître les messages comme provenant de nœuds identiques ou différents dans la grappe.

Débogage de la mise en grappe

Consultez les commandes suivantes pour le débogage de la mise en grappe :

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

Affiche les messages de débogage pour la mise en grappe.

- **debug service-module**

Affiche les messages de débogage pour les problèmes au niveau de la lame, y compris les problèmes de vérification de l'intégrité entre le superviseur et l'application.

- **show cluster info trace**

La commande **show cluster info trace** affiche les renseignements de débogage pour un dépannage ultérieur.

Consultez la sortie suivante pour la commande **show cluster info trace** :

```
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_KEEPLIVE from 80-1 at
MASTER
```

Par exemple, si vous voyez les messages suivants indiquant que deux nœuds ayant le même nom **local-unit** agissent en tant que nœud de contrôle, cela peut signifier que deux nœuds ont le même nom **local-unit** (vérifiez votre configuration) ou qu'un nœud reçoit ses propres messages de diffusion (vérifiez votre réseau).

```
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 msec,
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
```

Dépannage du VPN S2S distribué

Notifications de VPN distribuées

Vous recevrez des messages contenant les expressions identifiées lorsque les situations d'erreur suivantes se produisent sur une grappe exécutant le VPN distribué :

Situation	Notification
Si un nœud de données de grappe existant ou en cours de connexion n'est pas en mode VPN distribué lors de la tentative de connexion à la grappe :	Nouveau membre de la grappe (<i>member-name</i>) refusé en raison d'une incompatibilité de mode VPN. et Le nœud de contrôle (<i>control-name</i>) refuse la demande d'inscription de l'unité (<i>unit-name</i>) pour la raison suivante : les capacités du mode VPN ne sont pas compatibles avec la configuration du nœud de contrôle
Si la licence n'est pas configurée correctement sur un membre de la grappe pour le VPN distribué :	ERREUR : le nœud de contrôle a demandé le changement de mode VPN de la grappe sur distribué. Impossible de changer de mode en raison d'une licence d'opérateur manquante.
Si l'horodatage ou l'ID de membre est non valide dans le SPI d'un paquet IKEv2 reçu :	SPI expiré reçu ou SPI corrompu détecté
Si la grappe ne peut pas créer de session de sauvegarde :	Échec de la création de la sauvegarde pour une session IKEv2.
Erreur de traitement du contact initial IKEv2 :	Négociation IKEv2 abandonnée en raison d'une ERREUR : session de sauvegarde obsolète trouvée sur la sauvegarde
Problèmes de redistribution :	Échec de l'envoi du message de redistribution de session à <i>member-name</i> Échec de la réception de la réponse de déplacement de session de <i>member-name</i> (nœud de contrôle uniquement)
Si la topologie change lors de la redistribution des sessions :	Modification de la topologie de grappe détectée. Redistribution de la session VPN abandonnée.

Vous pouvez rencontrer l'une des situations suivantes :

- Les sessions VPN de site à site sont distribuées à un seul des châssis d'une grappe lorsque le commutateur Nexus 7K est configuré avec un port de couche 4 comme algorithme d'équilibrage de charges à l'aide de la commande **port-channel load-balance src-dst l4port**. Voici un exemple d'allocation de session de grappe :

```
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
```

Étant donné que le VPN IKEv2 de site à site utilise le port 500 pour les ports source et de destination, les paquets IKE ne sont envoyés qu'à l'une des liaisons du canal de port connecté entre le Nexus 7K et le châssis.

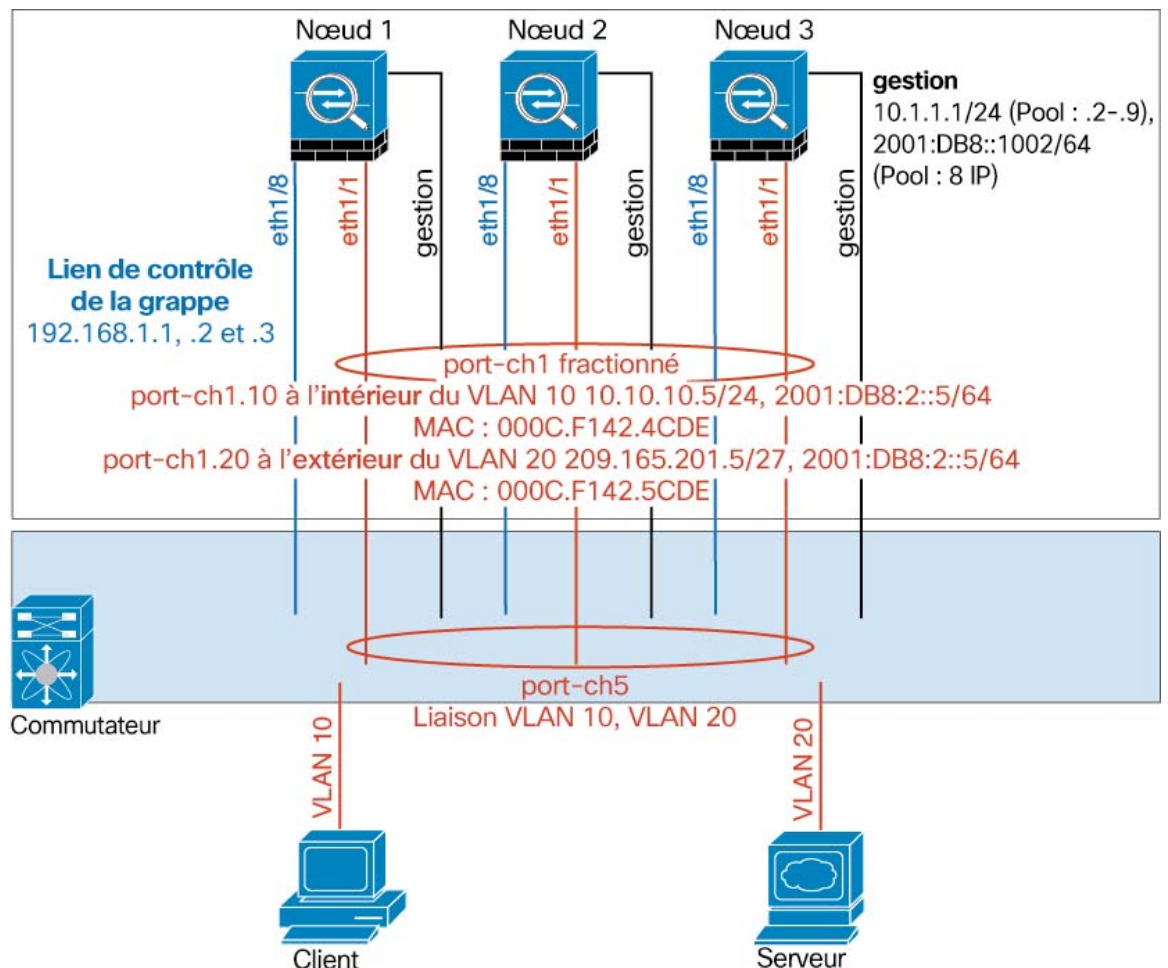
Modifiez l'algorithme d'équilibrage de charges Nexus 7K sur l'adresse IP et le port de couche 4 à l'aide de **port-channel load-balance src-dst ip-l4port**. Ensuite, les paquets IKE sont envoyés à toutes les liaisons et donc à tous les nœuds.

Pour un ajustement plus immédiat, sur le nœud de contrôle de la grappe, exécutez : **cluster redistribute vpn-sessiondb** pour redistribuer les sessions VPN actives aux nœuds de grappe des autres châssis.

Exemples de mise en grappe d'ASA

Ces exemples comprennent des déploiements typiques.

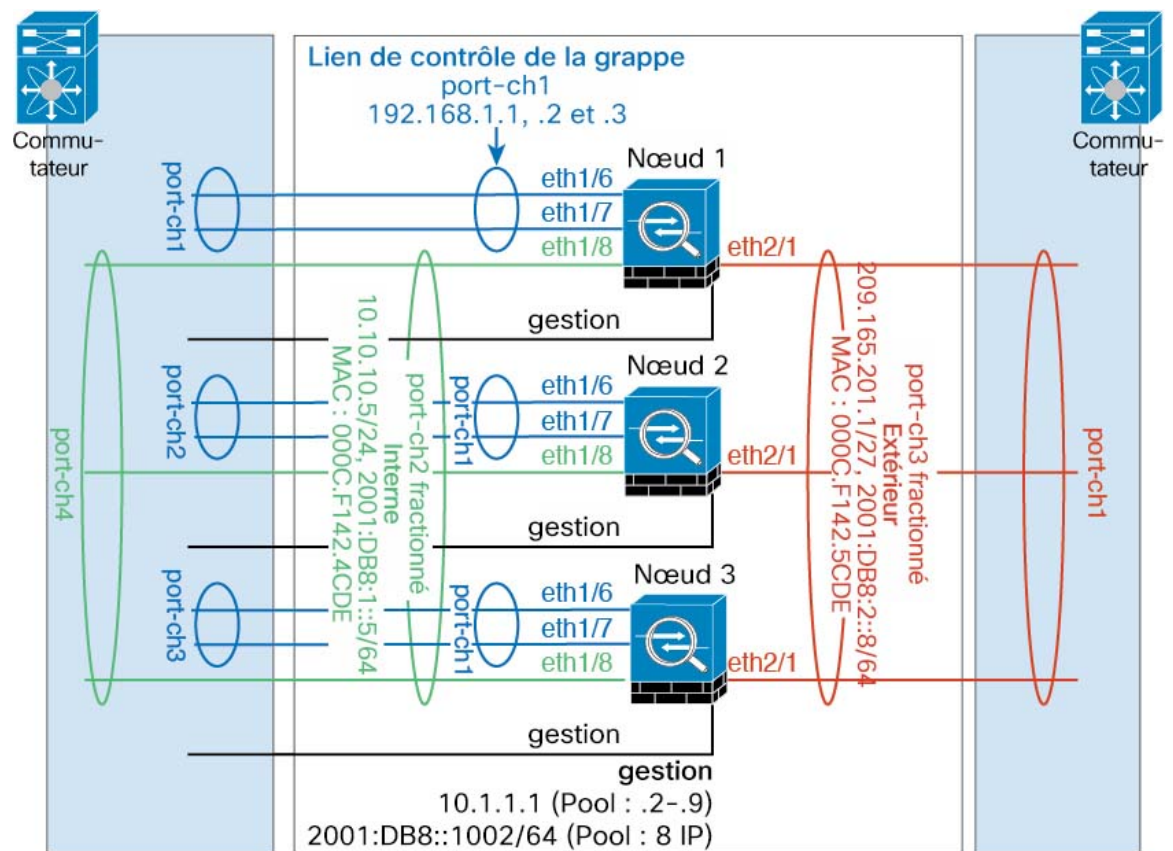
Pare-feu sur clé



Le trafic de données provenant de différents domaines de sécurité est associé à différents VLAN, par exemple, le VLAN 10 pour le réseau interne et le VLAN 20 pour le réseau externe. Chaque ASA dispose d'un seul port physique connecté au commutateur ou routeur externe. Le regroupement de liaisons est activé de sorte que tous les paquets sur la liaison physique soient encapsulés dans une norme 802.1q. L'ASA sert de pare-feu entre le VLAN 10 et le VLAN 20.

Lorsque vous utilisez des EtherChannels étendus, toutes les liaisons de données sont regroupées dans un seul EtherChannel du côté du commutateur. Si l'ASA n'est plus disponible, le commutateur rééquilibre le trafic entre les unités restantes.

Ségrégation du trafic



Vous pourriez souhaiter une séparation physique du trafic entre le réseau interne et le réseau externe.

Comme le montre le diagramme ci-dessus, il y a un EtherChannel étendu sur le côté gauche qui se connecte au commutateur interne et l'autre sur le côté droit au commutateur externe. Vous pouvez également créer des sous-interfaces VLAN sur chaque EtherChannel, au besoin.

Configuration OTV pour la mise en grappe inter-sites en mode routé

La réussite de la mise en grappe inter-sites pour le mode routé avec les EtherChannels étendus dépend de la configuration et de la supervision appropriées de l'OTV. L'OTV joue un rôle majeur en transférant les paquets sur l'interface DCI. L'OTV transfère les paquets de monodiffusion sur l'interface DCI uniquement lorsqu'il

apprend l'adresse MAC dans sa table de transfert. Si l'adresse MAC n'est pas apprise dans la table de transfert OTV, il abandonnera les paquets de monodiffusion.

Exemple de configuration OTV

```
//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

```

Modifications du filtre OTV requises en raison d'une défaillance du site

Si un site tombe en panne, les filtres doivent être supprimés de l'OTV, car il n'est plus nécessaire de bloquer l'adresse MAC globale. Certaines configurations supplémentaires sont nécessaires.

Vous devez ajouter une entrée statique pour l'adresse MAC globale ASA sur le commutateur OTV dans le site qui est fonctionnel. Cette entrée permettra à l'OTV à l'autre extrémité d'ajouter ces entrées sur l'interface de superposition. Cette étape est nécessaire, car si le serveur et le client ont déjà l'entrée ARP pour l'ASA, ce qui est le cas pour les connexions existantes, ils n'envoieront pas à nouveau l'ARP. Par conséquent, l'OTV n'aura pas la possibilité d'apprendre l'adresse MAC globale ASA dans sa table de transfert. Étant donné que l'OTV n'a pas l'adresse MAC globale dans sa table de transfert et que, conformément à sa conception, il ne diffusera pas de paquets de monodiffusion sur l'interface de superposition, il abandonnera les paquets de monodiffusion à l'adresse MAC globale du serveur, et les connexions existantes seront interrompues.

```

//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

```

Lorsque l'autre site est restauré, vous devez ajouter de nouveau les filtres et supprimer cette entrée statique sur l'OTV. Il est très important d'effacer le tableau des adresses MAC dynamiques sur les deux OTV pour effacer l'entrée de superposition pour l'adresse MAC globale.

Effacement du tableau d'adresses MAC

Lorsqu'un site tombe en panne et qu'une entrée statique pour l'adresse MAC globale est ajoutée à l'OTV, vous devez laisser l'autre OTV apprendre l'adresse MAC globale sur l'interface de superposition. Une fois l'autre site rétabli, ces entrées doivent être effacées. Assurez-vous d'effacer la table d'adresses MAC pour vous assurer que l'OTV ne conserve aucune de ces entrées dans sa table de transfert.

```
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
```

Supervision du cache OTV ARP

Le service OTV gère un cache ARP pour mandataire ARP pour les adresses IP qu'il a apprises sur l'interface OTV.

```
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#
```

Exemples de mise en grappe inter-sites

Les exemples suivants montrent les déploiements de grappe pris en charge.

Exemple de mode de routage EtherChannel étendu avec des adresses MAC et IP propres au site

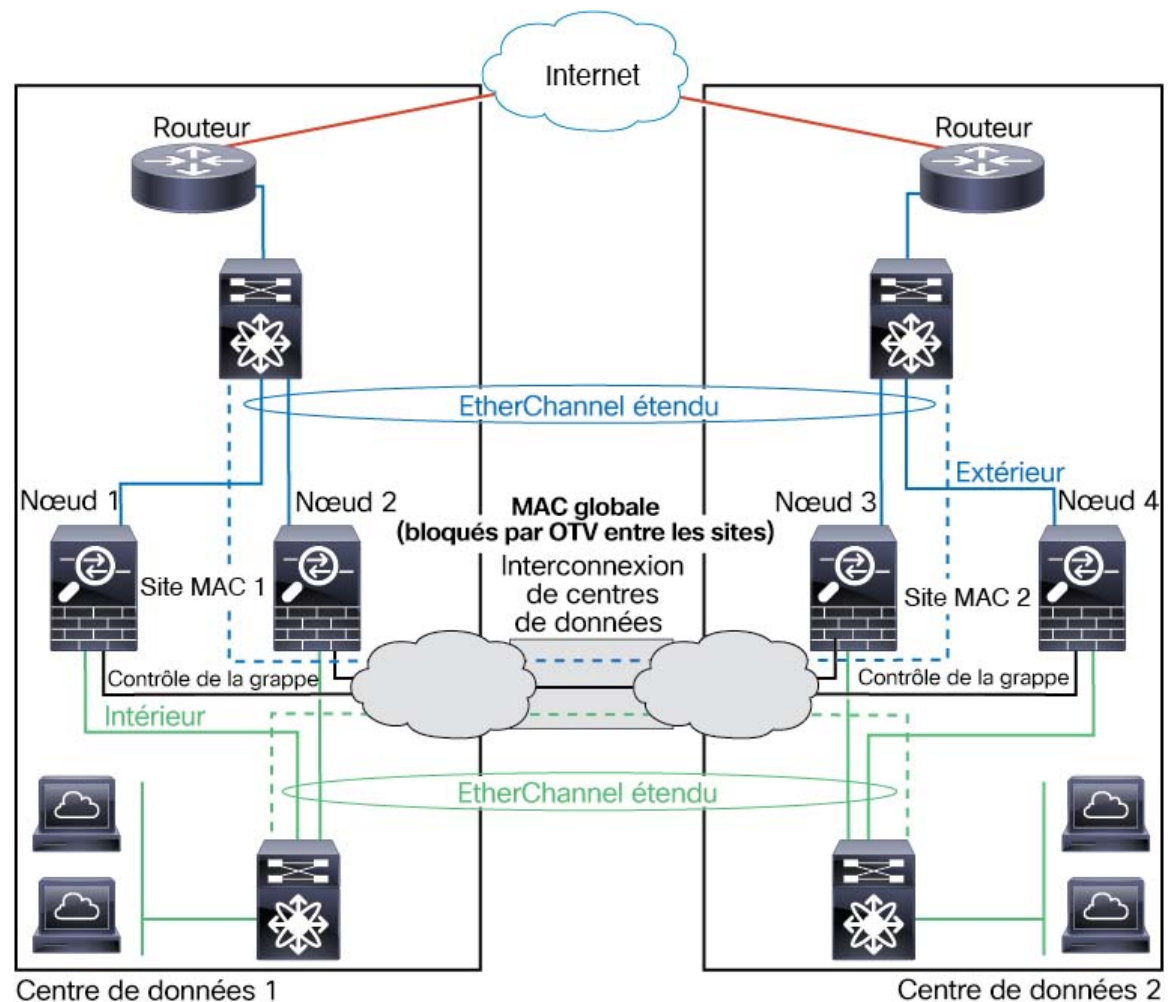
L'exemple suivant montre deux membres de la grappe dans chacun des deux centres de données, placés entre le routeur de passerelle et un réseau interne sur chaque site (insertion est-ouest). Les membres de la grappe sont connectés par la liaison de commande de grappe sur DCI. Les membres de la grappe de chaque site se connectent aux commutateurs locaux en utilisant des EtherChannels étendus pour le réseau interne et externe. Chaque EtherChannel est réparti sur tous les châssis de la grappe.

Les VLAN de données sont étendus entre les sites à l'aide de la virtualisation du transport par superposition (OTV) (ou d'un outil similaire). Vous devez ajouter des filtres bloquant l'adresse MAC globale pour empêcher le trafic de traverser l'interface DCI vers l'autre site lorsqu'il est destiné à la grappe. Si les nœuds de la grappe d'un site deviennent inaccessibles, vous devez supprimer les filtres pour que le trafic puisse être dirigé vers les nœuds de la grappe de l'autre site. Vous devez utiliser les adresses VACL pour filtrer l'adresse MAC globale. Pour certains commutateurs, comme le Nexus avec la carte de ligne de la gamme F3, vous devez également utiliser l'inspection ARP pour bloquer les paquets ARP de l'adresse MAC globale. L'inspection ARP nécessite que vous définissiez à la fois l'adresse MAC et l'adresse IP du site sur l'ASA. Si vous configurez uniquement l'adresse MAC du site N'oubliez pas de désactiver l'inspection ARP.

La grappe sert de passerelle pour les réseaux internes. Le MAC virtuel global, qui est partagé sur tous les nœuds de la grappe, est utilisé uniquement pour recevoir des paquets. Les paquets sortants utilisent une adresse MAC spécifique au site pour chaque grappe de contrôleurs de domaine. Cette fonctionnalité empêche les commutateurs d'apprendre la même adresse MAC globale à partir des deux sites sur deux ports différents, ce qui entraîne une oscillation MAC; au lieu de cela, ils connaissent uniquement l'adresse MAC du site.

Dans ce scénario :

- Tous les paquets de sortie envoyés par la grappe utilisent l'adresse MAC du site et sont localisés dans le centre de données.
- Tous les paquets entrants vers la grappe sont envoyés à l'aide de l'adresse MAC globale, de sorte qu'ils peuvent être reçus par n'importe quel nœud des deux sites; les filtres de l'OTV localisent le trafic dans le centre de données.



Exemple intersites nord-sud de mode transparent de l'EtherChannel étendu

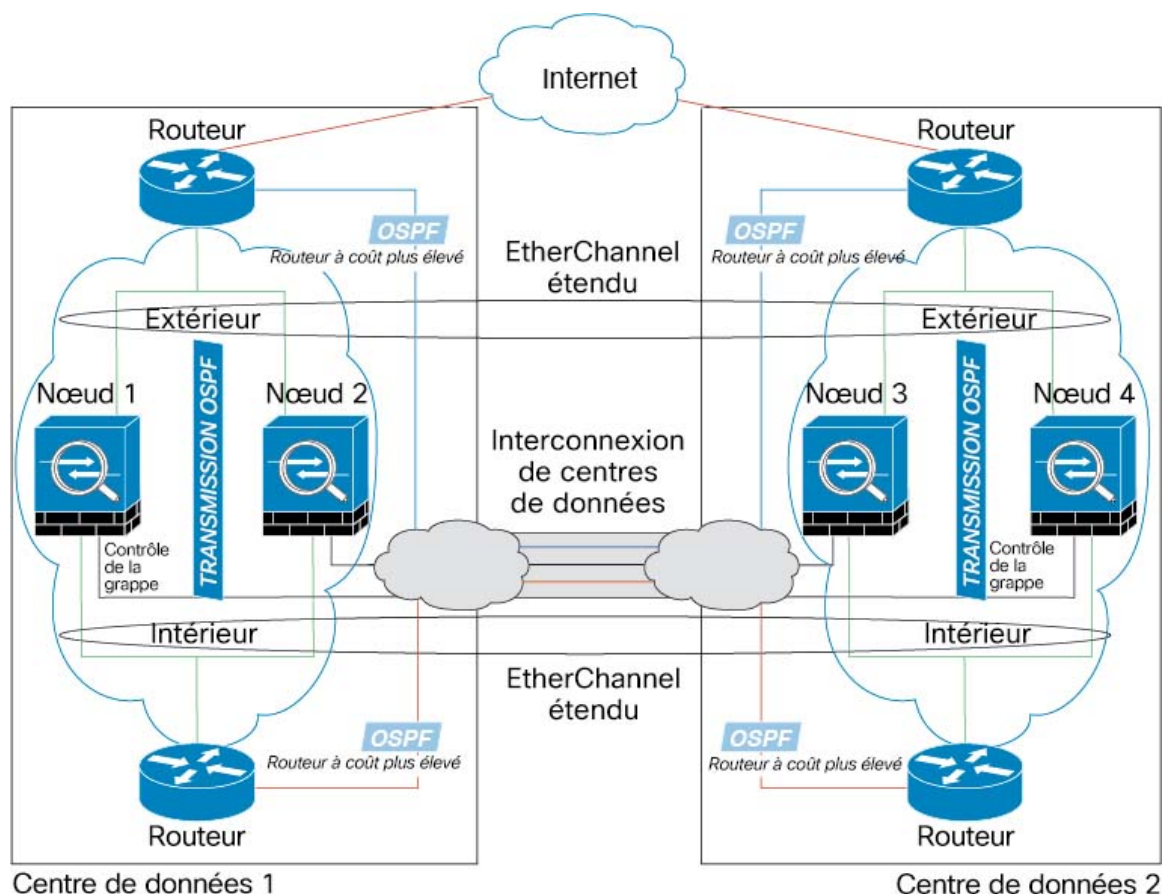
L'exemple suivant montre 2 membres de la grappe dans chacun des 2 centres de données placés entre les routeurs internes et externes (insertion nord-sud). Les membres de la grappe sont connectés par la liaison de commande de grappe sur DCI. Les membres de la grappe de chaque site se connectent aux commutateurs

locaux en utilisant des EtherChannels étendus pour l'intérieur et l'extérieur. Chaque EtherChannel est réparti sur tous les châssis de la grappe.

Les routeurs internes et externes de chaque centre de données utilisent le protocole OSPF, qui est transmis par les périphériques ASA transparents. Contrairement aux MAC, les adresses IP des routeurs sont uniques sur tous les routeurs. En attribuant un routage à coût plus élevé sur l'interface DCI, le trafic reste dans chaque centre de données, sauf si tous les membres de la grappe d'un site donné tombent en panne. La voie de routage la moins dispendieuse qui traverse les ASA doit traverser le même groupe de ponts sur chaque site pour que la grappe puisse maintenir des connexions dissymétriques. En cas de défaillance de tous les membres de la grappe sur un site, le trafic passe de chaque routeur sur l'interface DCI aux membres de la grappe sur l'autre site.

La mise en œuvre des commutateurs sur chaque site peut inclure :

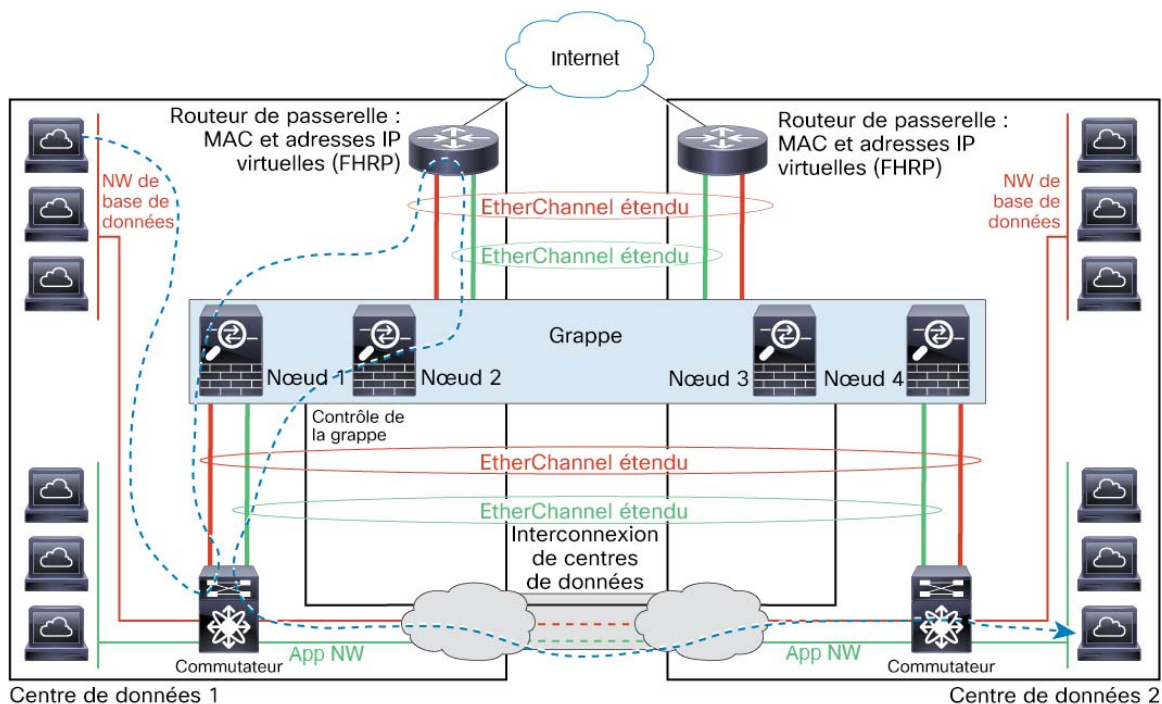
- VSS inter-sites, vPC, StackWise ou StackWise Virtual : dans ce scénario, vous installez un commutateur au centre de données 1, et l'autre au centre de données 2. Une option consiste à ce que les nœuds de la grappe de chaque centre de données se connectent uniquement au commutateur local, tandis que le trafic du commutateur redondant traverse le DCI. Dans ce cas, les connexions sont pour la plupart conservées localement à chaque centre de données. Vous pouvez éventuellement connecter chaque nœud aux deux commutateurs de la DCI si la DCI peut gérer le trafic supplémentaire. Dans ce cas, le trafic est réparti entre les centres de données, il est donc essentiel que l'interface DCI soit très robuste.
- Local VSS, vPC, StackWise ou StackWise Virtual sur chaque site : Pour une meilleure redondance des commutateurs, vous pouvez installer 2 paires de commutateurs redondants distincts sur chaque site. Dans ce cas, bien que les nœuds de la grappe aient toujours un EtherChannel étendu avec le châssis du centre de données 1 connecté uniquement aux deux commutateurs locaux et le châssis du centre de données 2 connecté à ces commutateurs locaux, l'EtherChannel étendu est essentiellement « divisé ». Chaque système de commutation local redondant voit l'EtherChannel étendu comme un EtherChannel local au site.



Exemple inter-sites est-ouest en mode transparent de l'EtherChannel étendu

L'exemple suivant montre deux membres de la grappe dans chacun des deux centres de données, placés entre le routeur de la passerelle et deux réseaux internes sur chaque site, le réseau d'applications et le réseau de base de données (insertion est-ouest). Les membres de la grappe sont connectés par la liaison de commande de grappe sur DCI. Les membres de la grappe de chaque site se connectent aux commutateurs locaux au moyen d'EtherChannels étendus pour les réseaux App et DB à l'intérieur et à l'extérieur. Chaque EtherChannel est réparti sur tous les châssis de la grappe.

Le routeur de la passerelle de chaque site utilise un FHRP comme le HSRP pour fournir les mêmes adresses MAC virtuelles et IP de destination sur chaque site. Une bonne pratique pour éviter le basculement involontaire des adresses MAC consiste à ajouter de manière statique les adresses MAC réelles des routeurs de passerelle au tableau d'adresses MAC de l'ASA à l'aide de la commande `mac-address-table static outside_interface mac_address`. Sans ces entrées, si la passerelle du site 1 communique avec la passerelle du site 2, ce trafic pourrait passer par l'ASA et tenter d'atteindre le site 2 à partir de l'interface interne, ce qui poserait des problèmes. Les VLAN de données sont étendus entre les sites à l'aide de la virtualisation du transport par superposition (OTV) (ou d'un outil similaire). Vous devez ajouter des filtres pour empêcher le trafic de traverser l'interface DCI vers l'autre site lorsqu'il est destiné au routeur de la passerelle. Si le routeur de la passerelle d'un site devient inaccessible, vous devez supprimer les filtres pour que le trafic puisse être dirigé vers le routeur de la passerelle de l'autre site.



Référence pour la mise en grappe

Cette section comprend des renseignements supplémentaires sur le fonctionnement de la mise en grappe.

Fonctionnalités et mise en grappe de l'ASA

Certaines fonctionnalités d'ASA ne sont pas prises en charge avec la mise en grappe d'ASA, et d'autres le sont uniquement sur le nœud de contrôle. Pour une utilisation correcte, d'autres fonctionnalités peuvent comporter des mises en garde.

Fonctionnalités non prises en charge par la mise en grappe

Ces fonctionnalités ne peuvent pas être configurées lorsque la mise en grappe est activée, et les commandes seront rejetées.

- Fonctionnalités de communication unifiée qui reposent sur le mandataire TLS
- VPN d'accès à distance (VPN SSL et VPN IPsec)
- Interfaces de tunnel virtuel (VTI)
- Routage IS-IS
- Les inspections d'application suivantes :
 - CTIQBE
 - H323, H225 et RAS

- Intercommunication IPsec
 - MGCP
 - MMP
 - RTSP
 - SCCP (Skinny)
 - WAAS
 - WCCP
-
- Filtre de trafic de réseau de zombies
 - Auto Update Server
 - Client DHCP, serveur et serveur mandataire. Le relais DHCP est pris en charge.
 - Équilibrage de la charge VPN
 - Basculement
 - Routage et pont intégrés
 - Détection de connexion inactive (DCD)
 - Mode FIPS

Fonctionnalités centralisées pour la mise en grappe

Les fonctionnalités suivantes ne sont prises en charge que sur le nœud de contrôle et ne sont pas adaptées à la grappe.



Remarque

Le trafic pour les fonctionnalités centralisées est acheminé des nœuds membres vers le nœud de contrôle par la liaison de commande de grappe.

Si vous utilisez la fonctionnalité de rééquilibrage, le trafic des fonctionnalités centralisées peut être rééquilibré vers des nœuds sans contrôle avant que le trafic ne soit classé comme fonctionnalité centralisée. Si cela se produit, le trafic est renvoyé au nœud de contrôle.

Pour les fonctionnalités centralisées, si le nœud de contrôle tombe en panne, toutes les connexions sont abandonnées et vous devez rétablir les connexions sur le nouveau nœud de contrôle.

- Les inspections d'application suivantes :
 - DCERPC
 - ESMTP
 - IM
 - NetBIOS
 - PPTP

- RADIUS
 - RSH
 - SNMP
 - SQLNET
 - SunRPC
 - TFTP
 - XDMCP
- Surveillance du routage statique
 - Autorisation et authentification pour l'accès réseau La comptabilité est décentralisée.
 - Services de filtrage
 - VPN de site à site

En mode centralisé, les connexions VPN sont établies avec le nœud de commande de grappe uniquement. Il s'agit du mode par défaut pour la mise en grappe VPN. Le VPN de site à site peut également être déployé en mode VPN distribué, dans lequel les connexions VPN S2S IKEv2 sont réparties sur plusieurs nœuds.
 - Traitement du protocole du plan de contrôle de multidiffusion IGMP (le transfert du plan de données est distribué dans la grappe)
 - Traitement du protocole du plan de contrôle de multidiffusion PIM (le transfert du plan de données est distribué dans la grappe)
 - Routage dynamique

Fonctionnalités appliquées aux unités individuelles

Ces fonctionnalités sont appliquées à chaque nœud ASA, plutôt qu'à la grappe dans son ensemble ou au nœud de contrôle.

- QoS : la politique QoS est synchronisée dans la grappe dans le cadre de la duplication de la configuration. Cependant, la politique est appliquée sur chaque nœud indépendamment. Par exemple, si vous configurez le contrôle en sortie, les valeurs de débit de conformité et de rafale de conformité s'appliquent au trafic sortant d'un ASA particulier. Dans une grappe à 3 nœuds et avec le trafic réparti uniformément, le taux de conformité devient en fait trois fois supérieur au débit de la grappe.
- Détection des menaces : la détection des menaces fonctionne sur chaque nœud indépendamment; par exemple, les statistiques principales sont propres au nœud. La détection de l'analyse de port, par exemple, ne fonctionne pas, car l'analyse du trafic sera équilibrée entre tous les nœuds et un nœud ne verra pas tout le trafic.
- Gestion des ressources : la gestion des ressources en mode contexte multiple est appliquée séparément sur chaque nœud en fonction de l'utilisation locale.
- Trafic LISP : le trafic LISP sur le port UDP 4342 est inspecté par chaque nœud de réception, mais aucun directeur n'est affecté. Chaque nœud ajoute au tableau EID qui est partagé dans la grappe, mais le trafic LISP lui-même ne participe pas au partage de l'état de la grappe.

AAA pour l'accès réseau et la mise en grappe

L'AAA pour l'accès réseau comprend trois composants : l'authentification, l'autorisation et la gestion de comptes. L'authentification et l'autorisation sont mises en œuvre en tant que fonctionnalités centralisées sur le nœud de contrôle de mise en grappe avec la duplication des structures de données sur les nœuds de données de grappe. Si un nœud de contrôle est choisi, le nouveau nœud de contrôle disposera de toute l'information nécessaire pour continuer le fonctionnement ininterrompu des utilisateurs authentifiés établis et de leurs autorisations associées. Les délais d'inactivité et absolue pour l'authentification des utilisateurs sont conservés lors d'un changement de nœud de contrôle se produit.

La comptabilité est mise en œuvre en tant que fonctionnalité distribuée dans une grappe. La comptabilité est effectuée flux par flux, de sorte que le nœud de la grappe possédant un flux enverra les messages de démarrage et d'arrêt de comptabilité au serveur AAA lorsque la comptabilité est configurée pour un flux.

Paramètres de connexion

Les limites de connexion s'appliquent à l'ensemble de la grappe (voir les commandes **set connection conn-max**, **set connection embryonic-conn-max**, **set connection per-client-embryonic-max**, et **set connection per-client-max**). Chaque nœud dispose d'une estimation des valeurs de compteur à l'échelle de la grappe en fonction des messages de diffusion. Pour des raisons d'efficacité, la limite de connexion configurée dans la grappe pourrait ne pas être appliquée exactement au nombre limite. Chaque nœud peut surévaluer ou sous-évaluer la valeur du compteur à l'échelle de la grappe à tout moment. Cependant, les informations seront mises à jour au fil du temps dans une grappe à charge équilibrée.

FTP et mise en grappe

- Si le canal de données et les flux du canal de contrôle FTP appartiennent à différents membres de la grappe, le propriétaire du canal de données enverra périodiquement des mises à jour du délai d'inactivité au propriétaire du canal de contrôle et mettra à jour la valeur du délai d'inactivité. Cependant, si le propriétaire du flux de contrôle est rechargé et que le flux de contrôle est réhébergé, la relation de flux parent/enfant ne sera plus maintenue; le délai d'inactivité du flux de contrôle ne sera pas mis à jour.
- Si vous utilisez l'AAA pour l'accès FTP, le flux du canal de contrôle est centralisé sur le nœud de contrôle.

Inspection ICMP

Le flux des paquets d'erreurs ICMP et ICMP dans la grappe varie selon que l'inspection d'erreurs ICMP ou ICMP est activée ou non. Sans inspection ICMP, ICMP est un flux unidirectionnel et il n'y a pas de prise en charge de flux directeur. Avec l'inspection ICMP, le flux ICMP devient bidirectionnel et est soutenu par un flux directeur/de sauvegarde. L'une des différences avec un flux ICMP inspecté réside dans le traitement par le directeur du paquet transféré : le directeur transmettra le paquet de réponse d'écho ICMP au propriétaire du flux au lieu de retourner le paquet au transitaire.

Routage multidiffusion et mise en grappe

L'unité de contrôle gère tous les paquets de routage de multidiffusion et les paquets de données jusqu'à ce que le transfert rapide soit établi. Une fois la connexion établie, chaque unité de données peut transférer des paquets de données en multidiffusion.

NAT et mise en grappe

La NAT peut affecter le débit global de la grappe. Les paquets NAT entrants et sortants peuvent être envoyés à différents ASA dans la grappe, car l'algorithme d'équilibrage de charge repose sur les adresses IP et les

ports, et la NAT fait en sorte que les paquets entrants et sortants aient des adresses IP ou des ports différents. Lorsqu'un paquet arrive à ASA qui n'est pas le propriétaire NAT, il est transféré sur la liaison de commande de grappe vers le propriétaire, ce qui entraîne un trafic important sur la liaison de commande de grappe. Notez que le nœud de réception ne crée pas de flux de transfert vers le propriétaire, car le propriétaire de la NAT peut ne pas créer de connexion pour le paquet en fonction des résultats des vérifications de sécurité et des politiques.

Si vous souhaitez toujours utiliser la NAT en grappe, tenez compte des directives suivantes :

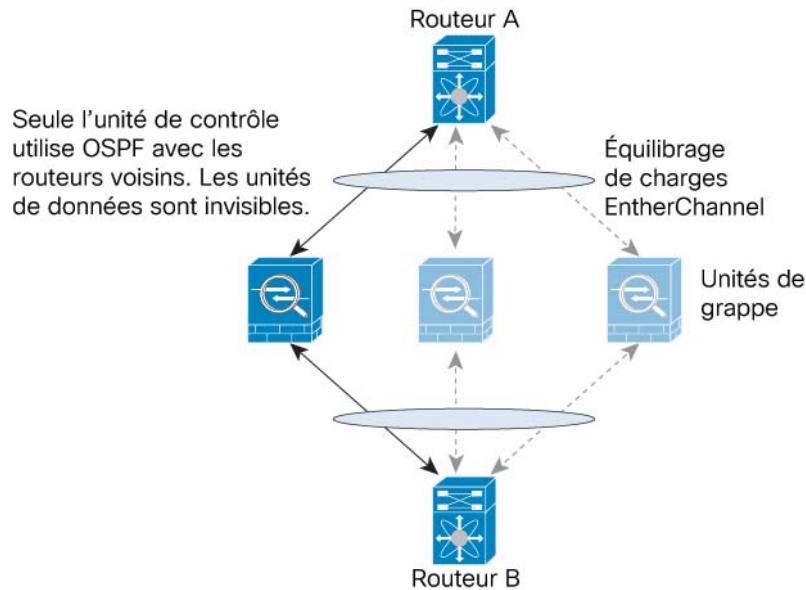
- PAT avec attribution de bloc de ports : Consultez les consignes suivantes pour cette fonctionnalité :
 - La limite maximale par hôte n'est pas une limite à l'échelle de la grappe et s'applique à chaque nœud individuellement. Ainsi, dans une grappe à 3 nœuds avec la limite maximale par hôte configurée à 1, si le trafic d'un hôte est équilibré en charge sur les 3 nœuds, 3 blocs avec 1 dans chaque nœud peuvent lui être alloués.
 - Les blocs de ports créés sur le nœud de sauvegarde à partir des pools de sauvegarde ne sont pas pris en compte lors de l'application de la limite maximale par hôte.
 - Les modifications des règles PAT à la volée, où l'ensemble PAT est modifié avec une toute nouvelle plage d'adresses IP, entraînera des échecs de création de sauvegarde xlate pour les demandes de sauvegarde xlate qui étaient encore en transit lorsque le nouveau ensemble est entré en vigueur. Ce comportement n'est pas spécifique à la fonctionnalité d'attribution de bloc de ports et il s'agit d'un problème transitoire d'ensemble PAT que l'on observe uniquement dans les déploiements en grappe où l'ensemble est distribué et le trafic est équilibré en charge sur les nœuds de la grappe.
 - Lorsque vous utilisez une grappe, vous ne pouvez pas simplement modifier la taille de l'allocation de bloc. La nouvelle taille n'est en vigueur qu'après le rechargement de chaque périphérique de la grappe. Pour éviter d'avoir à téléverser chaque périphérique, nous vous recommandons de supprimer toutes les règles d'attribution de blocage et d'effacer tous les xlates associés à ces règles. Vous pouvez ensuite modifier la taille du bloc et recréer les règles d'attribution des blocs.
- Distribution des adresses de l'ensemble NAT pour la PAT dynamique : lorsque vous configurez un ensemble PAT, la grappe divise chaque adresse IP de l'ensemble en blocs de ports. Par défaut, chaque bloc comporte 512 ports, mais si vous configurez des règles d'attribution de bloc de ports, votre paramètre de blocage est utilisé à la place. Ces blocs sont répartis uniformément entre les nœuds de la grappe, de sorte que chaque nœud ait un ou plusieurs blocs pour chaque adresse IP dans l'ensemble PAT. Ainsi, vous pourriez avoir aussi peu qu'une adresse IP dans un ensemble PAT pour une grappe, si cela est suffisant pour le nombre de connexions PAT que vous attendez. Les blocs de ports couvrent la plage de ports 1 024 à 65535, sauf si vous configurez l'option pour inclure les ports réservés, 1 à 1023, dans la règle NAT de l'ensemble PAT.
- Réutiliser un pool PAT dans plusieurs règles : pour utiliser le même pool PAT dans plusieurs règles, vous devez faire attention à la sélection d'interface dans les règles. Vous devez soit utiliser des interfaces spécifiques dans toutes les règles, soit utiliser « any » dans toutes les règles. Vous ne pouvez pas combiner des interfaces spécifiques et « any » dans les règles, sans quoi le système pourrait ne pas être en mesure de faire correspondre le trafic de retour vers le bon nœud dans la grappe. L'option la plus fiable est l'utilisation d'ensembles de PAT uniques par règle.
- Pas de tourniquet : le tourniquet pour un ensemble PAT n'est pas pris en charge avec la mise en grappe.
- Pas de PAT étendue : la PAT étendue n'est pas prise en charge avec la mise en grappe.
- Tableaux xlates dynamiques de NAT gérés par le nœud de contrôle : le nœud de contrôle conserve et reproduit le tableau xlate sur les nœuds de données. Lorsqu'un nœud de données reçoit une connexion

qui nécessite une NAT dynamique et que le xlate n'est pas dans le tableau, il le demande au nœud de contrôle. Le nœud de données est propriétaire de la connexion.

- Disques xlates périmés : le temps d'inactivité xlate sur le propriétaire de la connexion n'est pas mis à jour. Ainsi, le temps d'inactivité peut dépasser le délai d'inactivité. Une valeur de minuteur d'inactivité supérieure au délai d'expiration configuré avec un contrôle de référence de 0 est une indication d'un xlate périmé.
- Fonctionnalité PAT par session : bien qu'elle ne soit pas exclusive à la mise en grappe, la fonctionnalité PAT par session améliore l'évolutivité de PAT et, pour la mise en grappe, permet à chaque nœud de données de posséder des connexions PAT; en revanche, les connexions PAT multisessions doivent être transférées au nœud de contrôle et détenues par celui-ci. Par défaut, tout le trafic TCP et le trafic DNS UDP utilisent un xlate PAT par session, tandis que ICMP et tout le reste du trafic UDP utilise la multisession. Vous pouvez configurer des règles NAT par session pour modifier ces valeurs par défaut pour TCP et UDP, mais vous ne pouvez pas configurer la PAT par session pour ICMP. Par exemple, en raison de l'utilisation croissante du protocole Quic sur le port UDP/443 comme alternative de performance à HTTPS TLS sur TCP/443, vous devriez activer la traduction d'adresses par session (PAT) pour UDP/443. Pour le trafic qui bénéficie de la PAT multisession, comme le H.323, le SIP ou l'interface simplifiée, vous pouvez désactiver la PAT par session pour les ports TCP associés (les ports UDP de ces ports H.323 et SIP sont déjà multisession en par défaut). Référez-vous au guide de configuration du pare-feu pour obtenir plus de renseignements sur les serveurs mandataires TLS.
- Pas de PAT statique pour les inspections suivantes :
 - FTP
 - PPTP
 - RSH
 - SQLNET
 - TFTP
 - XDMCP
 - SIP
- Si vous avez un nombre extrêmement important de règles NAT, plus de dix mille, vous devez activer le modèle de validation transactionnelle à l'aide de la commande **asp rule-engine transactional-commit nat** dans l'interface de ligne de commande du périphérique. Sinon, le nœud pourrait ne pas être en mesure de rejoindre la grappe.

Routage et mise en grappe dynamiques

Le processus de routage ne s'exécute que sur l'unité de contrôle, et les routages sont appris par l'unité de contrôle et répliqués sur les serveurs secondaires. Si un paquet de routage arrive à une unité de données, il est redirigé vers l'unité de contrôle.

Illustration 1 : Routage dynamique

Une fois que les unités de données ont appris les routages de l'unité de contrôle, chaque unité prend les décisions de transfert indépendamment.

La base de données du LSA OSPF n'est pas synchronisée entre l'unité de contrôle et les unités de données. S'il y a un basculement de l'unité de contrôle, le routeur voisin détectera un redémarrage; le basculement n'est pas transparent. Le processus OSPF choisit une adresse IP comme ID de routeur. Bien que cela ne soit pas obligatoire, vous pouvez attribuer un ID de routeur statique pour vous assurer qu'un ID de routeur cohérent est utilisé dans la grappe. Consultez la fonctionnalité de transfert sans arrêt OSPF pour gérer l'interruption.

SCTP et mise en grappe

Une association SCTP peut être créée sur n'importe quel nœud (en raison de l'équilibrage de la charge); ses connexions multi-hébergement doivent résider sur le même nœud.

Inspection SIP et mise en grappe

Un flux de contrôle peut être créé sur n'importe quel nœud (en raison de l'équilibrage de la charge); ses flux de données enfants doivent résider sur le même nœud.

La configuration du mandataire TLS n'est pas prise en charge.

SNMP et mise en grappe

Un agent SNMP interroge chaque ASA en fonction de l'adresse IP locale de son . Vous ne pouvez pas interroger les données consolidées de la grappe.

Vous devez toujours utiliser l'adresse locale, et non l'adresse IP de la grappe principale pour l'interrogation SNMP. Si l'agent SNMP interroge l'adresse IP de la grappe principale, si un nouveau nœud de contrôle est choisi, l'interrogation du nouveau nœud de contrôle échouera.

Lorsque vous utilisez SNMPv3 avec la mise en grappe et que vous ajoutez un nouveau nœud de grappe après la formation initiale de la grappe, les utilisateurs SNMPv3 ne seront pas répliqués sur le nouveau nœud. Vous

devez les rajouter sur le nœud de contrôle pour forcer les utilisateurs à se reproduire vers le nouveau nœud, ou directement sur le nœud de données.

STUN et mise en grappe

L'inspection STUN est prise en charge dans les modes de basculement et de grappe, car les sténopés sont dupliqués. Cependant, l'ID de transaction n'est pas répliqué sur les nœuds. Dans le cas où un nœud tombe en panne après avoir reçu une demande STUN et qu'un autre nœud a reçu la réponse STUN, la réponse STUN sera abandonnée.

Syslog, NetFlow et la mise en grappe

- Syslog : chaque nœud de la grappe génère ses propres messages de journal système. Vous pouvez configurer la journalisation de sorte que chaque nœud utilise le même ID d'appareil ou un ID d'appareil différent dans le champ d'en-tête du message syslog. Par exemple, la configuration du nom d'hôte est répliquée et partagée par tous les nœuds de la grappe. Si vous configurez la journalisation pour utiliser le nom d'hôte comme ID d'appareil, les messages syslog générés par tous les nœuds semblent provenir d'un seul nœud. Si vous configurez la journalisation pour utiliser le nom de nœud local attribué dans la configuration de démarrage de la grappe comme ID d'appareil, les messages du journal système semblent provenir de nœuds différents.
- NetFlow : chaque nœud de la grappe génère son propre flux NetFlow. Le collecteur NetFlow peut traiter chaque appareil ASA uniquement comme un exportateur NetFlow distinct.

Cisco TrustSec et la mise en grappe

Seul le nœud de contrôle reçoit les informations des balises de groupe de sécurité (SGT). Le nœud de contrôle remplit ensuite la balise SGT pour les nœuds de données, et les nœuds de données peuvent prendre une décision de correspondance pour la balise SGT en fonction de la politique de sécurité.

VPN et mise en grappe sur le châssis Cisco Secure Firewall eXtensible Operating System (FXOS)

Une grappe ASA FXOS prend en charge l'un des deux modes mutuellement exclusifs pour le VPN de S2S, centralisé ou distribué :

- Mode VPN centralisé. Le mode par défaut. En mode centralisé, les connexions VPN sont établies avec l'unité de commande de grappe uniquement.

La fonctionnalité VPN est limitée à l'unité de contrôle et ne tire pas parti des capacités de haute disponibilité de la grappe. Si l'unité de contrôle tombe en panne, toutes les connexions VPN existantes sont perdues et les utilisateurs connectés au VPN subissent une interruption de service. Lorsqu'une nouvelle unité de contrôle est choisie, vous devez rétablir les connexions VPN.

Lorsque vous connectez un tunnel VPN à une adresse d'interface étendue, les connexions sont automatiquement transférées à l'unité de contrôle. Les clés et les certificats VPN sont répliqués sur toutes les unités.

- Mode de VPN distribué. Dans ce mode, les connexions VPN S2S IPsec IKEv2 sont réparties sur les membres d'une grappe ASA, ce qui offre une grande évolutivité. La distribution des connexions VPN entre les membres d'une grappe permet d'utiliser pleinement la capacité et le débit de la grappe, ce qui augmente considérablement la prise en charge des VPN au-delà des capacités VPN centralisées.

**Remarque**

Le mode de mise en grappe VPN centralisée prend en charge S2S IKEv1 et S2S IKEv2.

Le mode de mise en grappe du VPN distribué prend uniquement en charge S2S, IKEv2.

Le mode de mise en grappe du VPN distribué est pris en charge sur le Firepower 9300 uniquement.

Le VPN d'accès à distance n'est pas pris en charge en mode de mise en grappe de VPN centralisé ou distribué.

Facteur d'évolutivité de rendement

Lorsque vous combinez plusieurs unités dans une grappe, vous pouvez vous attendre à ce que les performances totales de la grappe atteignent environ 80 % du débit combiné maximal.

Par exemple, pour le débit TCP, le périphérique Firepower 9300 avec ses 3 modules SM-40 peut gérer environ 135 Gbit/s du trafic de pare-feu du monde réel lorsqu'il fonctionne seul. Pour le double châssis, le débit maximal combiné sera d'environ 80 % des 270 Gbit/s (2 châssis x 135 Gbit/s) : 216 Gbit/s.

Choix d'unité de contrôle

Les membres de la grappe communiquent sur le lien de commande de grappe pour élire une unité de contrôle comme suit :

1. Lorsque vous déployez la grappe, chaque unité diffuse une demande de sélection toutes les 3 secondes.
2. toute autre unité ayant un niveau de priorité plus élevée répondra à la demande de sélection; la priorité est définie lorsque vous déployez la grappe et n'est pas configurable.
3. Si, après 45 secondes, une unité ne reçoit pas de réponse d'une autre unité de priorité plus élevée, elle devient l'unité de contrôle.

**Remarque**

Si plusieurs unités sont à égalité pour la priorité la plus élevée, le nom de l'unité de la grappe suivi du numéro de série est utilisé pour déterminer l'unité de contrôle.

4. Si une unité se joint ultérieurement à la grappe avec une priorité plus élevée, elle ne devient pas automatiquement l'unité de contrôle; l'unité de contrôle existante conserve toujours ses fonctions d'unité de contrôle, sauf si elle arrête de répondre, auquel cas une nouvelle unité de contrôle est élue.
5. Dans un scénario de « split-brain » (processeur partagé), où il y a temporairement plusieurs unités de contrôle, l'unité ayant la priorité la plus élevée conserve le rôle tandis que les autres unités retournent aux rôles d'unité de données.

**Remarque**

Vous pouvez forcer manuellement une unité à devenir l'unité de contrôle. Pour les fonctions centralisées, si vous forcez le changement d'unité de contrôle, toutes les connexions sont abandonnées et vous devez rétablir les connexions sur la nouvelle unité de contrôle.

Haute disponibilité au sein de la grappe

La mise en grappe assure une disponibilité élevée en surveillant l'intégrité du châssis, des unités et de l'interface, et en reproduisant les états de connexion entre les unités.

Surveillance des applications du châssis

La surveillance de l'intégrité de l'application du châssis est toujours activée. Le superviseur Châssis Firepower 4100/9300 vérifie l'application ASA régulièrement (à chaque seconde). Si l'ASA est opérationnel et ne peut pas communiquer avec le superviseur Châssis Firepower 4100/9300 pendant 3 secondes, l'ASA génère un message syslog et quitte la grappe.

Si le superviseur Châssis Firepower 4100/9300 ne peut pas communiquer avec l'application après 45 secondes, il recharge l'ASA. Si l'ASA ne peut pas communiquer avec le superviseur, il se supprime de la grappe.

Surveillance de l'intégrité de l'unité

Chaque unité envoie périodiquement un paquet de diffusion keepaliveheartbeat sur la liaison de commande de grappe. Si le nœud de contrôle ne reçoit aucun paquet keepaliveheartbeat ou autre paquet d'un nœud de données au cours de la période d'expiration configurable, le nœud de contrôle supprime le nœud de données de la grappe. Si les nœuds de données ne reçoivent pas de paquets du nœud de contrôle, un nouveau nœud de contrôle est choisi parmi le nœud restant.

Si les nœuds ne peuvent pas se joindre sur le lien de commande de grappe en raison d'une défaillance du réseau et non parce qu'un nœud est réellement défaillant, la grappe peut entrer dans un scénario de « scission du cœur » où les nœuds de données isolés élient leurs propres nœuds de contrôle. Par exemple, si un routeur tombe en panne entre deux emplacements de grappe, le nœud de contrôle d'origine à l'emplacement 1 supprimera les nœuds de données de l'emplacement 2 de la grappe. Pendant ce temps, les nœuds de l'emplacement 2 élient leur propre nœud de contrôle et formeront leur propre grappe. Notez que le trafic symétrique peut échouer dans ce scénario. Une fois la liaison de commande de grappe restaurée, le nœud de contrôle qui a la priorité la plus élevée conservera le rôle de nœud de contrôle. Consultez [Choix d'unité de contrôle](#), à la page 86 pour de plus amples renseignements.

Surveillance d'interfaces

Chaque nœud surveille l'état de la liaison de toutes les interfaces matérielles utilisées et signale les modifications d'état au nœud de contrôle. Pour la mise en grappe sur plusieurs châssis, les EtherChannels étendus utilisent le protocole cLACP (Link Aggregation Control Protocol). Chaque châssis surveille l'état de la liaison et les messages du protocole cLACP pour déterminer si le port est toujours actif dans l'EtherChannel et informe l'application ASA si l'interface est en panne. Lorsque vous activez la surveillance de l'intégrité, les interfaces physiques sont surveillées par défaut (y compris l'interface principale EtherChannel pour les interfaces EtherChannel). Seules les interfaces nommées qui sont dans un état activé peuvent être surveillées. Par exemple, tous les ports membres d'un EtherChannel doivent tomber en panne avant qu'un EtherChannel *nommé* ne soit supprimé de la grappe (selon votre paramètre de regroupement de ports minimal). Vous pouvez éventuellement désactiver la surveillance par interface.

Si une interface surveillée tombe en panne sur un nœud particulier, mais qu'elle est active sur d'autres nœuds, ce nœud est supprimé de la grappe. Le délai avant la suppression par ASA d'un nœud de la grappe dépend du fait que le nœud est un membre établi ou qu'il rejoint la grappe. L'ASA ne surveille pas les interfaces pendant les 90 premières secondes où un nœud rejoint la grappe. Les changements d'état de l'interface pendant cette période n'entraîneront pas le retrait de ASA de la grappe. Pour un membre établi, le nœud est supprimé après 500 ms.

Pour la mise en grappe sur plusieurs châssis, si vous ajoutez ou supprimez un EtherChannel de la grappe, la surveillance de l'intégrité de l'interface est suspendue pendant 95 secondes pour que vous ayez le temps d'effectuer les modifications sur chaque châssis.

Surveillance de l'application Decorator

Lorsque vous installez une application décorateur sur une interface, comme l'application Radware DefensePro, ASA et l'application décorateur doivent être opérationnels pour rester dans la grappe. L'unité ne rejoint pas la grappe tant que les deux applications ne sont pas opérationnelles. Une fois dans la grappe, l'unité surveille l'intégrité de l'application du séparateur toutes les 3 secondes. Si l'application décorateur est en panne, l'unité est supprimée de la grappe.

État après l'échec

Lorsqu'un nœud de la grappe tombe en panne, les connexions hébergées par ce nœud sont transférées en toute transparence vers d'autres nœuds; Les renseignements d'état sur les flux de trafic sont partagés sur la liaison de commande de grappe du nœud de contrôle.

Si le nœud de contrôle échoue, un autre membre de la grappe ayant la priorité la plus élevée (numéro le plus bas) devient le nœud de contrôle.

ASA tente automatiquement de rejoindre la grappe, en fonction de l'événement d'échec.



Remarque

Lorsque ASA devient inactif et ne parvient pas à rejoindre automatiquement la grappe, toutes les interfaces de données sont fermées; Seule l'interface de gestion uniquement peut envoyer et recevoir du trafic. L'interface de gestion reste active en utilisant l'adresse IP que le nœud a reçue de l'ensemble d'adresses IP de la grappe. Cependant, si vous rechargez et que le nœud est toujours inactif dans la grappe, l'interface de gestion est désactivée. Vous devez utiliser le port de console pour toute autre configuration.

Rejoindre la grappe

Après le retrait d'un membre de la grappe, la façon dont il peut rejoindre la grappe dépend de la raison de sa suppression :

- Échec de la liaison de commande de grappe lors de la jonction : après avoir résolu le problème avec la liaison de commande de grappe, vous devez rejoindre manuellement la grappe en réactivant la mise en grappe au niveau du port de console de l'ASA en saisissant la commande **cluster group name**, puis **enable**.
- Échec de la liaison de commande de la grappe après avoir rejoint la grappe : l'ASA essaie automatiquement de la rejoindre toutes les 5 minutes, indéfiniment. Ce comportement est configurable.
- Échec de l'interface de données : l'ASA tente automatiquement de rejoindre à 5 minutes, puis à 10 minutes et enfin à 20 minutes. Si la jonction échoue après 20 minutes, l'ASA désactive la mise en grappe. Après avoir résolu le problème de l'interface de données, vous devez activer manuellement la mise en grappe au niveau du port de console de l'ASA en saisissant la commande **cluster group name**, puis **enable**. Ce comportement est configurable.
- Unité en échec : si l'unité a été supprimée de la grappe en raison d'un échec de vérification de l'intégrité de l'unité, la jonction avec la grappe dépend de la source de la défaillance. Par exemple, une panne de courant temporaire signifie que l'unité rejoindra la grappe au redémarrage, à condition que la liaison de commande de grappe soit active. L'unité tente de rejoindre la grappe toutes les 5 secondes.

- Échec de la communication Châssis-Application : lorsque l'ASA détecte que l'intégrité de l'application de châssis a été récupérée, l'ASA essaie de rejoindre la grappe immédiatement.
- Échec de l'application de décorateur : l'ASA rejoint la grappe lorsqu'il détecte que l'application de décorateur est sauvegardée.
- Erreur interne : les défaillances internes comprennent : le dépassement du délai de synchronisation de l'application, les statuts incohérents de l'application, etc. Une unité tentera de rejoindre la grappe automatiquement aux intervalles suivants : 5 minutes, 10 minutes, puis 20 minutes. Ce comportement est configurable.

Réplication de l'état de la connexion du chemin de données

Chaque connexion a un propriétaire et au moins un propriétaire secondaire dans la grappe. Le propriétaire de secours ne prend pas en charge la connexion en cas d'échec; au lieu de cela, il stocke les informations d'état TCP/UDP, de sorte que la connexion puisse être transférée de manière transparente à un nouveau propriétaire en cas de défaillance. Le propriétaire de la sauvegarde est généralement aussi le directeur.

Certains trafics nécessitent des informations d'état au-dessus de la couche TCP ou UDP. Consultez le tableau suivant pour connaître la prise en charge ou l'absence de prise en charge de ce type de trafic.

Tableau 1 : Fonctionnalités répliquées dans la grappe

Trafic	Soutien relatif à l'état	Notes
Temps de disponibilité	Oui	Assure le suivi de la disponibilité du système.
Table ARP	Oui	—
Tableau d'adresses MAC	Oui	—
Identité de l'utilisateur	Oui	Comprend les règles AAA (uauth).
Base de données du voisin IPv6	Oui	—
Routage dynamique	Oui	—
ID du moteur SNMP	Non	—
VPN distribué (site à site) pour Firepower 4100/9300	Oui	La session de sauvegarde devient la session active, puis une nouvelle session de sauvegarde est créée.

Gestion des connexions par la grappe

Les connexions peuvent être équilibrées vers plusieurs nœuds de la grappe. Les rôles de connexion déterminent la façon dont les connexions sont gérées, à la fois en fonctionnement normal et en situation de disponibilité élevée.

Rôles de connexion

Consultez les rôles suivants, définis pour chaque connexion :

- **Propriétaire** : généralement, le nœud qui reçoit initialement la connexion. Le propriétaire gère l'état TCP et traite les paquets. Une connexion n'a qu'un seul propriétaire. Si le propriétaire d'origine échoue, lorsque les nouveaux nœuds reçoivent des paquets de la connexion, le directeur choisit un nouveau propriétaire dans ces nœuds.
- **Propriétaire du sauvegarde** : nœud qui stocke les informations d'état TCP/UDP reçues du propriétaire, de sorte que la connexion puisse être transférée en toute transparence à un nouveau propriétaire en cas de défaillance. Le propriétaire de secours ne prend pas en charge la connexion en cas d'échec. Si le propriétaire devient indisponible, le premier nœud à recevoir les paquets de la connexion (selon l'équilibrage de la charge) contacte le propriétaire de secours pour obtenir les informations d'état pertinentes, afin qu'il puisse devenir le nouveau propriétaire.

Tant que le directeur (voir ci-dessous) n'est pas le même nœud que le propriétaire, le directeur est également le propriétaire secondaire. Si le propriétaire se choisit lui-même directeur, un propriétaire de secours distinct est choisi.

Pour la mise en grappe sur le périphérique Firepower 9300, qui peut inclure jusqu'à 3 nœuds de grappe dans un châssis, si le propriétaire de secours se trouve sur le même châssis que le propriétaire, un propriétaire de secours supplémentaire sera choisi dans un autre châssis pour protéger les flux d'une défaillance du châssis.

Si vous activez la localisation du directeur pour la mise en grappe inter-sites, il existe deux rôles de propriétaire de la sauvegarde : le rôle de sauvegarde locale et de sauvegarde globale. Le propriétaire choisit toujours une sauvegarde locale sur le même site que lui (en fonction de l'ID du site). La sauvegarde globale peut se trouver sur n'importe quel site et peut même se trouver sur le même nœud que la sauvegarde locale. Le propriétaire envoie des informations sur l'état de la connexion aux deux sauvegardes.

Si vous activez la redondance de sites et que le propriétaire de secours se trouve sur le même site que le propriétaire, un propriétaire de secours supplémentaire sera choisi dans un autre site pour protéger les flux d'une défaillance du site. La sauvegarde de châssis et la sauvegarde de site sont indépendantes. Par conséquent, dans certains cas, un flux comportera à la fois une sauvegarde de châssis et une sauvegarde de site.

- **Directeur** : nœud qui gère les demandes de recherche de propriétaire provenant des transitaires. Lorsque le propriétaire reçoit une nouvelle connexion, il choisit un directeur en fonction d'un hachage de l'adresse IP source/de destination et des ports (voir ci-dessous pour les détails du hachage ICMP) et envoie un message au directeur pour enregistrer la nouvelle connexion. Si les paquets arrivent à un nœud autre que le propriétaire, le nœud interroge le directeur sur quel nœud est le propriétaire afin qu'il puisse transférer les paquets. Une connexion n'a qu'un seul directeur. En cas de défaillance d'un directeur, le propriétaire en choisit un nouveau.

Tant que le directeur ne se trouve pas sur le même nœud que le propriétaire, le directeur est également le propriétaire suppléant (voir ci-dessus). Si le propriétaire se choisit lui-même directeur, un propriétaire de secours distinct est choisi.

Si vous activez la localisation de directeur pour la mise en grappe inter-sites, il y a deux rôles de directeur : le directeur local et le directeur global. Le propriétaire choisit toujours un directeur local sur le même site que lui (en fonction de l'ID du site). Le directeur global peut être sur n'importe quel site et peut même être le même nœud que le directeur local. En cas de défaillance du propriétaire d'origine, le directeur local choisit un nouveau propriétaire de connexion sur le même site.

Détails du hachage ICMP/ICMPv6 :

- Pour les paquets Echo, le port source correspond à l'identifiant ICMP et le port de destination est 0.

- Pour les paquets de réponse, le port source est 0 et le port de destination est l'identifiant ICMP.
- Pour les autres paquets, les ports source et de destination sont à 0.
- **Transitaire** : nœud qui transfère les paquets au propriétaire. Si un transitaire reçoit un paquet pour une connexion qu'il ne possède pas, il interroge le directeur à propos du propriétaire, puis établit un flux vers le propriétaire pour tout autre paquet qu'il reçoit pour cette connexion. Le directeur peut également être un transitaire. Si vous activez la localisation de directeur, le redirecteur interroge toujours le directeur local. Le transitaire n'interroge le directeur global que si le directeur local ne connaît pas le propriétaire, par exemple, si un membre de la grappe reçoit des paquets pour une connexion qui appartient à un autre site. Notez que si un transitaire reçoit le paquet SYN-ACK, il peut en dériver le propriétaire directement d'un témoin SYN dans le paquet, donc il n'a pas besoin d'interroger le directeur. (Si vous désactivez la répartition aléatoire de la séquence TCP, le témoin SYN n'est pas utilisé; une requête au directeur est requise.) Pour les flux de courte durée tels que DNS et ICMP, au lieu d'interroger, le redirecteur envoie immédiatement le paquet au directeur, qui l'envoie ensuite au propriétaire. Une connexion peut avoir plusieurs redirecteurs; le débit le plus efficace est obtenu par une bonne méthode d'équilibrage de la charge où il n'y a pas de redirecteurs et où tous les paquets d'une connexion sont reçus par le propriétaire.



Remarque

Nous vous déconseillons de désactiver la répartition aléatoire de la séquence TCP lors de l'utilisation de la mise en grappe. Il y a un faible risque que certaines sessions TCP ne soient pas établies, car le paquet SYN/ACK sera abandonné.

- **Propriétaire de fragment** : pour les paquets fragmentés, les nœuds de la grappe qui reçoivent un fragment déterminent le propriétaire du fragment à l'aide d'un hachage de l'adresse IP source du fragment, de l'adresse IP de destination et de l'ID de paquet. Tous les fragments sont ensuite transférés au propriétaire du fragment sur la liaison de commande de grappe. Les fragments peuvent être équilibrés en charge vers différents nœuds de grappe, car seul le premier fragment comprend le quintuple utilisé dans le hachage d'équilibrage de charge du commutateur. Les autres fragments ne contiennent pas les ports source et de destination et peuvent être répartis en charge vers d'autres nœuds de la grappe. Le propriétaire du fragment rassemble temporairement le paquet afin de pouvoir déterminer le directeur en fonction d'un hachage de l'adresse IP source/de destination et des ports. S'il s'agit d'une nouvelle connexion, le propriétaire du fragment s'enregistrera en tant que propriétaire de la connexion. S'il s'agit d'une connexion existante, le propriétaire du fragment transfère tous les fragments au propriétaire de la connexion fourni par la liaison de commande de grappe. Le propriétaire de la connexion rassemblera ensuite tous les fragments.

Connexions par traduction d'adresse de port

Lorsqu'une connexion utilise la traduction d'adresses de port (PAT), le type de PAT (par session ou multisession) influence quel membre de la grappe devient propriétaire de la nouvelle connexion :

- **PAT par session** : le propriétaire est le nœud qui reçoit le paquet initial dans la connexion.
Par défaut, le trafic TCP et DNS UDP utilise une PAT par session.
- **PAT multisession** : le propriétaire est toujours le nœud de contrôle. Si une connexion PAT multisession est initialement reçue par un nœud de données, le nœud de données transfère la connexion au nœud de contrôle.

Par défaut, le trafic UDP (à l'exception du trafic UDP DNS) et ICMP utilise la PAT multisession. Ces connexions appartiennent donc toujours au nœud de contrôle.

Nouvelle propriété de connexion

Exemple de flux de données pour TCP

1. SYN

2. SYN/ACK

3. SYN/ACK

4. Mise à jour de l'état

Intérieur

Extérieur

Client

Owner (responsable)

Directeur

Outil de transfert

Groupe

Serveur

Après l'étape 4, tous les paquets restants sont acheminés directement au propriétaire.

333480

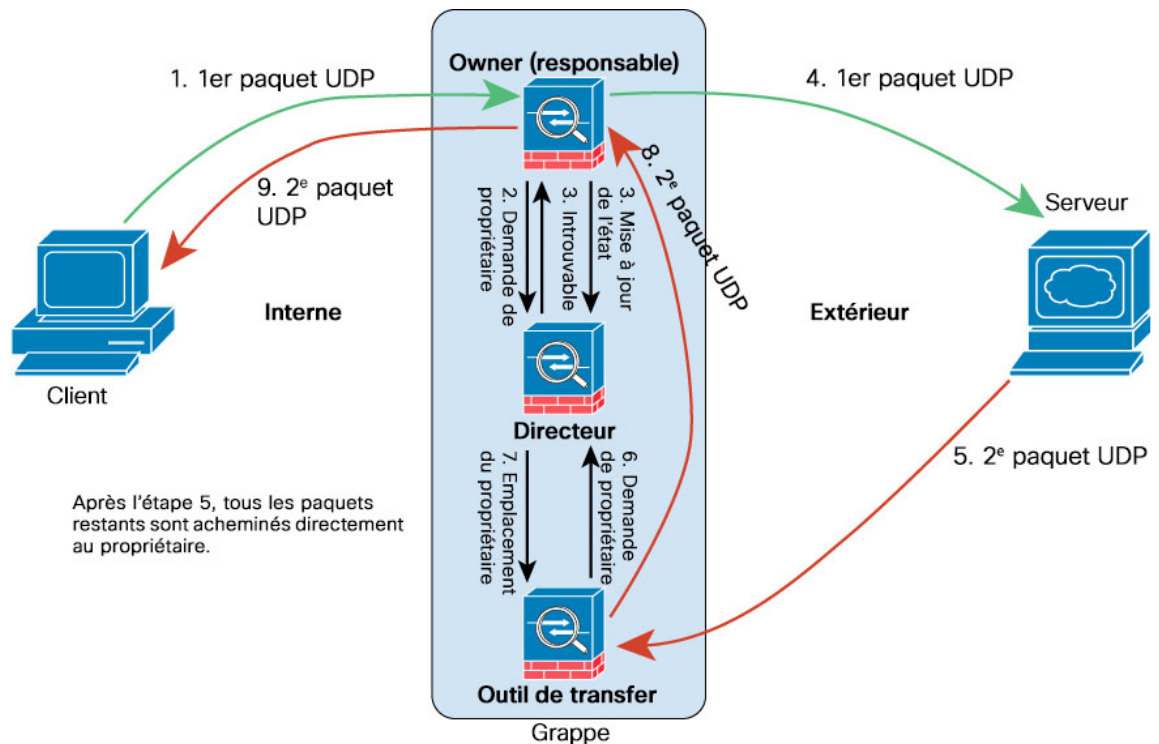
- 333480

6. Tous les paquets suivants livrés au transitaire seront transférés au propriétaire.
7. Si des paquets sont livrés à d'autres nœuds, il interrogera le directeur à propos du propriétaire et établira un flux.
8. Tout changement d'état du flux entraîne une mise à jour d'état du propriétaire au directeur.

Exemple de flux de données pour ICMP et UDP

L'exemple suivant montre l'établissement d'une nouvelle connexion.

1. Illustration 2 : Flux de données ICMP et UDP



Le premier paquet UDP provient du client et est remis à un ASA (selon la méthode d'équilibrage de la charge).

2. Le nœud qui a reçu le premier paquet interroge le nœud directeur qui est choisi en fonction d'un hachage de l'adresse IP et des ports source/de destination.
3. Le directeur ne trouve aucun flux existant, crée un flux de directeur et renvoie le paquet au nœud précédent. Autrement dit, le directeur a choisi un propriétaire pour ce flux.
4. Le propriétaire crée le flux, envoie une mise à jour d'état au directeur et transfère le paquet au serveur.
5. Le deuxième paquet UDP provient du serveur et est livré au redirecteur.
6. Le transitaire interroge le directeur pour fournir les renseignements sur la propriété. Pour les flux de courte durée tels que le DNS, au lieu d'interroger, le redirecteur envoie immédiatement le paquet au directeur, qui l'envoie ensuite au propriétaire.
7. Le directeur répond au transitaire avec les renseignements de propriété.

8. Le transitaire crée un flux de transfert pour enregistrer les informations sur le propriétaire et transfère le paquet au propriétaire.
9. Le propriétaire transfère le paquet au client.

Rééquilibrage des nouvelles connexions TCP dans la grappe

Si les capacités d'équilibrage de charges des routeurs en amont ou en aval entraînent une répartition non équilibrée des flux, vous pouvez configurer un rééquilibrage des nouvelles connexions afin que les nœuds ayant le plus de nouvelles connexions par seconde redirigent les nouveaux flux TCP vers d'autres nœuds. Aucun flux existant ne sera déplacé vers d'autres nœuds.

Comme cette commande rééquilibre uniquement en fonction des connexions par seconde, le nombre total de connexions établies sur chaque nœud n'est pas pris en compte, et le nombre total de connexions peut ne pas être égal.

Une fois qu'une connexion est déchargée sur un autre nœud, elle devient une connexion asymétrique.

Ne configurez pas le rééquilibrage de connexion pour les topologies inter-sites; il ne faut pas que les nouvelles connexions soient rééquilibrées entre les membres de la grappe sur un site différent.

Historique de la mise en grappe d'ASA sur le Firepower 4100/9300

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Suppression des propos tendancieux	9.19(1)	Les commandes, les sorties de commande et les messages syslog qui contenaient les termes « Master » (maître) et « Slave » (esclave) ont été remplacés par « Control » (contrôle) et « Data » (données). Commandes nouvelles/modifiées : cluster control-node, enable as-data-node, prompt, show cluster history, show cluster info
Amélioration de l'attribution des blocs de ports PAT pour la mise en grappe sur le Firepower 4100/9300	9.16(1)	L'allocation améliorée des blocs de ports PAT garantit que l'unité de contrôle conserve des ports en réserve pour les nœuds en cours de jonction et récupère de manière proactive les ports inutilisés. Pour optimiser au mieux l'allocation, vous pouvez définir le nombre maximum de nœuds que vous prévoyez d'avoir dans la grappe à l'aide de la commande cluster-member-limit. L'unité de contrôle peut ensuite allouer des blocs de ports au nombre de nœuds planifié sans avoir à réserver des ports pour des nœuds supplémentaires que vous ne comptez pas utiliser. La valeur par défaut est de 16 nœuds. Vous pouvez également surveiller le journal système 747046 pour vous assurer qu'il y a suffisamment de ports disponibles pour un nouveau nœud. Commandes nouvelles ou modifiées : cluster-member-limit, show nat pool cluster [summary], show nat pool ip detail
Améliorations de la commande show cluster history	9.16(1)	Nous avons ajouté des sorties supplémentaires pour la commande show cluster history. Commandes nouvelles ou modifiées : show cluster history brief, show cluster history latest, show cluster history reverse, show cluster history time

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Synchronisation de la configuration avec les unités de données en parallèle	9.14(1)	L'unité de contrôle synchronise maintenant les changements de configuration avec les unités de données en parallèle par défaut. Auparavant, la synchronisation se produisait de manière séquence. Commandes nouvelles ou modifiées : config-replicate-parallel
Messages pour un échec de jonction de grappe ou une éviction ajoutés à show cluster history	9.14(1)	De nouveaux messages ont été ajoutés à la commande show cluster history lorsqu'une unité de grappe ne parvient pas à rejoindre la grappe ou la quitter. Commandes nouvelles ou modifiées : show cluster history
Informations sur l'initiateur et le répondeur pour la détection des connexions inactives (DCD) et prise en charge du DCD dans une grappe.	9.13(1)	Si vous activez la détection des connexions inactives (DCD), vous pouvez utiliser la commande show conn detail pour obtenir des informations sur l'initiateur et le répondeur. La détection des connexions inactives vous permet de maintenir une connexion inactive, et la sortie show conn vous indique la fréquence à laquelle les points terminaux ont été sondés. En outre, DCD est désormais pris en charge dans une grappe. Commandes nouvelles ou modifiées : show conn (sortie uniquement).
Superviser la charge de trafic pour une grappe	9.13(1)	Vous pouvez désormais superviser la charge de trafic pour les membres de la grappe, y compris le nombre total de connexions, l'utilisation du CPU et de la mémoire, ainsi que les abandons de tampon. Si la charge est trop élevée, vous pouvez choisir de désactiver manuellement la mise en grappe sur l'unité si les unités restantes peuvent gérer la charge ou de régler l'équilibrage de charges sur le commutateur externe. Cette fonction est activée par défaut. Commandes nouvelles/modifiées : debug cluster load-monitor , load-monitor , show cluster info load-monitor
Jonction de grappe accélérée	9.13(1)	Lorsqu'une unité de données a la même configuration que l'unité de contrôle, elle ignorera la synchronisation de la configuration et se joindra plus rapidement. Cette fonction est activée par défaut. Cette fonctionnalité est configurée sur chaque unité et n'est pas répliquée de l'unité de contrôle à l'unité de données. Remarque Certaines commandes de configuration ne sont pas compatibles avec la jonction accélérée de grappes; si ces commandes sont présentes sur l'unité, même si la jonction accélérée de grappes est activée, la synchronisation de la configuration se produira toujours. Vous devez supprimer la configuration incompatible pour que la jonction accélérée de grappes fonctionne. Utilisez la commande show cluster info unit-join-acceleration incompatible-config pour afficher la configuration incompatible. Commandes nouvelles/modifiées : unit join-acceleration , show cluster info unit-join-acceleration incompatible-config

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Protocole ARP gratuit par site pour la mise en grappe	9.12(1)	L'ASA génère maintenant des paquets ARP gratuits (GARP) pour maintenir l'infrastructure de commutation à jour : le membre ayant la priorité la plus élevée sur chaque site génère périodiquement du trafic GARP pour les adresses MAC et IP globales. Lorsque vous utilisez des adresses MAC et IP par site, les paquets provenant de la grappe utilisent une adresse MAC et une adresse IP propres au site, tandis que les paquets reçus par la grappe utilisent une adresse MAC et une adresse IP globales. Si le trafic n'est pas généré périodiquement à partir de l'adresse MAC globale, vous pourriez rencontrer un délai d'expiration de l'adresse MAC sur vos commutateurs pour l'adresse MAC globale. Après un délai d'expiration, le trafic destiné à l'adresse MAC globale sera inondé dans l'ensemble de l'infrastructure de commutation, ce qui peut entraîner des problèmes de performance et de sécurité. Le GARP est activé par défaut lorsque vous définissez l'ID de site pour chaque unité et l'adresse MAC du site pour chaque EtherChannel étendu. Commandes nouvelles ou modifiées : site-periodic-garp interval
Groupement d'unités de grappe en parallèle par châssis Firepower 9300	9.10(1)	Pour Firepower 9300, cette fonctionnalité garantit que les modules de sécurité d'un châssis rejoignent la grappe simultanément afin que le trafic soit réparti uniformément entre les modules. Si un module se joint très avant les autres modules, il peut recevoir plus de trafic que souhaité, car les autres modules ne peuvent pas encore partager la charge. Commandes nouvelles ou modifiées : unit parallel-join
Adresse IP personnalisable par liaison de commande de grappe pour le Firepower 4100/9300	9.10(1)	Par défaut, la liaison de commande de grappe utilise le réseau 127.2.0.0/16. Vous pouvez maintenant définir le réseau lorsque vous déployez la grappe dans FXOS. Le châssis génère automatiquement l'adresse IP de l'interface de liaison de commande de grappe pour chaque unité en fonction de l'ID du châssis et de l'ID d'emplacement : 127.2.chassis_id.slot_id. Cependant, certains déploiements réseau ne permettent pas le passage du trafic 127.2.0.0/16. Par conséquent, vous pouvez maintenant définir un sous-réseau personnalisé /16 pour la liaison de commande de grappe dans FXOS, à l'exception des adresses de boucle avec retour (127.0.0.0/8) et de multidiffusion (224.0.0.0/4). Commandes FXOS nouvelles ou modifiées : set cluster-control-link network
Le délai de réponse de l'interface de grappe s'applique désormais aux interfaces qui passent d'un état inactif à un état opérationnel.	9.10(1)	Lorsqu'une mise à jour d'état d'interface se produit, l'ASA attend le nombre de millisecondes spécifié dans la commande health-check monitor-interface debounce-time ou dans l'écran Configuration > Device Management (Gestion des périphériques) > High Availability and Scalability (Haute disponibilité et évolutivité) > ASA Cluster (Grappe ASA) d'ASDM avant de marquer l'interface comme ayant échoué et que l'unité soit supprimée de la grappe. Cette fonctionnalité s'applique désormais aux interfaces qui passent d'un état inactif à un état opérationnel. Par exemple, dans le cas d'un EtherChannel qui passe d'un état inactif à un état opérationnel (par exemple, le commutateur a rechargé ou le commutateur a activé un EtherChannel), un délai de réponse plus long peut empêcher l'interface de sembler être défaillante sur une unité de la grappe juste parce qu'une autre unité de la grappe a été plus rapide à regrouper les ports. Nous n'avons pas modifié de commandes.

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Rejoindre automatiquement la grappe après une défaillance interne	9.9(2)	Auparavant, de nombreuses conditions d'erreur entraînaient le retrait d'une unité de la grappe et vous deviez rejoindre manuellement la grappe après avoir résolu le problème. Désormais, une unité tentera de rejoindre la grappe automatiquement aux intervalles suivants par défaut : 5 minutes, 10 minutes, puis 20 minutes. Ces valeurs sont configurables. Les défaillances internes comprennent : des états d'applications non uniformes; et ainsi de suite. Commandes nouvelles ou modifiées : health-check system auto-rejoin, show cluster info auto-join
Afficher les statistiques relatives au transport pour les messages du protocole de transport fiable de la grappe	9.9(2)	Vous pouvez désormais afficher l'utilisation du tampon de transport fiable par grappe d'unités afin de cerner les problèmes d'abandon de paquet lorsque le tampon est plein dans le plan de commande. Commande nouvelle ou modifiée : show cluster info transport cp detail
Le comportement de la commande cluster remove unit correspond au comportement no enable	9.9(1)	La commande cluster remove unit supprime maintenant une unité de la grappe jusqu'à ce que vous réactiviez manuellement la mise en grappe ou la rechargez, de manière similaire à la commande no enable. Auparavant, si vous redéployiez la configuration de démarrage à partir de FXOS, la mise en grappe était réactivée. Désormais, l'état désactivé persiste même dans le cas d'un redéploiement de la configuration de démarrage. Cependant, le rechargement de l'ASA réactivera la mise en grappe. Commande nouvelle ou modifiée : cluster remove unit
Amélioration de la détection des échecs de vérification de l'intégrité du châssis	9.9(1)	Vous pouvez désormais configurer un délai de rétention inférieur pour la vérification de l'intégrité du châssis : 100 ms. Le minimum précédent était de 300 ms. Notez que la durée minimum combinée (<i>intervalle x nombre de tentatives</i>) ne peut pas être inférieure à 600 ms. Commande nouvelle ou modifiée : app-agent heartbeat interval
Redondance inter-sites pour la mise en grappe	9.9(1)	La redondance inter-sites garantit qu'un propriétaire de secours pour un flux de trafic se trouvera toujours sur l'autre site que le propriétaire. Cette fonctionnalité offre une protection contre les défaillances du site. Commande nouvelle ou modifiée : site-redundancy, show asp cluster counter change, show asp table cluster chash-table, show conn flag
VPN de site à site distribué avec mise en grappe sur le Firewall 9300	9.9(1)	Une grappe ASA sur le Firewall 9300 prend en charge le VPN de site à site en mode distribué. Le mode distribué permet d'avoir de nombreuses connexions VPN IPsec IKEv2 de site à site réparties sur les membres d'une grappe ASA, pas seulement sur l'unité de contrôle (comme en mode centralisé). Cela étend considérablement la prise en charge des VPN au-delà des capacités VPN centralisées et offre une haute disponibilité. Le VPN S2S distribué s'exécute sur une grappe allant jusqu'à deux châssis, chacun contenant jusqu'à trois modules (six membres de grappe au total), chaque module prenant en charge jusqu'à 6 000 sessions actives (12 000 au total), pour un maximum d'environ 36 000 sessions actives (72 000 au total). Commandes ajoutées ou modifiées : cluster redistribute vpn-sessiondb, show cluster vpn-sessiondb, vpn mode, show cluster resource usage, show vpn-sessiondb, show connection detail, show crypto ikev2

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Détection améliorée des échecs de vérification de l'intégrité de l'unité de grappe	9.8(1)	Vous pouvez maintenant configurer un délai de rétention inférieur pour la vérification de l'intégrité de l'unité : 0,3 seconde au minimum. La valeur minimum précédente était de 0,8 seconde. Cette fonctionnalité modifie le schéma de messagerie de vérification de l'intégrité de l'unité pour passer des <i>keepalives</i> dans le plan de commande aux <i>heartbeats</i> dans le plan de données. L'utilisation de pulsations améliore la fiabilité et la réactivité de la mise en grappe en évitant les problèmes liés à la surcharge du processeur du plan de commande et aux retards de planification. Notez que la configuration d'un délai de rétention inférieur augmente l'activité de messagerie de la liaison de commande de grappe. Nous vous suggérons d'analyser votre réseau avant de configurer un faible délai de rétention. Par exemple, assurez-vous que l'envoi d'un message Ping d'une unité à une autre via la liaison de commande de grappe retourne dans le <i>holdtime/3</i>, car il y aura trois messages de pulsation au cours d'un intervalle de délai de rétention. Si vous rétrogradez votre logiciel ASA après avoir défini le délai de rétention sur 0,3–0,7, ce paramètre reviendra à la valeur par défaut de 3 secondes, car le nouveau paramètre n'est pas pris en charge. Nous avons modifié les commandes suivantes : health-check holdtime, show asp drop cluster counter, show cluster info health details
Délai anti-rebond configurable pour marquer une interface comme défaillante pour l'Châssis Firepower 4100/9300	9.8(1)	Vous pouvez maintenant configurer le délai antirebond avant que l'ASA considère une interface comme défaillante et que l'unité ne soit supprimée de la grappe. Cette fonctionnalité permet une détection plus rapide des défaillances d'interface. Notez que la configuration d'un délai antirebond inférieur augmente les risques de faux positifs. Lorsqu'une mise à jour d'état d'interface se produit, l'ASA attend le nombre de millisecondes spécifié avant de marquer l'interface comme défaillante, et l'unité est supprimée de la grappe. Le délai antirebond par défaut est de 500 ms, avec une plage de 300 ms à 9 secondes. Commande nouvelle ou modifiée : health-check monitor-interface debounce-time
Amélioration de la mise en grappe inter-sites pour l'ASA sur le Châssis Firepower 4100/9300	9.7(1)	Vous pouvez maintenant configurer l'ID de site pour chaque Châssis Firepower 4100/9300 lorsque vous déployez la grappe ASA. Auparavant, vous deviez configurer l'ID de site dans l'application ASA; cette nouvelle fonctionnalité facilite le déploiement initial. Notez que vous ne pouvez plus définir l'ID de site dans la configuration ASA. En outre, pour une meilleure compatibilité avec la mise en grappe inter-sites, nous vous recommandons de mettre à niveau vers ASA 9.7(1) et FXOS 2.1.1, ce qui comprend plusieurs améliorations de la stabilité et des performances. Nous avons modifié la commande suivante : site-id
Localisation de directeur : amélioration de la mise en grappe inter-sites pour les centres de données	9.7(1)	Pour améliorer les performances et conserver le trafic sur un site de mise en grappe inter-sites pour les centres de données, vous pouvez activer la localisation de directeur. Les nouvelles connexions sont généralement équilibrées et appartiennent aux membres de la grappe d'un site donné. Cependant, l'ASA attribue le rôle de directeur à un membre sur <i>n'importe quel</i> site. La localisation du directeur permet des rôles de directeur supplémentaires : un directeur local sur le même site que le propriétaire et un directeur global qui peut se trouver sur n'importe quel site. Le fait de maintenir le propriétaire et le directeur sur le même site améliore les performances. En cas de défaillance du propriétaire d'origine, le directeur local choisit un nouveau propriétaire de connexion sur le même site. Le directeur global est utilisé si un membre d'une grappe reçoit des paquets pour une connexion appartenant à un autre site. Nous avons introduit ou modifié les commandes suivantes : director-localization, show asp table cluster chash, show conn, show conn detail

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Prise en charge de 16 châssis pour le Firepower 4100	9.6(2)	Vous pouvez maintenant ajouter jusqu'à 16 châssis à la grappe pour le Firepower 4100. Nous n'avons pas modifié de commandes.
Prise en charge de Firepower 4100	9.6(1)	Avec FXOS 1.1.4, l'ASA prend en charge la mise en grappe inter-châssis sur le Firepower 4100 pour un maximum de 6 châssis. Nous n'avons pas modifié de commandes.
Prise en charge des adresses IP spécifiques au site en mode routé avec EtherChannel étendu	9.6(1)	Pour la mise en grappe inter-sites en mode routé avec les canaux EtherChannels étendus, vous pouvez désormais configurer l'adresse IP spécifique au site en plus des adresses MAC spécifiques au site. L'ajout d'adresses IP de site vous permet d'utiliser l'inspection ARP sur les périphériques Overlay Transport Virtualization (OTV) pour empêcher les réponses ARP de l'adresse MAC globale de circuler sur l'interconnexion de centre de données (DCI), ce qui peut causer des problèmes de routage. L'inspection ARP est requise pour certains commutateurs qui ne peuvent pas utiliser de VACL pour filtrer les adresses MAC. Nous avons modifié les commandes suivantes : mac-address, show interface
Mise en grappe inter-châssis pour 16 modules et mise en grappe inter-sites pour l'application ASA Firepower 9300	9.5(2.1)	Grâce à FXOS 1.1.3, vous pouvez désormais activer la mise en grappe inter-châssis et par extension inter-sites. Vous pouvez inclure jusqu'à 16 modules. Par exemple, vous pouvez utiliser 1 module dans 16 châssis, ou 2 modules dans 8 châssis, ou toute combinaison qui fournit un maximum de 16 modules. Nous n'avons pas modifié de commandes.
Adresses MAC spécifiques au site pour la prise en charge de la mise en grappe inter-sites pour EtherChannel étendu en mode de pare-feu routé	9.5(2)	Vous pouvez désormais utiliser la mise en grappe inter-sites pour les EtherChannels étendus en mode routé. Pour éviter le basculement des adresses MAC, configurez un ID de site pour chaque membre d'une grappe afin qu'une adresse MAC propre au site pour chaque interface puisse être partagée entre les unités d'un site. Nous avons introduit ou modifié les commandes suivantes : site-id, mac-address site-id, show cluster info, show interface
Personnalisation de la grappe ASA du comportement de reconnexion automatique en cas de défaillance d'une interface ou de la liaison de commande de grappe	9.5(2)	Vous pouvez désormais personnaliser le comportement de reconnexion automatique en cas de défaillance d'une interface ou de la liaison de commande de grappe. Nous avons introduit la commande suivante : health-check auto-rejoin
La grappe ASA prend en charge GTPv1 et GTPv2	9.5(2)	La grappe ASA prend désormais en charge l'inspection GTPv1 et GTPv2. Nous n'avons pas modifié de commandes.
Délai de réplication de la grappe pour les connexions TCP	9.5(2)	Cette fonctionnalité aide à éliminer le « travail inutile » lié aux flux de courte durée en retardant la création du flux directeur/de sauvegarde. Nous avons introduit la commande suivante : cluster replication delay

Nom de la caractéristique	Version	Renseignements sur les fonctionnalités
Inspection LISP pour la mobilité des flux inter-sites	9.5(2)	L'architecture Cisco Locator/Protocole de séparation d'ID (LISP) sépare l'identité de l'appareil de son emplacement dans deux espaces de numérotation différents, ce qui rend la migration du serveur transparente pour les clients. L'ASA peut inspecter le trafic LISP pour détecter les changements d'emplacement, puis utiliser ces renseignements pour un fonctionnement de mise en grappe transparent; les membres de la grappe ASA inspectent le trafic LISP passant entre le routeur de premier saut et le routeur de tunnel de sortie (ETR) ou le routeur de tunnel de sortie (ITR), puis changent le propriétaire du flux pour qu'il se trouve sur le nouveau site. Nous avons introduit ou modifié les commandes suivantes : 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
Les améliorations de la NAT de classe opérateur sont désormais prises en charge dans le basculement et la mise en grappe d'ASA	9.5(2)	Pour une PAT de niveau fournisseur de services ou à grande échelle, vous pouvez attribuer un bloc de ports pour chaque hôte, au lieu de demander à la NAT d'attribuer une traduction de port à la fois (Voir RFC 6888). Cette fonctionnalité est désormais prise en charge dans les déploiements de basculement et de grappe ASA. Nous avons modifié la commande suivante : show local-host
Niveau configurable pour la mise en grappe des entrées de trace	9.5(2)	Par défaut, tous les niveaux d'événements de mise en grappe sont inclus dans le tampon de trace, y compris de nombreux événements de bas niveau. Pour limiter la trace aux événements de niveau supérieur, vous pouvez définir le niveau de trace minimum pour la grappe. Nous avons introduit la commande suivante : trace-level
Mise en grappe d'ASA intra-châssis pour le Firepower 9300	9.4(1.150)	Vous pouvez mettre en grappe jusqu'à 3 modules de sécurité dans le châssis Firepower 9300. Tous les modules dans le châssis doivent appartenir à la grappe. Nous avons introduit les commandes suivantes : cluster replication delay, debug service-module, management-only individual, show cluster chassis

À propos de la traduction

Cisco peut fournir des traductions du présent contenu dans la langue locale pour certains endroits. Veuillez noter que des traductions sont fournies à titre informatif seulement et, en cas d'incohérence, la version anglaise du présent contenu prévaudra.