



Grappe ASA pour le Secure Firewall 3100

La mise en grappe vous permet de regrouper plusieurs ASA en un seul périphérique logique. 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 d'ASA, à la page 1](#)
- [Licences pour la mise en grappe d'ASA, à la page 5](#)
- [Exigences et conditions préalables à la mise en grappe ASA, à la page 6](#)
- [Lignes directrices de la mise en grappe, à la page 8](#)
- [Configurer la mise en grappe ASA, à la page 13](#)
- [Gérer les nœuds de la grappe, à la page 47](#)
- [Supervision de la grappe ASA, à la page 52](#)
- [Exemples de mise en grappe ASA, à la page 64](#)
- [Référence pour la mise en grappe, à la page 77](#)
- [Historique de la mise en grappe d'ASA pour le Secure Firewall , à la page 93](#)

À propos de la mise en grappe d'ASA

Cette section décrit l'architecture de mise en grappe et son fonctionnement.

Intégration de la grappe dans votre réseau

La grappe se compose de plusieurs pare-feu agissant comme une seule unité. Pour agir comme une grappe, les pare-feu ont besoin de l'infrastructure suivante :

- Réseau de fond de panier isolé à grande vitesse pour la communication intra-grappe, connu sous le nom de *liaison de commande de grappe*.
- Accès de gestion à chaque pare-feu pour la configuration et la surveillance.

Lorsque vous placez la grappe dans votre réseau, les routeurs en amont et en aval doivent être en mesure d'équilibrer la charge des données entrant et sortant de la grappe à l'aide d'EtherChannels étendus. Les interfaces de plusieurs membres de la grappe sont regroupées dans un seul EtherChannel; l'EtherChannel effectue l'équilibrage de la charge entre les unités.

Membres de la grappe

Les membres de la grappe collaborent pour partager la politique de sécurité et les flux de trafic. Cette section décrit la nature de chaque rôle de membre.

Configuration du démarrage

Sur chaque appareil, vous configurez une configuration de démarrage minimale qui inclut le nom de la grappe, l'interface de la liaison de commande de grappe et les autres paramètres de la grappe. Le premier nœud sur lequel vous activez la mise en grappe devient généralement le nœud de *contrôle*. Lorsque vous activez la mise en grappe sur les nœuds suivants, ils rejoignent la grappe en tant que nœuds de *données*.

Rôles des nœuds de contrôle et de données

Un membre de la grappe est le nœud de contrôle. Si plusieurs nœuds de la grappe sont mis en ligne en même temps, le nœud de contrôle est déterminé par le paramètre de priorité dans la configuration de démarrage. la priorité est réglée entre 1 et 100, 1 étant la priorité la plus élevée. Tous les autres membres sont des nœuds de données. En règle générale, lorsque vous créez une grappe pour la première fois, le premier nœud que vous ajoutez devient le nœud de contrôle simplement parce que c'est le seul nœud de la grappe à ce moment-là.

Vous devez effectuer toute la configuration (à l'exception de la configuration de démarrage) sur le nœud de contrôle uniquement; la configuration est ensuite reproduite sur les nœuds de données. Dans le cas de ressources physiques, telles que les interfaces, la configuration du nœud de contrôle est reflétée sur tous les nœuds de données. Par exemple, si vous configurez Ethernet 1/2 comme interface interne et Ethernet 1/1 comme interface externe, ces interfaces sont également utilisées sur les nœuds de données en tant qu'interfaces interne et externe.

Certaines fonctionnalités ne sont pas évolutives en grappe, et le nœud de contrôle gère tout le trafic pour ces fonctionnalités.

Interfaces de la grappe

Vous pouvez configurer des interfaces de données ou en tant qu'interfaces individuelles. Consultez [À propos des interfaces de grappe, à la page 14](#) pour de plus amples renseignements.



Remarque

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

Liaison de commande de grappe

Chaque unité doit dédier au moins une interface matérielle comme liaison de commande de grappe. Consultez [Liaison de commande de grappe, à la page 14](#) pour de plus amples renseignements.

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 de la grappe 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

Pour l'interface de gestion, nous vous conseillons d'utiliser l'une des interfaces de gestion dédiées. Vous pouvez configurer les interfaces de gestion en tant qu'interfaces individuelles (pour les modes routé et transparent) ou en tant qu'interface EtherChannel étendue.

Nous vous conseillons d'utiliser des interfaces individuelles pour la gestion, même si vous utilisez des EtherChannels étendus pour vos interfaces de données. Les interfaces individuelles vous permettent de vous connecter directement à chaque unité si nécessaire, tandis qu'une interface EtherChannel étendue ne permet qu'une connexion à distance à l'unité de contrôle actuelle.



Remarque vous ne pouvez pas activer le routage dynamique pour l'interface de gestion. Vous devez utiliser une route statique.

Pour une interface individuelle, l'adresse IP de la grappe principale est une adresse fixe pour la grappe qui appartient toujours à l'unité de contrôle actuelle. Pour chaque interface, 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. L'adresse IP locale est utilisée pour le routage et est également utile pour la résolution de problèmes.

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.

Pour une interface EtherChannel étendue, vous ne pouvez configurer qu'une seule adresse IP, et cette adresse IP est toujours associée à l'unité de contrôle. Vous ne pouvez pas vous connecter directement à une unité de données à l'aide de l'interface EtherChannel; nous vous conseillons de configurer l'interface de gestion en tant qu'interface individuelle afin que vous puissiez vous connecter à chaque unité. Notez que vous pouvez utiliser un EtherChannel local au périphérique pour la gestion.

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.

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 ASA, à la page 6](#)
- Lignes directrices inter-sites : [Lignes directrices de la mise en grappe, à la page 8](#)
- Configurer la mobilité des flux de grappes : [Configurer la mobilité des flux de grappes, à la page 42](#)
- Activer la localisation du directeur : [Activer la localisation du directeur, à la page 40](#)
- Activer la redondance de site : [Activer la localisation du directeur, à la page 40](#)
- Exemples inter-sites : [Exemples de mise en grappe inter-sites, à la page 74](#)

Licences pour la mise en grappe d'ASA

Smart Software Manager standard et sur site

Chaque unité requiert la licence Essentials (activée par défaut) et la même licence de chiffrement. Nous vous conseillons d'accorder une licence à chaque unité avec le serveur de licences *avant* d'activer la mise en grappe pour éviter les problèmes de non-concordance de licences et, lors de l'utilisation de la licence de cryptage renforcé, des problèmes de chiffrement de la liaison de commande de grappe.

La fonctionnalité de mise en grappe elle-même ne nécessite aucune licence. Il n'y a pas de coût supplémentaire pour la licence de contexte sur les unités de données.

La licence de cryptage renforcé est automatiquement activée pour les clients qualifiés lorsque vous appliquez le jeton d'enregistrement. Pour la licence facultative de cryptage renforcé (3DES/AES) activée dans la configuration ASA, consultez ci-dessous.

Dans la configuration de la licence ASA, la licence Essentials est toujours activée par défaut sur toutes les unités. 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 : chaque unité demande une licence Essentials au serveur.
- Contexte : seule l'unité de contrôle demande la licence de contexte au serveur. La licence Essentials comprend 2 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 appareils Secure Firewall 3100 dans la grappe. La licence Essentials comprend 2 contextes; pour 6 unités, ces licences ajoutent jusqu'à 12 contextes. Vous configurez une licence de 20 contextes supplémentaires sur l'unité de contrôle. Par conséquent, la licence de grappe agrégée comprend 32 contextes. Comme la limite de plateformes pour un châssis est de 100, la licence combinée autorise un maximum de 100 contextes; les 32 contextes sont dans la limite. Par conséquent, vous pouvez configurer jusqu'à 32 contextes sur l'unité de contrôle; chaque unité de données aura également 32 contextes par la réplication de la configuration.
- Vous avez 3 appareils Secure Firewall 3100 dans la grappe. La licence Essentials comprend 2 contextes; pour 3 unités, ces licences ajoutent jusqu'à 6 contextes. Vous configurez une licence de 100 contextes supplémentaires sur l'unité de contrôle. Par conséquent, la licence de grappe agrégée comprend 106 contextes. Comme la limite de plateformes pour une unité est de 100, la licence combinée autorise un maximum de 100 contextes; les 106 contextes sont au-dessus de la limite. Par conséquent, vous ne pouvez configurer que 100 contextes sur l'unité de contrôle; chaque unité de données aura également 100 contextes par la réplication de la configuration. Dans ce cas, vous devez uniquement configurer la licence de contexte de l'unité de contrôle pour 94 contextes.
- Cryptage renforcé (3DES) : 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. Seule l'unité de contrôle demande cette licence et toutes les unités peuvent l'utiliser en raison de l'agrégation des licences.

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 des droits d'utilisation toutes les 35 secondes 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.

Exigences et conditions préalables à la mise en grappe ASA

Exigences du modèle

- Secure Firewall 3100 : maximum de ,8 unités

Configuration matérielle et logicielle requise pour l'ASA

Pour toutes les unités d'une grappe :

- Doivent être du même modèle et disposer de la même mémoire DRAM. Vous n'avez pas besoin d'avoir la même quantité de mémoire flash.
- Doit exécuter le logiciel identique, sauf lors d'une mise à niveau d'image. La mise à niveau rapide est prise en charge. 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.
- Doivent être dans le même mode de contexte de sécurité, unique ou multiple.
- (Mode de contexte unique) Doivent utiliser le même mode de pare-feu (routé ou transparent).
- Les nouveaux membres de la grappe doivent utiliser le même paramètre de chiffrement SSL (la commande **ssl encryption**) que l'unité de contrôle pour la communication de la liaison de commande de grappe initiale avant la réplique de la configuration.

Exigences du commutateur

- Assurez-vous d'achever la configuration du commutateur avant de configurer la mise en grappe sur les ASA.
- Pour obtenir la liste des commutateurs pris en charge, consultez [Compatibilité Cisco ASA](#).

Exigences ASA

- Fournissez à chaque unité une adresse IP unique avant de les joindre au réseau de gestion.
 - Consultez le chapitre de démarrage pour en savoir plus sur la connexion à l'ASA et la définition de l'adresse IP de gestion.
 - À l'exception de l'adresse IP utilisée par l'unité de contrôle (généralement la première unité que vous ajoutez à la grappe), ces adresses IP de gestion sont uniquement destinées à un usage temporaire.
 - Après qu'une unité de données ait rejoint la grappe, la configuration de l'interface de gestion est remplacée par celle répliquée à partir de l'unité de contrôle.

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 ($2/2 \times 5$ 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 ($3/2 \times 10$ 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 ($1/2 \times 10$ 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)).

Autres exigences

Nous vous conseillons d'utiliser un serveur de terminaux pour accéder à tous les ports de console des unités membres de la grappe. Pour la configuration initiale et la gestion continue (par exemple, lorsqu'une unité tombe en panne), un serveur de terminaux est utile pour la gestion à distance.

Lignes directrices de la mise en grappe

Mode contextuel

Le mode doit correspondre sur chaque unité membre.

Mode pare-feu

Pour le mode unique, le mode de pare-feu doit correspondre sur toutes les unités.

Basculement

Le basculement n'est pas pris en charge avec la mise en grappe.

IPv6

La liaison de commande de grappe est uniquement prise en charge avec IPv4.

Commutateurs

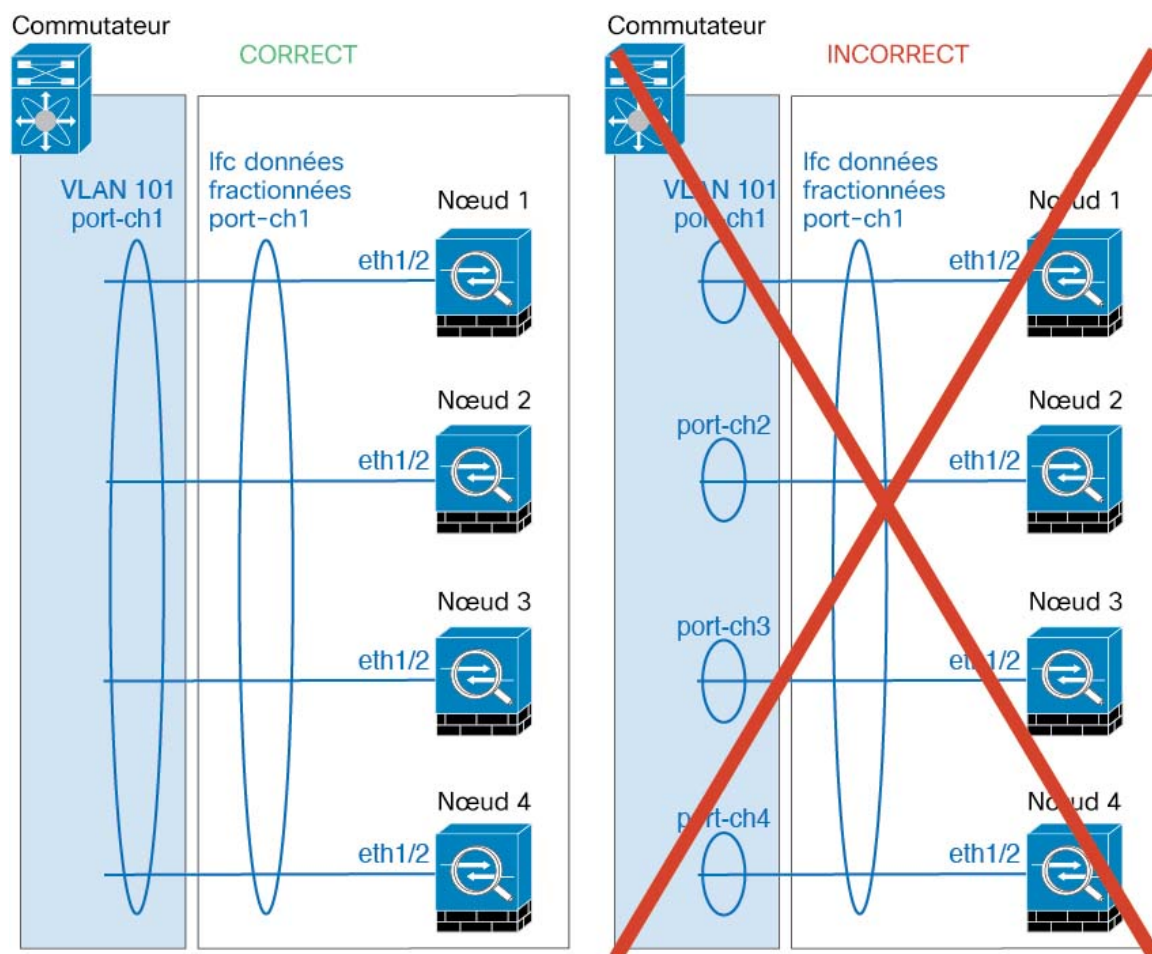
- Assurez-vous que les commutateurs connectés correspondent aux unités de transfert maximales MTU des interfaces de données et de l'interface de liaison de commande de grappe. Vous devez configurer la MTU de l'interface de la liaison de commande de grappe pour qu'elle soit au moins 100 octets supérieure 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 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.
- 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 conseillons 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.
- Vous devez désactiver la fonctionnalité de convergence progressive LACP sur toutes les interfaces EtherChannel face à la grappe pour les commutateurs Cisco Nexus.

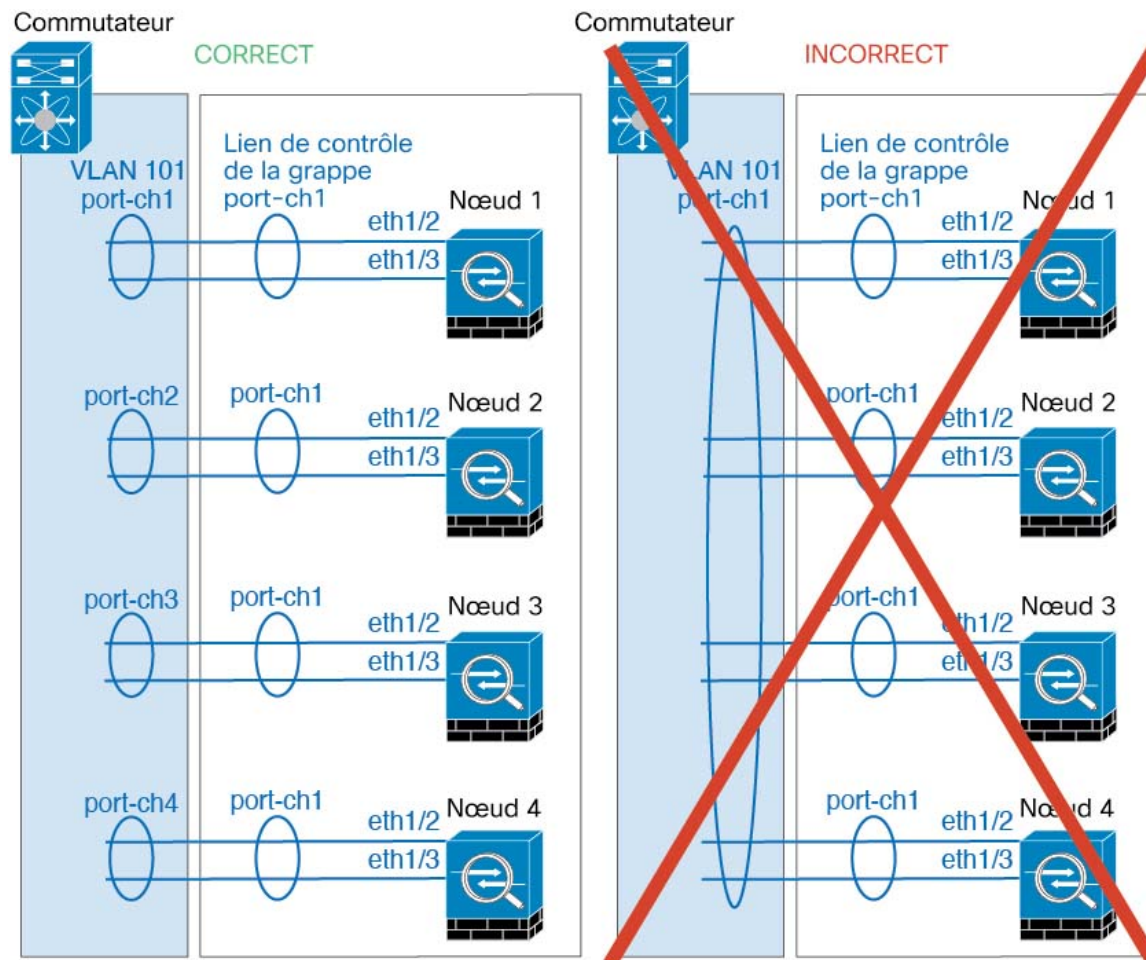
EtherChannels

- 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.



Lignes directrices 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 réhébergement.).
- 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 le pare-feu 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 bloquer votre connexion; 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 un EtherChannel étendu, lorsque le port du serveur de journal système est en panne et que le serveur ne gère 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.
- Il faut du temps pour répliquer les modifications sur toutes les unités d'une grappe. Si vous apportez une modification importante, par exemple, ajouter une règle de contrôle d'accès qui utilise des groupes d'objets (qui, lorsqu'ils sont déployés, sont divisés en plusieurs règles), le temps nécessaire pour effectuer la modification peut dépasser le délai d'expiration des unités de la grappe avec un message de réussite. Si cela se produit, vous pourriez voir le message « Échec de la réplication de la commande ». Vous pouvez ignorer le message.

Valeurs par défaut pour la mise en grappe

- Lorsque vous utilisez des EtherChannels étendus, l'ID système cLACP est généré automatiquement et la priorité du système est 1 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.
- La fonction de jonction automatique de la grappe en cas d'échec de la liaison de commande de grappe offre des tentatives illimitées toutes les 5 minutes.
- La fonction de jonction automatique de la grappe pour une interface de données défaillante effectue 3 essais toutes les 5 minutes, l'intervalle croissant étant fixé à 2.
- 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.
- Un délai de duplication de connexion de 5 secondes est activé par défaut pour le trafic HTTP.

Configurer la mise en grappe ASA

Pour configurer la mise en grappe, effectuez les tâches suivantes.

**Remarque**

Pour activer ou désactiver la mise en grappe, vous devez utiliser une connexion de console (pour l'interface de ligne de commande) ou une connexion ASDM.

Câbler les unités et configurer les interfaces

Avant de configurer la mise en grappe, câblez le réseau de liaisons de commande de grappe, le réseau de gestion et les réseaux de données. Configurez ensuite vos interfaces.

À propos des interfaces de grappe

Vous pouvez configurer des interfaces de données canaux EtherChannels étendus.

Chaque unité doit également dédier au moins une interface matérielle comme liaison de commande de grappe.

Liaison de commande de grappe

Chaque unité doit dédier au moins une interface matérielle comme liaison de commande de grappe. Nous vous recommandons d'utiliser un EtherChannel pour le lien de commande de grappe, si disponible.

Présentation du trafic de liaison de commande de grappe

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.

Interfaces et réseau de la liaison de commande de grappe

Vous pouvez utiliser n'importe quelle interface de données pour la liaison de commande de grappe, avec les exceptions suivantes :

- Vous ne pouvez pas utiliser une sous-interface VLAN comme liaison de commande de grappe.
- Vous ne pouvez pas utiliser une interface de gestion Management x/x comme liaison de commande de grappe, seule ou en tant qu'EtherChannel.

Vous pouvez utiliser un EtherChannel.

Chaque liaison de commande de grappe possède une adresse IP sur le même sous-réseau. Ce sous-réseau doit être isolé de tout autre trafic et ne doit inclure que les interfaces de liaison de commande de grappe ASA.

Pour une grappe de deux membres, ne connectez pas directement la liaison de commande de grappe d'un nœud ASA à l'autre ASA. 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.

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.



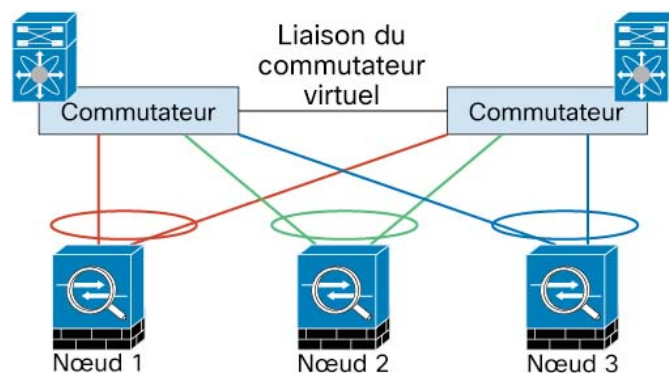
Remarque

Si votre grappe génère un trafic asymétrique (rééquilibreur) 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 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.

Échec de la liaison de commande de grappe

Si le protocole de liaison de commande de grappe tombe en panne pour une unité, la mise en grappe est désactivée; les interfaces de données sont fermées. Après avoir corrigé la liaison de commande de grappe, vous devez rejoindre manuellement la grappe en réactivant la mise en grappe.

**Remarque**

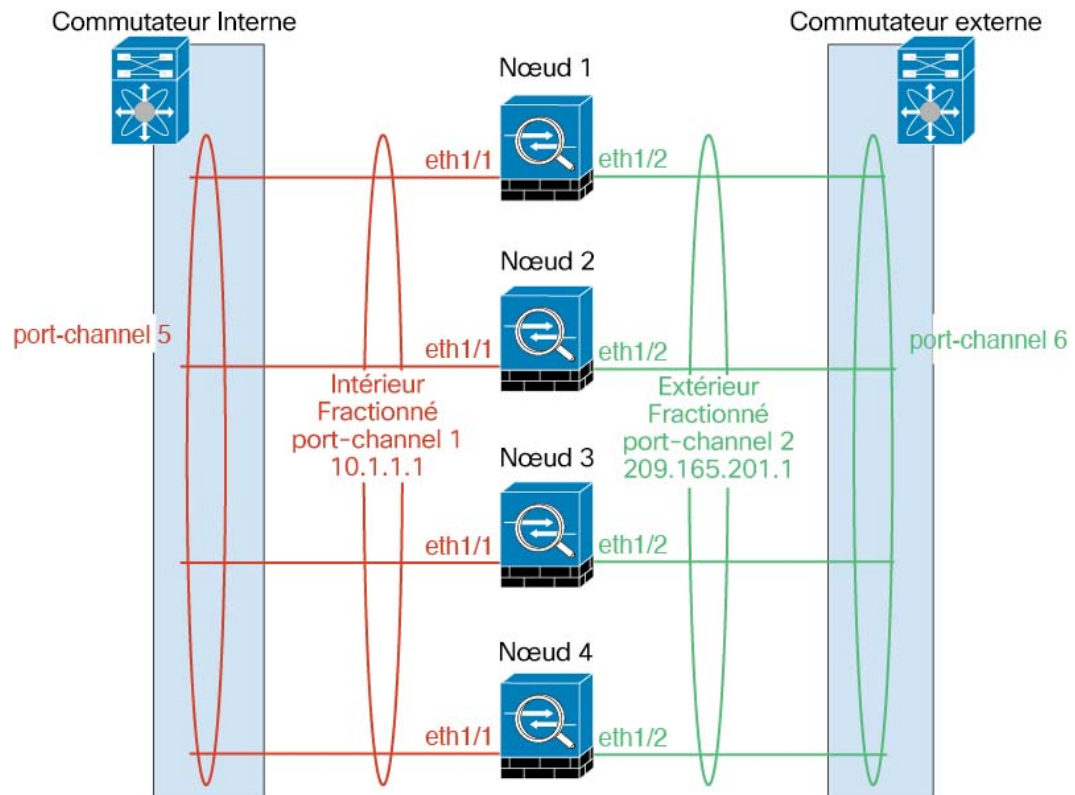
Lorsque l'ASA devient inactif, toutes les interfaces de données sont fermées; seule l'interface de gestion uniquement peut envoyer et recevoir du trafic. L'interface de gestion reste active en utilisant l'adresse IP que l'unité a reçue de l'ensemble d'adresses IP de la grappe. Cependant, si vous rechargez et que l'unité est toujours inactive dans la grappe, l'interface de gestion n'est pas accessible (car elle utilise alors l'adresse IP principale, qui est la même que l'unité de contrôle). Vous devez utiliser le port de console pour toute autre configuration.

EtherChannels étendus

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.



Lignes directrices pour le débit maximal

Pour atteindre un débit maximal, nous vous recommandons ce qui suit :

- Utilisez un algorithme de hachage pour l'équilibrage de la charge qui est « symétrique », ce qui signifie que les paquets des deux directions auront le même hachage et seront envoyés au même ASA dans l'EtherChannel étendu. Nous vous recommandons d'utiliser l'adresse IP source et de destination (par défaut) ou les ports source et destination comme algorithme de hachage.
- Utilisez le même type de cartes de ligne lors de la connexion des ASA au commutateur afin que les algorithmes de hachage appliqués à tous les paquets soient les mêmes.

Équilibrage de la charge

Le lien EtherChannel est sélectionné à l'aide d'un algorithme de hachage propriétaire, en fonction des adresses IP source ou de destination et des numéros de ports TCP et UDP.



Remarque

Sur le commutateur, nous vous recommandons d'utiliser l'un des algorithmes suivants : **source-dest-ip** ou **source-dest-ip-port** (consultez la commande **d'équilibrage de la charge du canal de port** du système d'exploitation Cisco Nexus ou Cisco IOS). N'utilisez pas le mot clé **vlan** dans l'algorithme d'équilibrage de charges, car cela pourrait entraîner une répartition inégale du trafic vers les nœuds d'une grappe.

Le nombre de liaisons dans l'EtherChannel affecte l'équilibrage de la charge.

L'équilibrage de charge symétrique n'est pas toujours possible. Si vous configurez la NAT, les paquets de transfert et de retour auront des adresses IP et/ou des ports différents. Le trafic de retour sera envoyé vers une autre unité en fonction du hachage, et la grappe devra rediriger la majeure partie du trafic de retour vers la bonne unité.

Redondance EtherChannel

L'EtherChannel a une redondance intégrée. Il surveille l'état du protocole de ligne de toutes les liaisons. Si une liaison échoue, le trafic est rééquilibré entre les liaisons restantes. Si tous les liens de l'EtherChannel échouent sur une unité particulière, mais que d'autres unités sont toujours actives, cette unité est supprimée de la grappe.

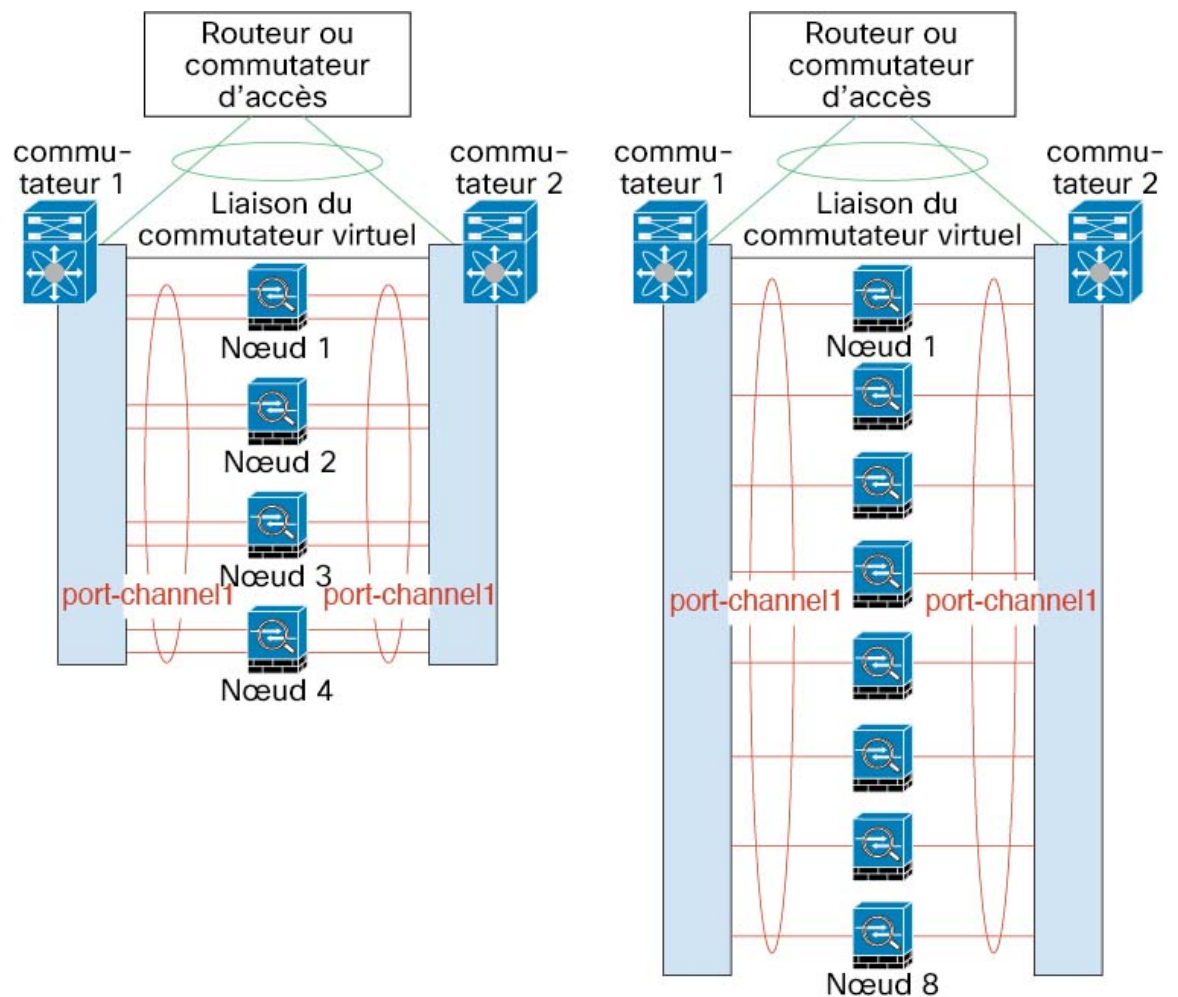
Connexion à un système de commutateurs redondants

Vous pouvez inclure plusieurs interfaces pour chaque ASA dans l'EtherChannel étendu. Plusieurs interfaces par ASA sont particulièrement utiles pour la connexion aux deux commutateurs dans un système VSS, vPC, StackWise ou StackWise Virtual.

Selon vos commutateurs, vous pouvez configurer jusqu'à 32 liens actifs dans l'EtherChannel étendu. Cette fonctionnalité nécessite que les deux commutateurs du vPC prennent en charge les canaux EtherChannels avec 16 liens actifs chacun (par exemple, le module Cisco Nexus 7000 avec le module Ethernet de 10 Gigabit de la gamme F2).

Pour les commutateurs qui prennent en charge 8 liens actifs dans l'EtherChannel, vous pouvez configurer jusqu'à 16 liens actifs dans l'EtherChannel étendu lors de la connexion à deux commutateurs dans un système redondant.

La figure suivante montre un EtherChannel de 16 liens actifs dans une grappe de 4 nœuds et une grappe de 8 nœuds.



Câbler les unités de grappe et configurer les équipements en amont et en aval

Avant de configurer la mise en grappe, câblez le réseau de liaisons de commande de grappe, le réseau de gestion et les réseaux de données.

Procédure

Étape 1 Câblez le réseau de liaisons de commande de grappe, le réseau de gestion et les réseaux de données.

Remarque

Au minimum, un réseau de liaison de commande de grappe actif est requis avant de configurer les nœuds pour qu'ils se joignent à la grappe.

Étape 2 Vous devez également configurer les équipements en amont et en aval. Par exemple, si vous utilisez des EtherChannels, vous devez configurer les équipements en amont et en aval pour les EtherChannels.

Configurer le mode d'interface de grappe de chaque unité

Avant d'activer la mise en grappe, vous devez convertir le pare-feu pour utiliser les EtherChannels étendus. Étant donné que la mise en grappe limite les types d'interfaces que vous pouvez utiliser, ce processus vous permet de vérifier votre configuration existante pour détecter les interfaces incompatibles, puis vous empêche de configurer des interfaces non prises en charge.

Avant de commencer

- Vous devez définir le mode séparément sur chaque ASA que vous souhaitez ajouter à la grappe.
- Vous pouvez toujours configurer l'interface de gestion uniquement en tant qu'interface individuelle (conseillé). L'interface de gestion peut être une interface individuelle même en mode de pare-feu transparent.
- si vous configurez l'interface de gestion en tant qu'interface individuelle, vous ne pouvez pas activer le routage dynamique pour l'interface de gestion. Vous devez utiliser une route statique.

Procédure

Étape 1 Affichez toute configuration incompatible afin de pouvoir forcer le mode d'interface et corriger votre configuration ultérieurement; le mode n'est pas modifié avec cette commande :

cluster interface-mode spanned check-details

Exemple :

```
ciscoasa(config)# cluster interface-mode spanned check-details
```

Étape 2 Définissez le mode d'interface pour la mise en grappe :

cluster interface-mode spanned force

Exemple :

```
ciscoasa(config)# cluster interface-mode spanned force
```

Il n'y a pas de paramètres par défaut; vous devez choisir explicitement le mode. Si vous n'avez pas défini le mode, vous ne pouvez pas activer la mise en grappe.

L'option **force** change le mode sans vérifier si votre configuration comporte des paramètres incompatibles. Vous devez résoudre manuellement tous les problèmes de configuration après avoir changé de mode. Étant donné que toute configuration d'interface ne peut être corrigée qu'après avoir défini le mode, nous vous conseillons d'utiliser l'option **force** afin de pouvoir au moins partir de la configuration existante. Vous pouvez réexécuter l'option **check-details** après avoir défini le mode pour obtenir plus de conseils.

Sans l'option **force**, en cas de configuration incompatible, vous êtes invité à effacer votre configuration et à la recharger, ce qui vous obligera à vous connecter au port de console pour reconfigurer votre accès de gestion. Si votre configuration est compatible (ce qui est rare), le mode est modifié et la configuration est conservée. Si vous ne souhaitez pas effacer votre configuration, vous pouvez quitter la commande en tapant **n**.

Pour supprimer le mode d'interface, saisissez la commande **no cluster interface-mode**.

Configurer les interfaces sur le nœud de contrôle

Vous devez modifier toute interface actuellement configurée avec une adresse IP pour qu'elle soit prête pour la grappe avant d'activer la mise en grappe. Pour les autres interfaces, vous pouvez les configurer avant ou après avoir activé la mise en grappe; nous vous recommandons de préconfigurer toutes vos interfaces afin que la configuration complète soit synchronisée avec les nouveaux membres de la grappe.

Cette section décrit comment configurer les interfaces pour qu'elles soient compatibles avec la mise en grappe.

Configurer l'interface de gestion en tant qu'interface individuelle

Les interfaces individuelles sont des interfaces de routage normales, chacune ayant sa propre adresse IP prise dans un ensemble d'adresses IP. L'adresse IP de la grappe principale est une adresse fixe pour la grappe qui appartient toujours au nœud de contrôle.

nous conseillons de configurer l'interface de gestion en tant qu'interface individuelle. Les interfaces de gestion individuelles vous permettent de vous connecter directement à chaque unité si nécessaire, tandis qu'une interface EtherChannel étendue permet uniquement la connexion au nœud de contrôle.

Avant de commencer

- Pour le mode de contexte multiple, effectuez cette procédure dans chaque contexte. SI vous n'êtes pas déjà en mode de configuration du contexte, saisissez la commande **changeto context name**.
- (Facultatif) Configurez l'interface en tant qu'interface EtherChannel locale du périphérique et/ou configurez les sous-interfaces.
 - Dans le cas d'un EtherChannel, celui-ci est local à l'unité et n'est pas un EtherChannel étendu.

Procédure

Étape 1 Configurez un ensemble d'adresses IP locales (IPv4 et/ou IPv6), dont une sera affectée à chaque unité de grappe pour l'interface :

(IPv4)

ip local pool *poolname first-address — last-address [mask mask]*

(IPv6)

ipv6 local pool *poolname ipv6-address/prefix-length number_of_addresses*

Exemple :

```
ciscoasa(config)# ip local pool ins 192.168.1.2-192.168.1.9
ciscoasa(config-if)# ipv6 local pool insipv6 2001:DB8:45:1002/64 8
```

Incluez au moins autant d'adresses qu'il y a d'unités dans la grappe. Si vous prévoyez d'étendre la grappe, incluez des adresses supplémentaires. L'adresse IP principale de la grappe qui appartient à l'unité principale

actuelle ne fait *pas* partie de cet ensemble; assurez-vous de réserver une adresse IP sur le même réseau pour l'adresse IP principale de la grappe.

Vous ne pouvez pas déterminer à l'avance l'adresse locale exacte attribuée à chaque unité; pour voir l'adresse utilisée sur chaque unité, saisissez la commande **show ip[v6] local pool poolname**. Chaque membre de la grappe se voit attribuer un ID de membre lorsqu'il rejoint la grappe. L'ID détermine l'adresse IP locale utilisée à partir de l'ensemble.

Étape 2 Entrez le mode de configuration d'interface :

interface *interface_id*

Exemple :

```
ciscoasa(config)# interface management 1/1
```

Étape 3 Définissez l'interface en mode gestion uniquement afin qu'elle ne transmette pas de trafic :

management-only

Par défaut, les interfaces de type de gestion sont configurées en tant que gestion uniquement. En mode transparent, cette commande est toujours activée pour une interface de type de gestion.

Étape 4 Nommez l'interface :

nameif *name*

Exemple :

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

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 5 Définissez l'adresse IP principale de la grappe et identifiez l'ensemble de grappes :

(IPv4)

ip address *ip_address* [*mask*] **cluster-pool** *poolname*

(IPv6)

ipv6 address *ipv6-address/prefix-length* **cluster-pool** *poolname*

Exemple :

```
ciscoasa(config-if)# ip address 192.168.1.1 255.255.255.0 cluster-pool ins
ciscoasa(config-if)# ipv6 address 2001:DB8:45:1002::99/64 cluster-pool insipv6
```

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. Vous pouvez configurer une adresse IPv4 et/ou une adresse IPv6.

Les configurations automatiques DHCP, PPPoE et IPv6 ne sont pas prises en charge; vous devez configurer manuellement les adresses IP. La configuration manuelle de l'adresse Link-Local n'est pas non plus prise en charge.

Étape 6 Définissez le niveau de sécurité, où *number* est un entier compris entre 0 (le plus bas) et 100 (le plus élevé) :

security-level *number*

Exemple :

```
ciscoasa(config-if)# security-level 100
```

Étape 7

Activez l'interface :

no shutdown

Exemples

L'exemple suivant configure les interfaces Ethernet 1/3 et Ethernet 1/4 en tant qu'EtherChannel local d'appareil, puis configurent l'EtherChannel en tant qu'interface individuelle :

```
ip local pool mgmt 10.1.1.2-10.1.1.9
ipv6 local pool mgmtipv6 2001:DB8:45:1002/64 8

interface ethernet 1/3
channel-group 1 mode active
no shutdown

interface ethernet 1/4
channel-group 1 mode active
no shutdown

interface port-channel 1
nameif management
ip address 10.1.1.1 255.255.255.0 cluster-pool mgmt
ipv6 address 2001:DB8:45:1002::99/64 cluster-pool mgmtipv6
security-level 100
management-only
```

Configurer des EtherChannel étendus

Un EtherChannel étendu couvre tous les ASA de la grappe et assure l'équilibrage de charges dans le cadre de l'opération EtherChannel.

Avant de commencer

- Vous devez être en mode d'interface EtherChannel étendu.
- 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. Voir [Configurer la BVI \(Bridge Virtual Interface\)](#).
- Lorsque vous utilisez des EtherChannels étendus, 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 une unité qui n'est pas une unité active dans la grappe.

Procédure

Étape 1 Spécifiez l'interface que vous souhaitez ajouter au groupe de canaux :

interface *physical_interface*

Exemple :

```
ciscoasa(config)# interface ethernet 1/1
```

L'ID *physical_interface* comprend le type, le logement et le numéro de port sous la forme type slot/port. Cette première interface du groupe de canaux détermine le type et la vitesse de toutes les autres interfaces du groupe.

Étape 2 Affectez cette interface à un EtherChannel :

channel-group *channel_id* **mode active**

Exemple :

```
ciscoasa(config-if)# channel-group 1 mode active
```

Le *channel_id* est compris entre 1 et 48. Si l'interface de canal de port pour cet ID de canal n'existe pas encore dans la configuration, elle sera ajoutée automatiquement :

interface port-channel *channel_id*

Seul le mode **actif** est pris en charge pour les EtherChannels étendus.

Étape 3 Activez l'interface :

no shutdown

Étape 4 (Facultatif) Ajoutez des interfaces supplémentaires à l'EtherChannel en répétant le processus.

Exemple :

```
ciscoasa(config)# interface ethernet 1/2
ciscoasa(config-if)# channel-group 1 mode active
ciscoasa(config-if)# no shutdown
```

Les interfaces multiples dans l'EtherChannel par unité sont utiles pour se connecter à des commutateurs dans un VSS, vPC, StackWise ou StackWise Virtual.

Étape 5 Spécifiez l'interface de canal de port :

interface port-channel *channel_id*

Exemple :

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

Cette interface a été créée automatiquement lorsque vous avez ajouté une interface au groupe de canaux.

Étape 6 (Facultatif) Si vous créez des sous-interfaces VLAN sur cet EtherChannel, 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 7

(Mode de contexte multiple) Affectez l'interface à un contexte. Saisissez ensuite :

```
changeto context name
interface port-channel channel_id
```

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 8

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 9

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 10 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 11 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*

Exemple :

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

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*.

La *mac_address* 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.

Étape 12 (Mode routé) Pour la mise en grappe inter-sites, configurez une adresse MAC et une adresse IP spécifiques au site 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.dddd 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.

Créer la configuration de démarrage

Chaque nœud de la grappe nécessite une configuration de démarrage pour rejoindre la grappe.

Configurer les paramètres de démarrage du nœud de contrôle

Chaque nœud de la grappe nécessite une configuration de démarrage pour rejoindre la grappe. En règle générale, le premier nœud que vous configurez pour rejoindre la grappe sera le nœud de contrôle. Après avoir activé la mise en grappe, après une période de sélection, la grappe choisit un nœud de contrôle. Avec un seul nœud dans la grappe au départ, ce nœud deviendra le nœud de contrôle. Les nœuds que vous ajouterez ensuite à la grappe seront des nœuds de données.

Avant de commencer

- Sauvegardez vos configurations au cas où vous souhaiteriez quitter la grappe ultérieurement et deviez restaurer votre configuration.
- Pour le mode de contexte multiple, effectuez ces procédures dans l'espace d'exécution du système. Pour passer du contexte à l'espace d'exécution du système, saisissez la commande **changeto system**.
- Vous devez utiliser le port de console pour activer ou désactiver la mise en grappe. Vous ne pouvez pas utiliser Telnet ou SSH.
- Lorsque vous ajoutez un nœud à une grappe en cours d'exécution, il se peut que vous constatiez des interruptions temporaires et limitées de paquets/connexions; c'est un comportement prévu.
- Déterminez à l'avance la taille de la liaison de commande de grappe. Consultez [Dimensionner la liaison de commande de grappe, à la page 15](#).

Procédure

Étape 1

Activez l'interface de liaison de commande de grappe avant de rejoindre la grappe.

Vous identifierez ultérieurement cette interface comme la liaison de commande de grappe lorsque vous activerez la mise en grappe.

Nous vous conseillons de combiner plusieurs interfaces de liaison de commande de grappe dans un EtherChannel si vous avez suffisamment d'interfaces. L'EtherChannel est local à l'ASA et n'est pas un EtherChannel étendu.

La configuration de l'interface de liaison de commande de grappe n'est pas répliquée du nœud de contrôle aux nœuds de données; cependant, vous devez utiliser la même configuration sur chaque nœud. Comme cette configuration n'est pas répliquée, vous devez configurer les interfaces de liaison de commande de grappe séparément sur chaque nœud.

- Vous ne pouvez pas utiliser une sous-interface VLAN comme liaison de commande de grappe.

- Vous ne pouvez pas utiliser une interface de gestion Management x/x comme liaison de commande de grappe, seule ou en tant qu'EtherChannel.

- a) Entrez le mode de configuration d'interface :

interface *interface_id*

Exemple :

```
ciscoasa(config)# interface ethernet 1/6
```

- b) (Facultatif, pour un EtherChannel) Affectez cette interface physique à un EtherChannel :

channel-group *channel_id* **mode on**

Exemple :

```
ciscoasa(config-if)# channel-group 1 mode on
```

Le *channel_id* est compris entre 1 et 48. Si l'interface de canal de port pour cet ID de canal n'existe pas encore dans la configuration, elle sera ajoutée automatiquement :

interface port-channel *channel_id*

Nous vous conseillons d'utiliser le mode activé pour les interfaces membres de la liaison de commande de grappe afin de réduire le trafic inutile sur la liaison de commande de grappe. La liaison de commande de grappe n'a pas besoin du surdébit du trafic LACP, car il s'agit d'un réseau isolé et stable. **Remarque :** nous vous recommandons de régler les EtherChannels de *données* en mode actif.

- c) Activez l'interface :

no shutdown

Il vous suffit d'activer l'interface; ne configurez pas de nom pour l'interface ni d'autres paramètres.

- d) (Pour un EtherChannel) Répétez l'opération pour chaque interface supplémentaire que vous souhaitez ajouter à l'EtherChannel :

Exemple :

```
ciscoasa(config)# interface ethernet 1/7
ciscoasa(config-if)# channel-group 1 mode on
ciscoasa(config-if)# no shutdown
```

Étape 2

Spécifiez le nœud de transfert maximal de l'interface de liaison de commande de grappe, qui doit être supérieur d'au moins 100 octets à la MTU la plus élevée des interfaces de données.

mtu cluster *bytes*

Exemple :

```
ciscoasa(config)# mtu cluster 9198
```

Définissez la MTU entre 1 400 et 9 198 octets, mais pas entre 2 561 et 8 362. En raison de la gestion du groupe de blocs, cette taille MTU n'est pas optimale pour le fonctionnement du système. Par défaut, la MTU est de 1500 octets. Nous vous suggérons de régler la MTU de liaison de commande de grappe au maximum de .

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 9198 octets, la MTU de l'interface de données la plus élevée peut s'établir à 9098, tandis que la liaison de commande de grappe peut être définie sur 9198.

Cette commande est une commande de configuration globale, mais fait également partie de la configuration de démarrage qui n'est pas répliquée entre les nœuds.

Étape 3

Nommez la grappe et passez en mode de configuration de grappe :

cluster group *name*

Exemple :

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

Le nom doit être une chaîne ASCII comptant de 1 à 38 caractères. Vous ne pouvez configurer qu'un seul groupe de grappes par nœud. Tous les membres de la grappe doivent utiliser le même nom.

Étape 4

Nommez ce membre de la grappe :

local-unit *unit_name*

Utilisez une chaîne ASCII unique de 1 à 38 caractères. Chaque nœud doit avoir un nom unique. Un nœud ayant un nom en double ne sera pas autorisé dans la grappe.

Exemple :

```
ciscoasa(cfg-cluster)# local-unit node1
```

Étape 5

Spécifiez l'interface de liaison de commande de grappe, de préférence un EtherChannel :

cluster-interface *interface_id* **ip** *ip_address mask*

Exemple :

```
ciscoasa(cfg-cluster)# cluster-interface port-channel2 ip 192.168.1.1 255.255.255.0
INFO: Non-cluster interface config is cleared on Port-Channel2
```

Les sous-interfaces et les interfaces de gestion ne sont pas autorisées.

Spécifiez une adresse IPv4 pour l'adresse IP; IPv6 n'est pas pris en charge pour cette interface. Cette interface ne peut pas avoir de **nameif** configuré.

Pour chaque nœud, spécifiez une adresse IP différente sur le même réseau.

Étape 6

Si vous utilisez la mise en grappe inter-sites, définissez l'ID de site pour ce nœud afin qu'il utilise une adresse MAC propre au site :

site-id *number*

Exemple :

```
ciscoasa(cfg-cluster)# site-id 1
```

Le *nombre* est compris entre 1 et 8.

Étape 7 Définissez la priorité de ce nœud pour les sélections de nœud de contrôle :

priority *priority_number*

Exemple :

```
ciscoasa(cfg-cluster)# priority 1
```

La priorité est comprise entre 1 et 100, 1 représentant la priorité la plus élevée.

Étape 8 (Facultatif) Définissez une clé d'authentification pour le trafic de contrôle sur la liaison de commande de grappe :

key *shared_secret*

Exemple :

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

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 commande n'influe pas sur le trafic du chemin de données, y compris sur la mise à jour de l'état de connexion et les paquets transférés, qui sont toujours envoyés en clair.

Étape 9 (Facultatif) Spécifiez manuellement l'ID système cLACP et la priorité du système :

clacp system-mac {*mac_address* | **auto**} [**system-priority** *number*]

Exemple :

```
ciscoasa(cfg-cluster)# clacp system-mac 000a.0000.aaaa
```

Lorsqu'il utilise des EtherChannels étendus, l'ASA utilise cLACP pour négocier l'EtherChannel avec le commutateur voisin. Les ASA d'une grappe collaborent à la négociation du cLACP afin qu'ils apparaissent comme un seul périphérique (virtuel) pour le commutateur. Un paramètre de négociation cLACP est un ID système, qui se présente sous la forme d'une adresse MAC. Tous les ASA de la grappe utilisent le même ID système : généré automatiquement par le nœud de contrôle (par défaut) et répliqué sur toutes les unités secondaires; ou spécifié manuellement dans cette commande sous le format *H.H.H*, où H est un chiffre hexadécimal de 16 bits. (Par exemple, l'adresse MAC 00-0A-00-00-AA-AA est saisie sous la forme 000A.0000.AAAA.) Vous souhaitez peut-être configurer manuellement l'adresse MAC à des fins de dépannage, par exemple, afin de pouvoir utiliser une adresse MAC facilement identifiée. Généralement, vous utilisez l'adresse MAC générée automatiquement.

La priorité du système, comprise entre 1 et 65 535, est utilisée pour déterminer quel nœud est chargé de prendre une décision de regroupement. Par défaut, l'ASA utilise la priorité 1, qui est la priorité la plus élevée. La priorité doit être supérieure à celle du commutateur.

Cette commande ne fait pas partie de la configuration de démarrage, et est répliquée du nœud de contrôle aux nœuds de données. Cependant, vous ne pouvez pas modifier cette valeur après avoir activé la mise en grappe.

Étape 10 Activez la mise en grappe :

enable [**noconfirm**]

Exemple :

```
ciscoasa(cfg-cluster)# enable
INFO: Clustering is not compatible with following commands:
```

```

policy-map global_policy
  class inspection_default
    inspect skinny
policy-map global_policy
  class inspection_default
    inspect sip
Would you like to remove these commands? [Y]es/[N]o:Y

INFO: Removing incompatible commands from running configuration...
Cryptochecksum (changed): f16b7fc2 a742727e e40bc0b0 cd169999
INFO: Done

```

Lorsque vous saisissez la commande **enable**, l'ASA analyse la configuration en cours d'exécution à la recherche de commandes incompatibles, en vue de cerner des fonctionnalités qui ne sont pas prises en charge avec la grappe, y compris les commandes pouvant être présentes dans la configuration par défaut. Vous êtes invité à supprimer les commandes incompatibles. Si vous répondez **No** (Non), la mise en grappe n'est pas activée. Utilisez le mot clé **noconfirm** pour contourner la confirmation et supprimer automatiquement les commandes incompatibles.

Pour le premier nœud activé, une sélection de nœud de contrôle se produit. Comme le premier nœud doit être le seul membre de la grappe jusqu'à présent, il deviendra le nœud de contrôle. N'effectuez aucune modification de configuration pendant cette période.

Pour désactiver la mise en grappe, saisissez la commande **no enable**.

Remarque

Si vous désactivez la mise en grappe, toutes les interfaces de données sont fermées et seule l'interface de gestion uniquement est active.

Exemples

Dans l'exemple suivant, une interface de gestion est configurée, un EtherChannel local à un périphérique pour la liaison de commande de grappe, puis active la mise en grappe pour l'ASA appelé « node1 », qui deviendra le nœud de contrôle parce qu'il sera ajouté en premier à la grappe :

```

ip local pool mgmt 10.1.1.2-10.1.1.9
ipv6 local pool mgmtipv6 2001:DB8::1002/32 8
interface management 1/1
  nameif management
  ip address 10.1.1.1 255.255.255.0 cluster-pool mgmt
  ipv6 address 2001:DB8::1001/32 cluster-pool mgmtipv6
  security-level 100
  management-only
  no shutdown

interface ethernet 1/6
  channel-group 1 mode on
  no shutdown

interface ethernet 1/7
  channel-group 1 mode on
  no shutdown

cluster group pod1
  local-unit node1
  cluster-interface port-channel1 ip 192.168.1.1 255.255.255.0

```

```
priority 1
key chuntheunavoidable
enable noconfirm
```

Configurer les paramètres de démarrage du nœud de données

Exécutez la procédure suivante pour configurer les nœuds de données.

Avant de commencer

- Vous devez utiliser le port de console pour activer ou désactiver la mise en grappe. Vous ne pouvez pas utiliser Telnet ou SSH.
- Sauvegardez vos configurations au cas où vous souhaiteriez quitter la grappe ultérieurement et deviez restaurer votre configuration.
- Pour le mode de contexte multiple, effectuez cette procédure dans l'espace d'exécution du système. Pour passer du contexte à l'espace d'exécution du système, saisissez la commande **changeto system**.
- Si vous avez des interfaces dans votre configuration qui n'ont pas été configurées pour la mise en grappe (par exemple, l'interface de gestion Management 1/1 de configuration par défaut), vous pouvez rejoindre la grappe en tant que nœud de données (sans possibilité de devenir le nœud de contrôle dans une sélection actuelle).
- Lorsque vous ajoutez un nœud à une grappe en cours d'exécution, il se peut que vous constatiez des interruptions temporaires et limitées de paquets/connexions; c'est un comportement prévu.

Procédure

Étape 1 Configurez la même interface de liaison de commande de grappe que celle que vous avez configurée pour le nœud de contrôle.

Exemple :

```
ciscoasa(config)# interface ethernet 1/6
ciscoasa(config-if)# channel-group 1 mode on
ciscoasa(config-if)# no shutdown
ciscoasa(config)# interface ethernet 1/7
ciscoasa(config-if)# channel-group 1 mode on
ciscoasa(config-if)# no shutdown
```

Étape 2 Spécifiez la même MTU que vous avez configurée pour le nœud de contrôle :

Exemple :

```
ciscoasa(config)# mtu cluster 9198
```

Étape 3 Saisissez le même nom de grappe que celui que vous avez configuré pour le nœud de contrôle :

Exemple :

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

Étape 4 Nommez ce membre de la grappe avec une chaîne unique :

local-unit *unit_name*

Exemple :

```
ciscoasa(cfg-cluster)# local-unit node2
```

Spécifiez une chaîne ASCII de 1 à 38 caractères.

Chaque nœud doit avoir un nom unique. Un nœud ayant un nom en double ne sera pas autorisé dans la grappe.

Étape 5 Spécifiez la même interface de liaison de commande de grappe que vous avez configurée pour le nœud de contrôle, mais spécifiez une adresse IP différente sur le même réseau pour chaque nœud :

cluster-interface *interface_id* **ip** *ip_address mask*

Exemple :

```
ciscoasa(cfg-cluster)# cluster-interface port-channel2 ip 192.168.1.2 255.255.255.0
INFO: Non-cluster interface config is cleared on Port-Channel2
```

Spécifiez une adresse IPv4 pour l'adresse IP; IPv6 n'est pas pris en charge pour cette interface. Cette interface ne peut pas avoir de **nameif** configuré.

Étape 6 Si vous utilisez la mise en grappe inter-sites, définissez l'ID de site pour ce nœud afin qu'il utilise une adresse MAC propre au site :

site-id *number*

Exemple :

```
ciscoasa(cfg-cluster)# site-id 1
```

Le **number** est compris entre 1 et 8.

Étape 7 Définissez la priorité de ce nœud pour les sélections de nœuds de contrôle, généralement à une valeur plus élevée que celle du nœud de contrôle :

priority *priority_number*

Exemple :

```
ciscoasa(cfg-cluster)# priority 2
```

Définissez la priorité entre 1 et 100, 1 représentant la priorité la plus élevée.

Étape 8 Définissez la même clé d'authentification que vous avez définie pour le nœud de contrôle :

Exemple :

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

Étape 9 Activez la mise en grappe :**enable as-data-node**

Vous pouvez éviter toute incompatibilité de configuration (principalement l'existence d'interfaces non encore configurées pour la mise en grappe) en utilisant la commande **enable as-data-node**. Cette commande garantit que le nœud de données rejoint la grappe sans possibilité de devenir le nœud de contrôle dans une sélection actuelle. Sa configuration est remplacée par celle synchronisée à partir du nœud de contrôle.

Pour désactiver la mise en grappe, saisissez la commande **no enable**.

Remarque

Si vous désactivez la mise en grappe, toutes les interfaces de données sont fermées et seule l'interface de gestion est active.

Exemples

L'exemple suivant comprend la configuration d'un nœud de données, le nœud2 :

```
interface ethernet 1/6

channel-group 1 mode on
no shutdown

interface ethernet 1/7

channel-group 1 mode on
no shutdown

cluster group pod1

local-unit node2
cluster-interface port-channel1 ip 192.168.1.2 255.255.255.0
priority 2
key chuntheunavoidable
enable as-data-node
```

Personnaliser l'opération de mise en grappe

Vous pouvez personnaliser la surveillance de l'intégrité de la mise en grappe, le délai de réplication de la connexion TCP, la mobilité des flux et d'autres optimisations.

Exécutez ces procédures sur le nœud 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 le nœud de contrôle. Pour passer du contexte à l'espace d'exécution du système, saisissez la commande **changeto system**.

Procédure

-
- Étape 1** Entrez en mode de configuration de la grappe :
- cluster group** *name*
- Étape 2** (Facultatif) Activez la réplication de la console des nœuds de données vers le nœud 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 la console, les nœuds de données envoient les messages de la console au nœud de contrôle de sorte que vous ne devez surveiller qu'un seul port de console pour la grappe.
- Étape 3** 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)
-

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

Cette procédure permet de configurer la surveillance de l'intégrité des nœuds 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, ID redondant 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*
- Exemple :**
- ```
ciscoasa(config)# cluster group test
ciscoasa(cfg-cluster)#
```
- Étape 2** Personnalisez la fonctionnalité de contrôle de l'intégrité des nœuds de la grappe.

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

Pour déterminer l'intégrité des nœuds, les nœuds de la grappe ASA envoient des messages heartbeat `keepaliveheartbeat` sur la liaison de commande de grappe aux autres nœuds. Si un nœud ne reçoit pas de messages heartbeat d'un nœud homologue pendant la période d'attente, le nœud homologue est considéré comme non réactif ou mort.

- **holdtime timeout** : détermine l'intervalle de temps entre les messages d'état heartbeat des nœuds, entre 0,3 et 45 secondes; la valeur par défaut est de 3 secondes.
- **vss-enabled** : envoie des messages heartbeat sur toutes les interfaces EtherChannel dans la liaison de commande de grappe pour s'assurer qu'au moins l'un des commutateurs peut les recevoir. Si vous configurez la liaison de commande de grappe en tant EtherChannel (conseillé) et qu'elle est connectée à une paire VSS, vPC, StackWise ou StackWise Virtual, vous devrez peut-être activer l'option **vss-enabled**. Pour certains commutateurs, lorsqu'un nœud du système redondant s'arrête ou démarre, les interfaces membres d'EtherChannel connectées à ce commutateur peuvent sembler être en ligne avec l'ASA, mais elles ne font pas passer de trafic du côté du commutateur. L'ASA peut être supprimé par erreur de la grappe si vous définissez le délai d'attente ASA selon une faible valeur (par exemple 0,8 seconde) et qu'il envoie des messages `keepalive` sur l'une de ces interfaces EtherChannel.

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 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 tous les nœuds, 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**

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 un nœud particulier, mais qu'il y a des ports actifs sous la même interface logique sur d'autres nœuds, le nœud est supprimé 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 le nœud 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** (non) de cette commande. Vous pouvez désactiver la surveillance de l'intégrité des interfaces non essentielles, par exemple, l'interface de gestion.

- **interface\_id** : désactive la surveillance de tout ID de canal de port, ID redondant 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.

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 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 tous les nœuds, vous pouvez réactiver le contrôle d'intégrité.

**Exemple :**

```
ciscoasa(cfg-cluster)# no health-check monitor-interface management1/1
```

**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. Le temps total maximum pendant lequel le nœud tente de rejoindre la grappe est limité à 14400 minutes (10 jours) à partir du moment de la dernière défaillance.
- **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 le nœud ne soit supprimé de la grappe.

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

**Exemple :**

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

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 en échec, et le nœud est supprimé de la grappe. 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 antirebond plus long peut empêcher l'interface de sembler être défaillante sur un nœud de la grappe simplement parce qu'un autre nœud de la grappe a été plus rapide à regrouper les ports.

**Étape 6** (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 le nœud si les nœuds restants 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 le nœud.

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
```

**Exemple**

L'exemple suivant configure le délai de rétention du contrôle d'intégrité à 0,3 seconde; active VSS; désactive la surveillance sur l'interface Ethernet 1/2, qui est utilisée pour la gestion; définit la reconnexion automatique pour les interfaces de données à 4 tentatives commençant à 2 minutes, en multipliant la durée de 3 fois l'intervalle précédent; et définit la reconnexion automatique pour la liaison de commande de grappe à 6 tentatives toutes les 2 minutes.

```
ciscoasa(config)# cluster group test
ciscoasa(cfg-cluster)# health-check holdtime .3 vss-enabled
ciscoasa(cfg-cluster)# no health-check monitor-interface ethernet1/2
ciscoasa(cfg-cluster)# health-check data-interface auto-rejoin 4 2 3
ciscoasa(cfg-cluster)# health-check cluster-interface auto-rejoin 6 2 1
```

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

Vous pouvez configurer le rééquilibrage des connexions. Pour en savoir plus, consultez [Rééquilibrage des nouvelles connexions TCP dans la grappe, à la page 93](#)

Activez le délai de réplication de la 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 un nœud 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 un autre nœud 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

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.

En mode de contexte multiple, configurez ce paramètre dans le contexte.

#### Étape 2

Entrez en mode de configuration de la grappe :

**cluster group name**

#### Étape 3

(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 connexions soient rééquilibrées entre les membres de la grappe sur un site différent.

## 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 pour le membre de la grappe dans la configuration de démarrage.
- 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 pour le membre de la grappe dans la configuration de démarrage.

## 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 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. 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**.

## 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 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

- Affectez chaque unité de grappe à un ID de site en fonction de [Configurer les paramètres de démarrage du nœud de contrôle, à la page 27](#) et [Configurer les paramètres de démarrage du nœud de données, à la page 32](#).

- 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 :

- a) 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.

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

```
policy-map type inspect lisp inspect_map_name
```

```
parameters
```

- c) 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.

- d) 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 :

- a) 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.

- b) 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
```

# Gérer les nœuds de la grappe

Après avoir déployé la grappe, vous pouvez modifier la configuration et gérer les nœuds de cette dernière.

## Devenir un nœud 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 undonnées

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**

## Quitter la grappe

Si vous souhaitez quitter complètement la grappe, vous devez supprimer toute la configuration de démarrage de la grappe. Étant donné que la configuration actuelle sur chaque nœud est la même (synchronisée à partir de l'unité active), quitter la grappe implique également de restaurer une configuration antérieure à la grappe à partir d'une sauvegarde ou d'effacer votre configuration et de recommencer pour éviter les conflits d'adresses IP.

**Avant de commencer**

Vous devez utiliser le port de console; lorsque vous supprimez la configuration de la grappe, toutes les interfaces sont fermées, y compris l'interface de gestion et la liaison de commande de grappe. De plus, vous ne pouvez pas activer ni désactiver la mise en grappe à partir d'une connexion distante d'interface de ligne de commande.

**Procédure**

**Étape 1** Pour un nœud de données, désactivez la mise en grappe :

**cluster group** *cluster\_name* **no enable**

**Exemple :**

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

```
ciscoasa(cfg-cluster)# no enable
```

Vous ne pouvez pas modifier la configuration lorsque la mise en grappe est activée sur un nœud de données.

**Étape 2** Effacez la configuration de la grappe :

**clear configure cluster**

L'ASA désactive toutes les interfaces, y compris l'interface de gestion et la liaison de commande de grappe.

**Étape 3** Désactivez le mode d'interface de la grappe :

**no cluster interface-mode**

Le mode n'est pas stocké dans la configuration et doit être réinitialisé manuellement.

**Étape 4** Si vous avez une configuration de sauvegarde, copiez-la dans la configuration en cours d'exécution :

**copy backup\_cfg running-config**

**Exemple :**

```
ciscoasa(config)# copy backup_cluster.cfg running-config
Source filename [backup_cluster.cfg]?
Destination filename [running-config]?
ciscoasa(config)#
```

**Étape 5** Enregistrez la configuration pour le démarrage :

**write memory**

**Étape 6** Si vous n'avez pas de configuration de sauvegarde, reconfigurez l'accès de gestion. Par exemple, assurez-vous de modifier les adresses IP de l'interface et de restaurer le bon nom d'hôte.

## Modifier le nœud 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 nœuds de la grappe ou à un nœud en particulier, procédez comme suit. L'envoi d'une commande **show** à tous les nœuds collecte toutes les sorties et les affiche sur la console du nœud actuel. 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 nœuds ou, si vous spécifiez le nom du nœud, à un nœud en particulier :

**cluster exec** [*unit node\_name*] *command*

**Exemple :**

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

Pour afficher les noms de nœud, saisissez **cluster exec unit ?** (unité d'exécution de grappe?) (pour voir tous les noms à l'exception du nœud actuel), ou saisissez la commande **show cluster info**.

## Exemples

Pour copier simultanément le même fichier de capture de tous les nœuds de la grappe vers un serveur TFTP, saisissez la commande suivante sur le nœud de contrôle :

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

Plusieurs fichiers PCAP, un pour chaque nœud, sont copiés sur le serveur TFTP. Le nom du fichier de capture de destination est automatiquement associé au nom du nœud, par exemple

capture1\_asa1.pcap, capture1\_asa2.pcap, etc. Dans cet exemple, asa1 et asa2 sont des noms de nœuds de grappe.

L'exemple de sortie suivant pour la commande **cluster exec show port-channel summary** affiche les renseignements de l'EtherChannel pour chaque nœud de la grappe :

```
ciscoasa# cluster exec show port-channel summary
control node(LOCAL):*****
Number of channel-groups in use: 2
Group Port-channel Protocol Span-cluster Ports
-----+-----+-----+-----+-----
1 Po1 LACP Yes Gi0/0 (P)
2 Po2 LACP Yes Gi0/1 (P)
slave:*****
Number of channel-groups in use: 2
Group Port-channel Protocol Span-cluster Ports
-----+-----+-----+-----+-----
1 Po1 LACP Yes Gi0/0 (P)
2 Po2 LACP Yes Gi0/1 (P)
```

## Supervision de la grappe ASA

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 [details]]**

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 nœuds et de la grappe dans son ensemble. Le mot clé **details** affiche le nombre d'échecs de messages de pulsation.

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

```
ciscoasa# show cluster info
Cluster stbu: On
 This is "C" in state DATA_NODE
 ID : 0
 Site ID : 1
 Version : 9.4(1)
 Serial No.: P3000000025
 CCL IP : 10.0.0.3
 CCL MAC : 000b.fcf8.c192
 Last join : 17:08:59 UTC Sep 26 2011
 Last leave: N/A
 Other members in the cluster:
 Unit "D" in state DATA_NODE
 ID : 1
 Site ID : 1
 Version : 9.4(1)
 Serial No.: P3000000001
 CCL IP : 10.0.0.4
 CCL MAC : 000b.fcf8.c162
```

```

 Last join : 19:13:11 UTC Sep 23 2011
 Last leave: N/A
Unit "A" in state CONTROL_NODE
 ID : 2
 Site ID : 2
 Version : 9.4(1)
 Serial No.: JAB0815R0JY
 CCL IP : 10.0.0.1
 CCL MAC : 000f.f775.541e
 Last join : 19:13:20 UTC Sep 23 2011
 Last leave: N/A
Unit "B" in state DATA_NODE
 ID : 3
 Site ID : 2
 Version : 9.4(1)
 Serial No.: P3000000191
 CCL IP : 10.0.0.2
 CCL MAC : 000b.fcf8.c61e
 Last join : 19:13:50 UTC Sep 23 2011
 Last leave: 19:13:36 UTC Sep 23 2011

```

#### • show cluster info auto-join

Indique si le nœud 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 le nœud est désactivé de façon permanente ou si le nœud 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: Control node has application down that data node 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 un nœud de grappe n'a pas pu se joindre ou pourquoi un nœud 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]**

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 les écrans suivantes pour superviser le trafic de la grappe :

- **show conn [detail], 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 quel nœud 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.

La commande **show conn detail** indique également quels flux sont soumis à la mobilité des flux.

L'exemple suivant illustre une sortie de la commande **show conn detail** :

```
ciscoasa/ASA2/data node# show conn detail
12 in use, 13 most used
Cluster stub connections: 0 in use, 46 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
 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
ESP outside: 10.1.227.1/53744 NP Identity Ifc: 10.1.226.1/30604, , flags c, idle 0s,
uptime
1m21s, timeout 30s, bytes 7544, cluster sent/rcvd bytes 0/0, owners (0,255) Traffic
received
at interface outside Locally received: 7544 (93 byte/s) Traffic received at interface
NP
Identity Ifc Locally received: 0 (0 byte/s) UDP outside: 10.1.227.1/500 NP Identity
Ifc:
10.1.226.1/500, flags -c, idle 1m22s, uptime 1m22s, timeout 2m0s, bytes 1580, cluster
sent/rcvd bytes 0/0, cluster sent/rcvd total bytes 0/0, owners (0,255) Traffic received
at
interface outside Locally received: 864 (10 byte/s) Traffic received at interface NP
Identity
Ifc Locally received: 716 (8 byte/s)
```

Pour dépanner le flux de connexion, observez d'abord les connexions sur tous les nœuds en saisissant la commande **cluster exec show conn** sur n'importe quel nœud. Recherchez les flux qui ont les indicateurs suivants : directeur (Y), sauvegarde (y) et transitaire (z). L'exemple suivant montre une connexion SSH de 172.18.124.187:22 à 192.168.103.131:44727 sur les trois ASA; l'ASA 1 a l'indicateur z indiquant qu'il s'agit d'un transitaire pour la connexion, l'ASA 3 a l'indicateur y indiquant qu'il est le directeur

de la connexion et l'ASA 2 n'a pas d'indicateurs spéciaux indiquant qu'il est le propriétaire. Dans le sens sortant, les paquets de cette connexion pénètrent dans l'interface interne de l'ASA 2 et quittent l'interface externe. Dans le sens entrant, les paquets de cette connexion pénètrent dans l'interface externe sur l'ASA 1 et l'ASA 3, sont acheminés sur la liaison de commande de grappe vers l'ASA 2, puis quittent l'interface interne sur l'ASA 2.

```
ciscoasa/ASA1/control node# cluster exec show conn
ASA1 (LOCAL):*****
18 in use, 22 most used
Cluster stub connections: 0 in use, 5 most used
TCP outside 172.18.124.187:22 inside 192.168.103.131:44727, idle 0:00:00, bytes
37240828, flags z
```

```
ASA2:*****
12 in use, 13 most used
Cluster stub connections: 0 in use, 46 most used
TCP outside 172.18.124.187:22 inside 192.168.103.131:44727, idle 0:00:00, bytes
37240828, flags UIO
```

```
ASA3:*****
10 in use, 12 most used
Cluster stub connections: 2 in use, 29 most used
TCP outside 172.18.124.187:22 inside 192.168.103.131:44727, idle 0:00:03, bytes 0,
flags Y
```

- **show cluster info [conn-distribution | packet-distribution | loadbalance | flow-mobility counters]**

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

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

La commande **show cluster info flow-mobility counters** affiche les renseignements sur le mouvement de l'EID et le mouvement du propriétaire du flux. Consultez la sortie suivante pour **show cluster info flow-mobility counters** :

```
ciscoasa# show cluster info flow-mobility counters
EID movement notification received : 4
EID movement notification processed : 4
Flow owner moving requested : 2
```

- **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 | traffic | user-identity | xlate} [options]**

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) 0xc1ce5c49
access-list 101 line 8 extended permit tcp host 192.168.1.177 host 192.168.43.192
(hitcnt=3, 0, 1, 1, 1) 0xb6f59512
access-list 101 line 9 extended permit tcp host 192.168.1.177 host 192.168.43.44
(hitcnt=0, 0, 0, 0, 0) 0xdc104200
access-list 101 line 10 extended permit tcp host 192.168.1.112 host 192.168.43.44
(hitcnt=429, 109, 107, 109, 104)
0xce4f281d
access-list 101 line 11 extended permit tcp host 192.168.1.170 host 192.168.43.238
(hitcnt=3, 1, 0, 0, 2) 0x4143a818
access-list 101 line 12 extended permit tcp host 192.168.1.170 host 192.168.43.169
(hitcnt=2, 0, 1, 0, 1) 0xb18dfea4
access-list 101 line 13 extended permit tcp host 192.168.1.170 host 192.168.43.229
(hitcnt=1, 1, 0, 0, 0) 0x21557d71
access-list 101 line 14 extended permit tcp host 192.168.1.170 host 192.168.43.106
(hitcnt=0, 0, 0, 0, 0) 0x7316e016
access-list 101 line 15 extended permit tcp host 192.168.1.170 host 192.168.43.196
(hitcnt=0, 0, 0, 0, 0) 0x013fd5b8
access-list 101 line 16 extended permit tcp host 192.168.1.170 host 192.168.43.75
(hitcnt=0, 0, 0, 0, 0) 0x2c7dba0d
```

Pour afficher le nombre agrégé des connexions utilisées pour tous les nœuds, saisissez :

```
ciscoasa# show cluster conn count
Usage Summary In Cluster:*****
200 in use (cluster-wide aggregated)
cl2 (LOCAL):*****
100 in use, 100 most used

cl1:*****
100 in use, 100 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.

## 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.

## Supervision des interfaces de grappe

Consultez les commandes suivantes pour superviser les interfaces de la grappe :

- **show cluster interface-mode**

Affiche le mode d'interface de la grappe.

- **show port-channel**

Comprend des renseignements indiquant si un canal de port est étendu.

- **show lacp cluster {system-mac | system-id}**

Affiche l'ID système et la priorité cLACP.

- **debug lacp cluster [all | ccp | misc | protocol]**

Affiche les messages de débogage pour cLACP.

- **show interface**

Affiche l'utilisation de l'adresse MAC du site lorsqu'elle est utilisée :

```
ciscoasa# show interface port-channel1.3151
Interface Port-channel1.3151 "inside", is up, line protocol is up
Hardware is EtherChannel/LACP, BW 1000 Mbps, DLY 10 usec
VLAN identifier 3151
MAC address aaaa.1111.1234, MTU 1500
Site Specific MAC address aaaa.1111.aaaa
IP address 10.3.1.1, subnet mask 255.255.255.0
Traffic Statistics for "inside":
132269 packets input, 6483425 bytes
1062 packets output, 110448 bytes
98530 packets dropped
```

## 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 | transport]**

Affiche les messages de débogage pour la mise en grappe.

- **debug cluster flow-mobility**

Affiche les événements liés à la mobilité des flux de mise en grappe.

- **debug lisp eid-notify-intercept**

Affiche les événements lorsque le message eid-notify est intercepté.

- **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_KEEPA_LIVE from 80-1 at
CONTROL_NODE
```

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 msecs,
Detected another Control_node, leave and re-join as Data_node) to FSM. Current state
CONTROL_NODE
May 23 07:27:23.113 [INFO]State machine changed from state CONTROL_NODE to DISABLED
```

## Exemples de mise en grappe ASA

Ces exemples incluent toutes les configurations ASA liées aux grappes pour les déploiements typiques.

### Exemple de configuration ASA et du commutateur

Les exemples de configuration suivants connectent les interfaces suivantes entre l'ASA et le commutateur :

| Interface ASA | Interface du commutateur |
|---------------|--------------------------|
| Ethernet 1/2  | GigabitEthernet 1/0/15   |
| Ethernet 1/3  | GigabitEthernet 1/0/16   |
| Ethernet 1/4  | GigabitEthernet 1/0/17   |
| Ethernet 1/5  | GigabitEthernet 1/0/18   |

### Configuration ASA

#### Mode d'interface sur chaque unité

```
cluster interface-mode spanned force
```

#### Configuration de démarrage de l'unité de contrôle ASA1

```
interface Ethernet1/6
 channel-group 1 mode on
 no shutdown
!
interface Ethernet1/7
 channel-group 1 mode on
 no shutdown
!
interface Port-channel1
 description Clustering Interface
!
cluster group Moya
 local-unit A
 cluster-interface Port-channel1 ip 10.0.0.1 255.255.255.0
```

```

priority 10
key emphyri0
enable noconfirm

```

### Configuration de démarrage de l'unité de données ASA2

```

interface Ethernet1/6
 channel-group 1 mode on
 no shutdown
!
interface Ethernet1/7
 channel-group 1 mode on
 no shutdown
!
interface Port-channel1
 description Clustering Interface
!
cluster group Moya
 local-unit B
 cluster-interface Port-channel1 ip 10.0.0.2 255.255.255.0
 priority 11
 key emphyri0
 enable as-data-node

```

### Configuration de l'interface de l'unité de contrôle

```

ip local pool mgmt-pool 10.53.195.231-10.53.195.232

interface Ethernet1/2
 channel-group 10 mode active
 no shutdown
!
interface Ethernet1/3
 channel-group 10 mode active
 no shutdown
!
interface Ethernet1/4
 channel-group 11 mode active
 no shutdown
!
interface Ethernet1/5
 channel-group 11 mode active
 no shutdown
!
interface Management1/1
 management-only
 nameif management
 ip address 10.53.195.230 cluster-pool mgmt-pool
 security-level 100
 no shutdown
!
interface Port-channel10
 mac-address aaaa.bbbb.cccc
 nameif inside
 security-level 100
 ip address 209.165.200.225 255.255.255.224
!
interface Port-channel11
 mac-address aaaa.dddd.cccc

```

```
nameif outside
security-level 0
ip address 209.165.201.1 255.255.255.224
```

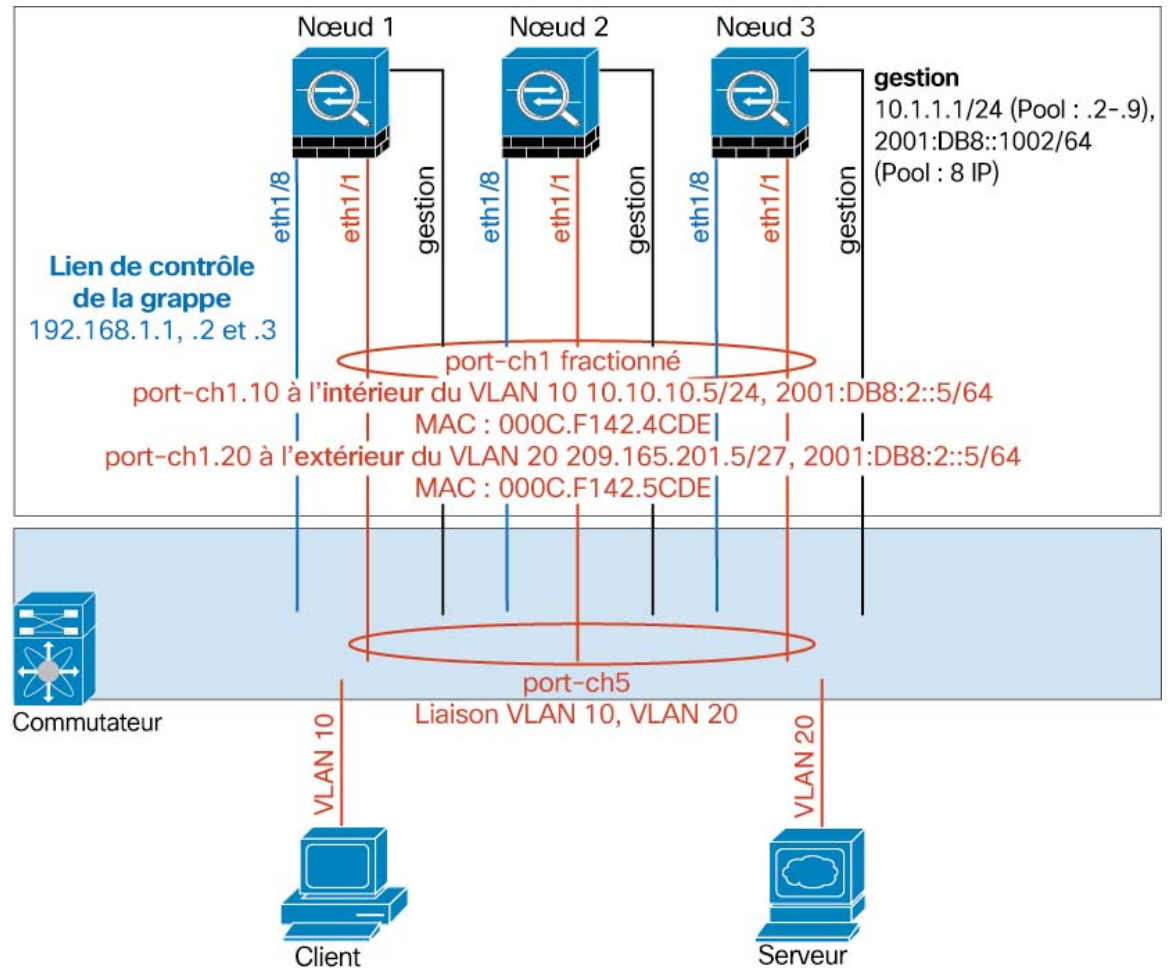
## Configuration du commutateur Cisco IOS

```
interface GigabitEthernet1/0/15
switchport access vlan 201
switchport mode access
spanning-tree portfast
channel-group 10 mode active
!
interface GigabitEthernet1/0/16
switchport access vlan 201
switchport mode access
spanning-tree portfast
channel-group 10 mode active
!
interface GigabitEthernet1/0/17
switchport access vlan 401
switchport mode access
spanning-tree portfast
channel-group 11 mode active
!
interface GigabitEthernet1/0/18
switchport access vlan 401
switchport mode access
spanning-tree portfast
channel-group 11 mode active

interface Port-channel10
switchport access vlan 201
switchport mode access

interface Port-channel11
switchport access vlan 401
switchport mode access
```

## 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.

### Mode d'interface sur chaque unité

```
cluster interface-mode spanned force
```

### Configuration de démarrage de l'unité de contrôle de l'unité 1

```
interface ethernet1/8
```

```
no shutdown
description CCL

cluster group cluster1
local-unit asa1
cluster-interface ethernet1/8 ip 192.168.1.1 255.255.255.0
priority 1
key chuntheunavoidable
enable noconfirm
```

### Configuration de démarrage de l'unité de données de l'unité 2

```
interface ethernet1/8
no shutdown
description CCL

cluster group cluster1
local-unit asa2
cluster-interface ethernet1/8 ip 192.168.1.2 255.255.255.0
priority 2
key chuntheunavoidable
enable as-data-node
```

### Configuration de démarrage de l'unité de données de l'unité 3

```
interface ethernet1/8
no shutdown
description CCL

cluster group cluster1
local-unit asa3
cluster-interface ethernet1/8 ip 192.168.1.3 255.255.255.0
priority 3
key chuntheunavoidable
enable as-data-node
```

### Configuration de l'interface de l'unité de contrôle

```
ip local pool mgmt 10.1.1.2-10.1.1.9
ipv6 local pool mgmtipv6 2001:DB8::1002/64 8

interface management 1/1
nameif management
ip address 10.1.1.1 255.255.255.0 cluster-pool mgmt
ipv6 address 2001:DB8::1001/32 cluster-pool mgmtipv6
security-level 100
management-only
no shutdown

interface ethernet1/1
channel-group 1 mode active
no shutdown

interface port-channel 1

interface port-channel 1.10
vlan 10
```

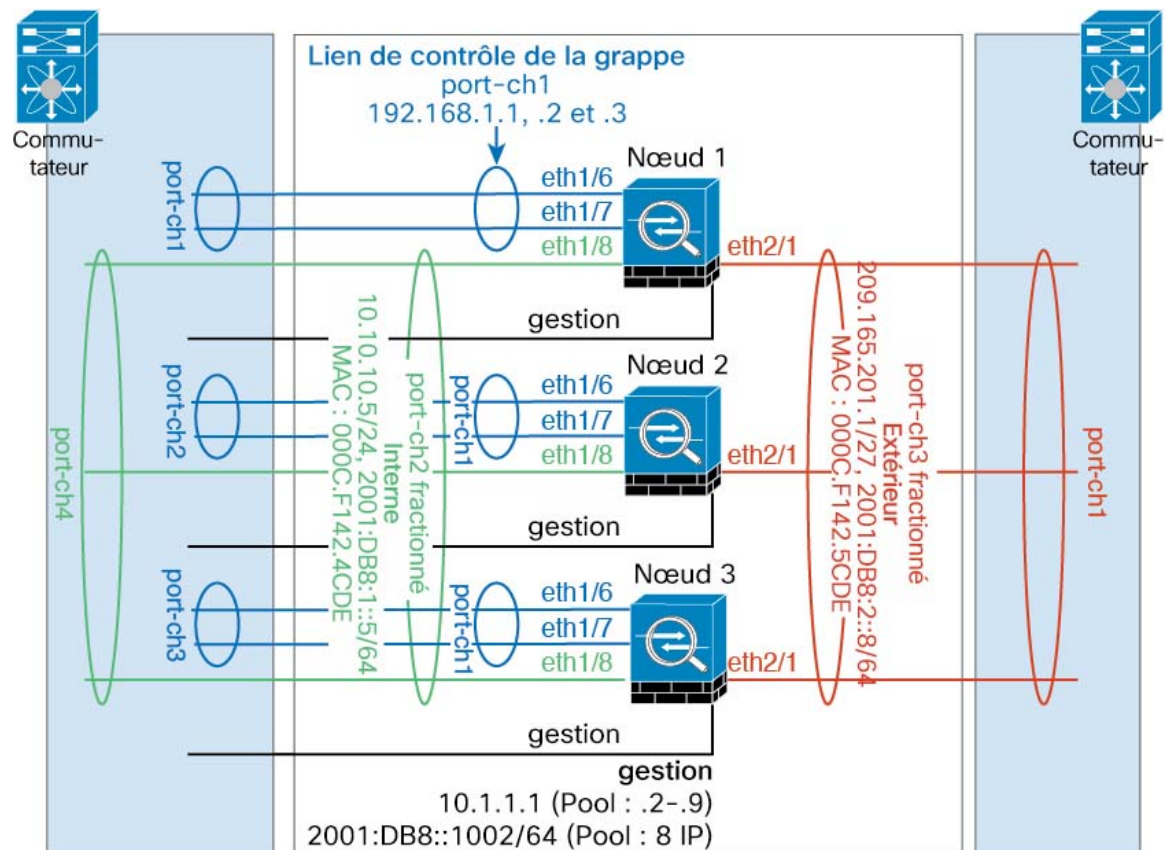
```

nameif inside
ip address 10.10.10.5 255.255.255.0
ipv6 address 2001:DB8:1::5/64
mac-address 000C.F142.4CDE

interface port-channel 1.20
vlan 20
nameif outside
ip address 209.165.201.1 255.255.255.224
ipv6 address 2001:DB8:2::8/64
mac-address 000C.F142.5CDE

```

## 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.

### Mode d'interface sur chaque unité

```
cluster interface-mode spanned force
```

**Configuration de démarrage de l'unité de contrôle de l'unité 1**

```
interface ethernet 1/6
 channel-group 1 mode on
 no shutdown

interface ethernet 1/7
 channel-group 1 mode on
 no shutdown

interface port-channel 1
 description CCL

cluster group cluster1
 local-unit asal
 cluster-interface port-channel1 ip 192.168.1.1 255.255.255.0
 priority 1
 key chuntheunavoidable
 enable noconfirm
```

**Configuration de démarrage de l'unité de données de l'unité 2**

```
interface ethernet 1/6
 channel-group 1 mode on
 no shutdown

interface ethernet 1/7
 channel-group 1 mode on
 no shutdown

interface port-channel 1
 description CCL

cluster group cluster1
 local-unit asa2
 cluster-interface port-channel1 ip 192.168.1.2 255.255.255.0
 priority 2
 key chuntheunavoidable
 enable as-data-node
```

**Configuration de démarrage de l'unité de données de l'unité 3**

```
interface ethernet 1/6
 channel-group 1 mode on
 no shutdown

interface ethernet 1/7
 channel-group 1 mode on
 no shutdown

interface port-channel 1
 description CCL

cluster group cluster1
 local-unit asa3
 cluster-interface port-channel1 ip 192.168.1.3 255.255.255.0
 priority 3
 key chuntheunavoidable
```

```
enable as-data-node
```

### Configuration de l'interface de l'unité de contrôle

```
ip local pool mgmt 10.1.1.2-10.1.1.9
ipv6 local pool mgmtipv6 2001:DB8::1002/64 8

interface management 1/1
 nameif management
 ip address 10.1.1.1 255.255.255.0 cluster-pool mgmt
 ipv6 address 2001:DB8::1001/32 cluster-pool mgmtipv6
 security-level 100
 management-only
 no shutdown

interface ethernet 1/8
 channel-group 2 mode active
 no shutdown

interface port-channel 2
 nameif inside
 ip address 10.10.10.5 255.255.255.0
 ipv6 address 2001:DB8:1::5/64
 mac-address 000C.F142.4CDE

interface ethernet 2/1
 channel-group 3 mode active
 no shutdown

interface port-channel 3
 nameif outside
 ip address 209.165.201.1 255.255.255.224
 ipv6 address 2001:DB8:2::8/64
 mac-address 000C.F142.5CDE
```

## 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'enverront 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

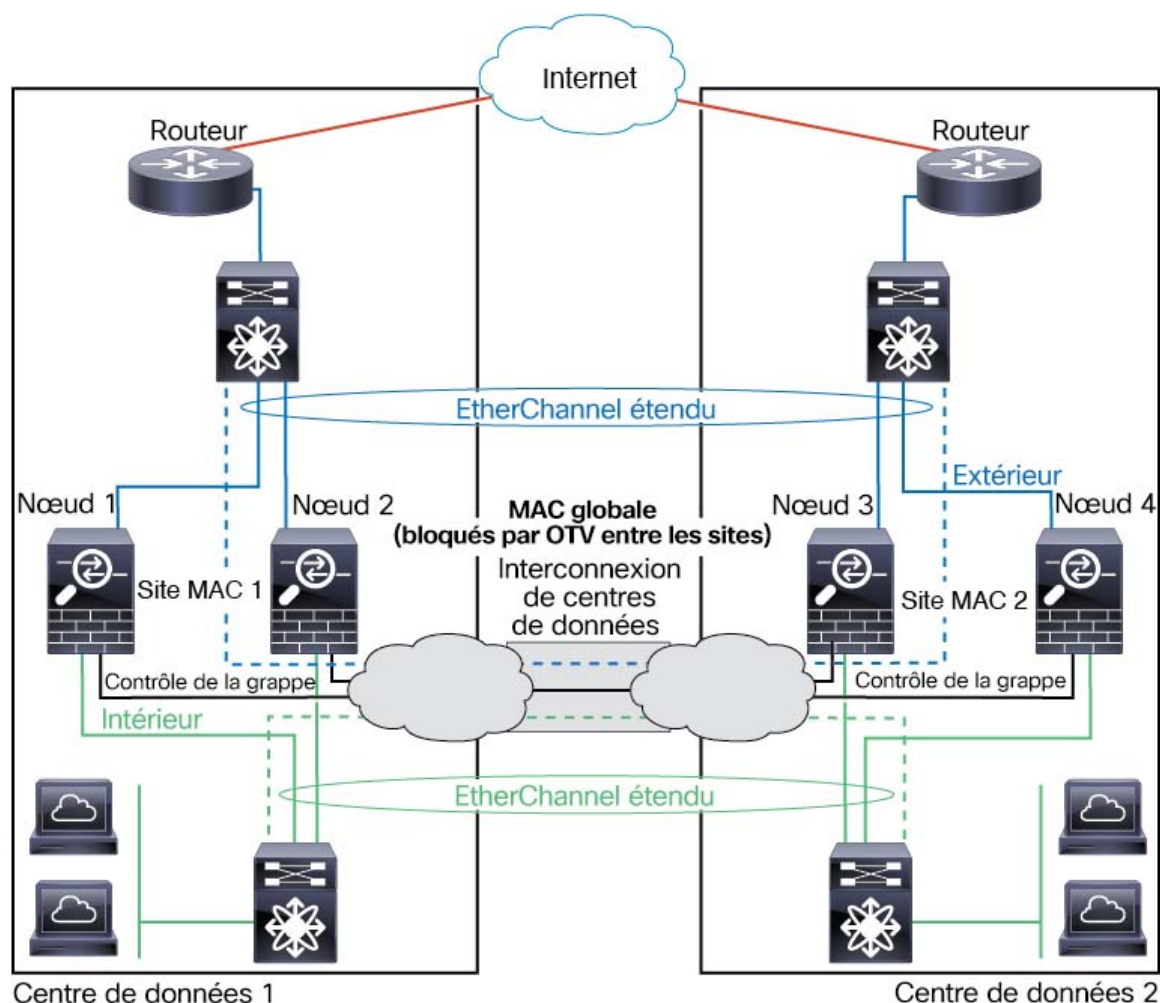
L'exemple suivant montre deux membres de la grappe dans chacun des deux centres de données, placés entre le routeur de passerelle et un réseau interne sur chaque site (insertion est-ouest). Les membres de la grappe sont connectés par la liaison de commande de grappe sur DCI. Les membres de la grappe 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

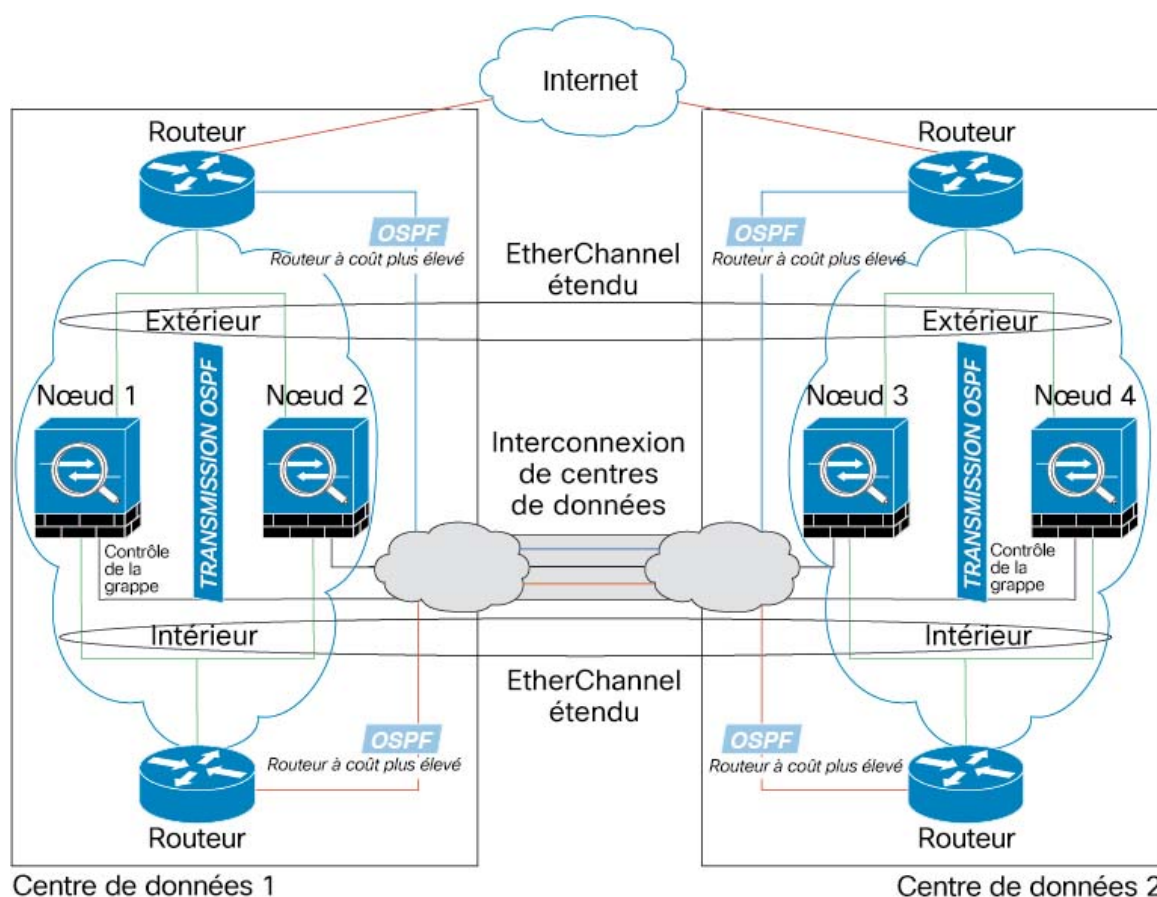
L'exemple suivant montre 2 membres de la grappe dans chacun des 2 centres de données placés entre les routeurs internes et externes (insertion nord-sud). Les membres de la grappe sont connectés par la liaison de commande de grappe sur DCI. Les membres de la grappe de chaque site se connectent aux commutateurs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

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

- 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.

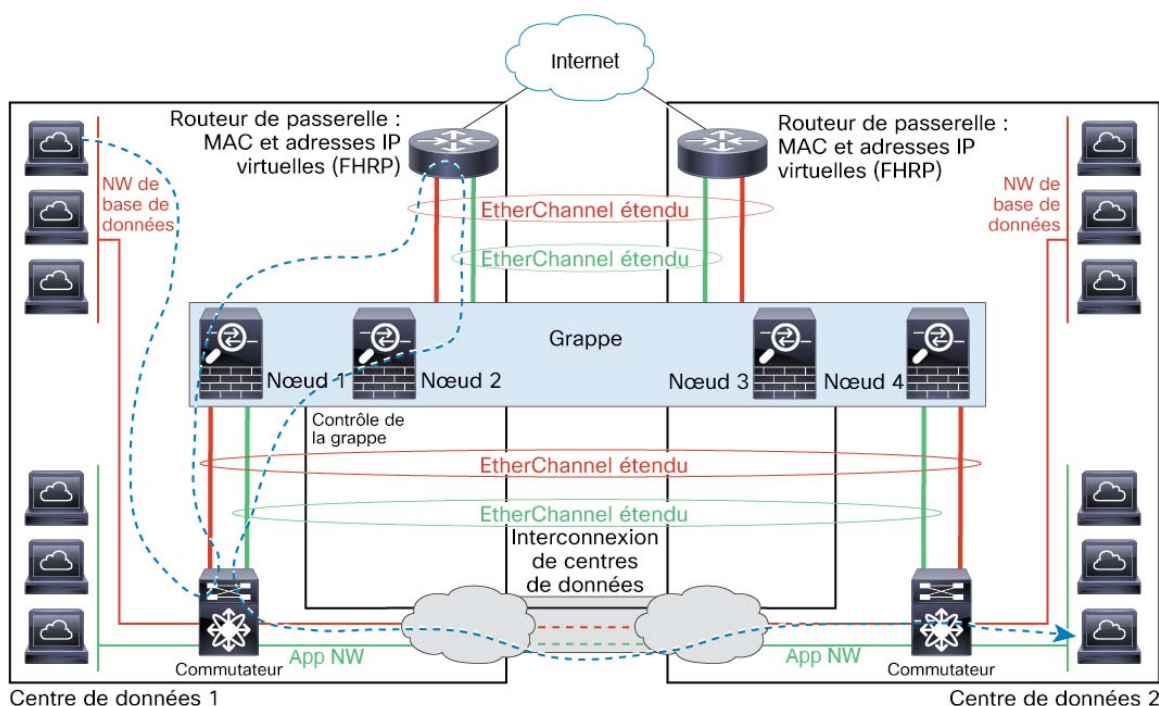


## Exemple inter-sites est-ouest en mode transparent de l'EtherChannel étendu

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)
- 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
- 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 rééquilibrage 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
- 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 nœuds individuels

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 et mise en grappe

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 et mise en grappe

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

### Routage multidiffusion en mode EtherChannel étendu

En mode EtherChannel étendu : 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

Cette section décrit comment utiliser le routage dynamique avec la mise en grappe.

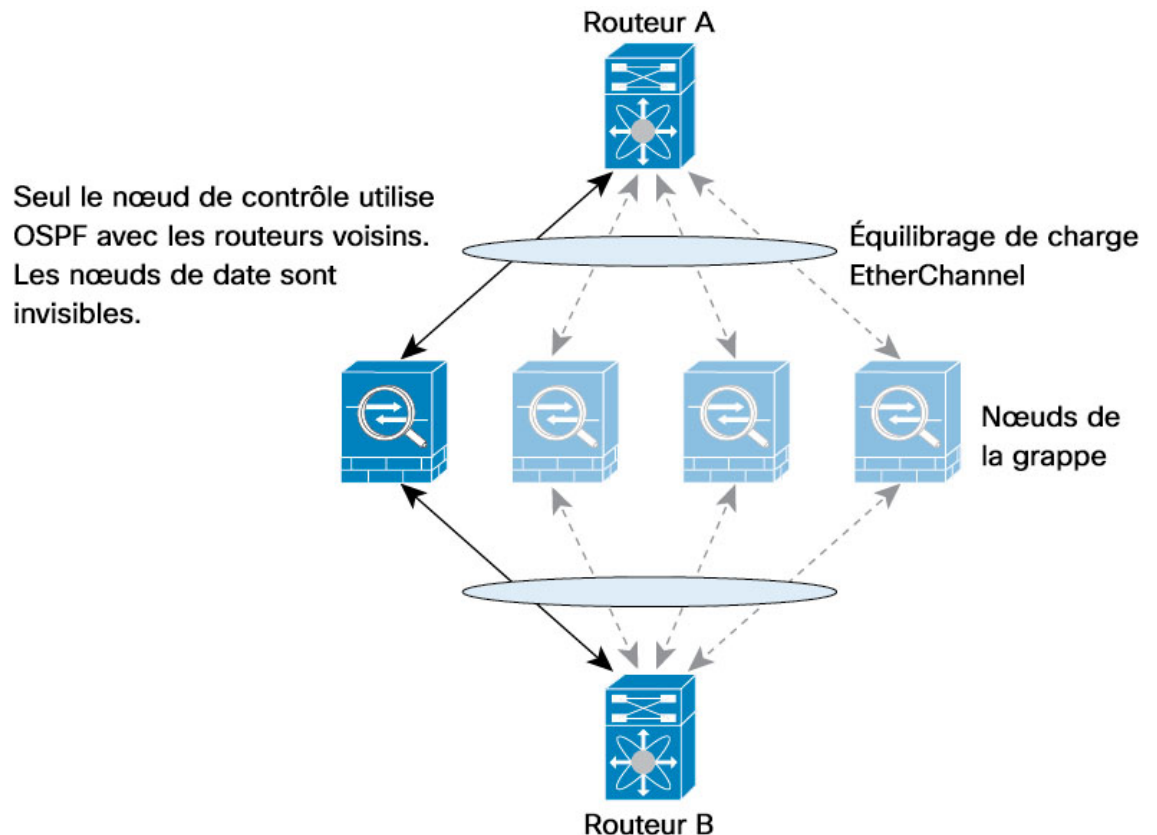
### Routage dynamique en mode EtherChannel étendu



**Remarque** IS-IS n'est pas pris en charge en mode EtherChannel étendu.

Le processus de routage ne s'exécute que sur le nœud de contrôle, et les routes sont apprises par le nœud de contrôle et répliquées sur les nœuds de données. Si un paquet de routage arrive à un nœud de données, il est redirigé vers le nœud de contrôle.

*Illustration 1 : Routage dynamique en mode EtherChannel étendu*



Une fois que le nœud de données a appris les routes du nœud de contrôle, chaque nœud prend des décisions de transfert indépendamment.

La base de données du LSA OSPF n'est pas synchronisée du nœud de contrôle avec les nœuds de données. S'il y a basculement du nœud 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

Le VPN de site à site est une fonctionnalité centralisée; seul le nœud de contrôle prend en charge les connexions VPN.



### Remarque

L'accès VPN à distance n'est pas pris en charge avec la mise en grappe.

La fonctionnalité VPN est limitée au nœud de contrôle et ne tire pas parti des capacités de haute disponibilité de la grappe. Si le nœud de contrôle tombe en panne, toutes les connexions VPN existantes sont perdues, et les utilisateurs de VPN verront une perturbation de service. Lorsqu'un nouveau nœud de contrôle est choisi, vous devez rétablir les connexions VPN.

Lorsque vous connectez un tunnel VPN à une adresse EtherChannel étendu, les connexions sont automatiquement transférées au nœud de contrôle.

Les clés et les certificats liés au VPN sont répliqués sur tous les nœuds.

## 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, si votre modèle peut gérer environ 10 Gbit/s de trafic lorsqu'il est exécuté seul, pour une grappe de 8 unités, le débit combiné maximal sera d'environ 80 % de 80 Gbit/s (8 unités x 10 Gbit/s) : 64 Gbit/s.

## Choix du nœud de contrôle

Les nœuds de la grappe communiquent sur la liaison de commande de grappe pour élire un nœud de contrôle comme suit :

1. Lorsque vous activez la mise en grappe pour un nœud (ou lorsqu'il démarre avec la mise en grappe déjà activée), il diffuse une demande de sélection toutes les 3 secondes.
2. Tous les autres nœuds ayant une priorité plus élevée répondent à la demande de sélection; la priorité est réglée entre 1 et 100, 1 étant la priorité la plus élevée.
3. Si, après 45 secondes, un nœud ne reçoit pas de réponse d'un autre nœud de priorité plus élevée, il devient le nœud de contrôle.

**Remarque**

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

4. Si un nœud se joint ultérieurement à la grappe avec une priorité plus élevée, il ne devient pas automatiquement le nœud de contrôle; le nœud de contrôle existant demeure toujours le nœud de contrôle, sauf s'il s'arrête de répondre, moment auquel un nouveau nœud de contrôle est sélectionné.
5. Dans un scénario de « discernement partagé », où il y a temporairement plusieurs nœuds de contrôle, le nœud ayant la priorité la plus élevée conserve le rôle tandis que les autres nœuds retournent aux rôles de nœud de données.

**Remarque**

Vous pouvez forcer manuellement un nœud à devenir le nœud de contrôle. Pour les fonctionnalités centralisées, si vous forcez un changement de nœud de contrôle, toutes les connexions sont abandonnées et vous devez rétablir les connexions sur le nouveau nœud de contrôle.

## Haute disponibilité au sein de la grappe

La mise en grappe assure une disponibilité élevée en surveillant l'intégrité des nœuds et de l'interface et en reproduisant les états de la connexion entre les nœuds.

### Surveillance de l'intégrité du nœud

Chaque nœud envoie périodiquement un paquet de diffusion heartbeat sur la liaison de commande de grappe. Si le nœud de contrôle ne reçoit aucun paquet heartbeat ou autre paquet d'un nœud de données au cours du délai 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 élu parmi les nœuds restants.

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 éliront 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 éliront 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 du nœud de contrôle, à la page 85](#) pour de plus amples renseignements.

### Surveillance d'interfaces

Chaque nœud surveille l'état de la liaison de toutes les interfaces matérielles désignées utilisées et signale les modifications d'état au nœud de contrôle.

- Spanned EtherChannel (EtherChannel étendu) : Utilise le protocole cLACP (cluster Link Aggregation Control Protocol). Chaque nœud surveille l'état de la liaison et les messages du protocole cLACP pour déterminer si le port est toujours actif dans l'EtherChannel. L'état est signalé au nœud de contrôle.

Lorsque vous activez la surveillance de l'intégrité, toutes les interfaces physiques (y compris les interfaces EtherChannel principales) sont surveillées par défaut; vous pouvez éventuellement désactiver la surveillance par interface. Seules les interfaces nommées peuvent être surveillées. Par exemple, l'EtherChannel désigné doit échouer pour être considéré comme en échec, ce qui signifie que tous les ports membres d'un EtherChannel ne doivent pas déclencher la suppression de la grappe (selon le paramètre minimal de groupement de ports).

Un nœud est supprimé de la grappe en cas de défaillance de ses interfaces surveillées. Le délai avant que ASA ne supprime un membre de la grappe dépend du fait que le nœud est un membre établi ou en train de se joindre à la grappe. si l'interface est en panne sur un membre établi, le ASA supprime le membre après 9 secondes. 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 les non-EtherChannels, le nœud est supprimé après 500 ms, quel que soit l'état membre.

## 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 nœud 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 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 en saisissant la commande **cluster group name**, puis **enable**. Ce comportement est configurable.
- Nœud en échec : si le nœud a été supprimé de la grappe en raison d'un échec de vérification de l'intégrité du nœud, la jonction avec la grappe dépend de la source de la défaillance. Par exemple, une panne de courant temporaire signifie que le nœud rejoindra la grappe au redémarrage, à condition que la liaison

de commande de grappe soit active et que la mise en grappe soit toujours activée avec la commande **enable**. L'ASA tente de rejoindre la grappe toutes les 5 secondes.

- Erreur interne : les défaillances internes comprennent : le dépassement du délai de synchronisation de l'application, les statuts incohérents de l'application, etc. Un nœud tentera de rejoindre la grappe automatiquement aux intervalles suivants : 5 minutes, 10 minutes, puis 20 minutes. Ce comportement est configurable.

Consultez [Configurer les paramètres de démarrage du nœud de contrôle](#), à la page 27.

## 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                                    | <b>Non</b>               | —                                                                                                      |
| 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.

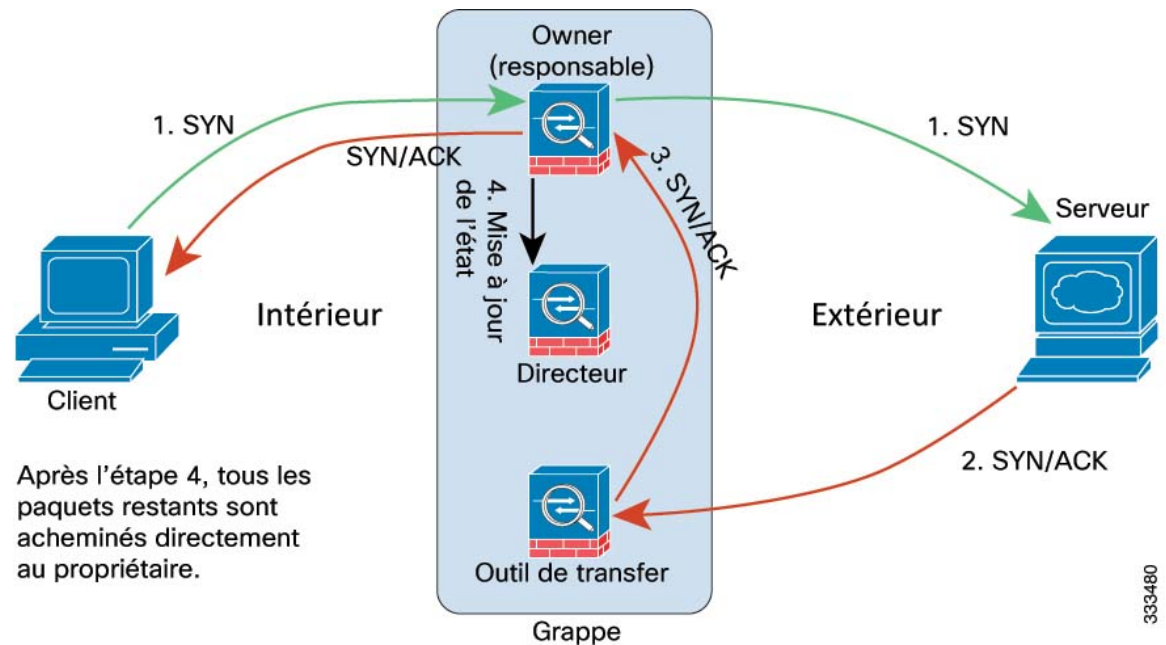
Vous pouvez modifier les valeurs par défaut de PAT par session pour les protocoles TCP et UDP afin que les connexions à ces protocoles soient gérées par session ou multisession selon la configuration. Pour ICMP, vous ne pouvez pas modifier la PAT multisession par défaut. Référez-vous au guide de configuration du pare-feu pour obtenir plus de renseignements sur les serveurs mandataires TLS.

## Nouvelle propriété de connexion

Lorsqu'une nouvelle connexion est acheminée vers un nœud de la grappe par l'équilibrage de la charge, ce nœud possède les deux sens de connexion. Si des paquets de connexion arrivent à un nœud différent, ils sont acheminés au nœud propriétaire sur la liaison de commande de grappe. Pour de meilleures performances, un bon équilibrage de la charge externe est nécessaire pour que les deux directions d'un flux arrivent au même nœud et pour que les flux soient répartis uniformément entre les nœuds. Si un flux inverse arrive sur un autre nœud, il est redirigé vers le nœud d'origine.

## Exemple de flux de données pour TCP

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



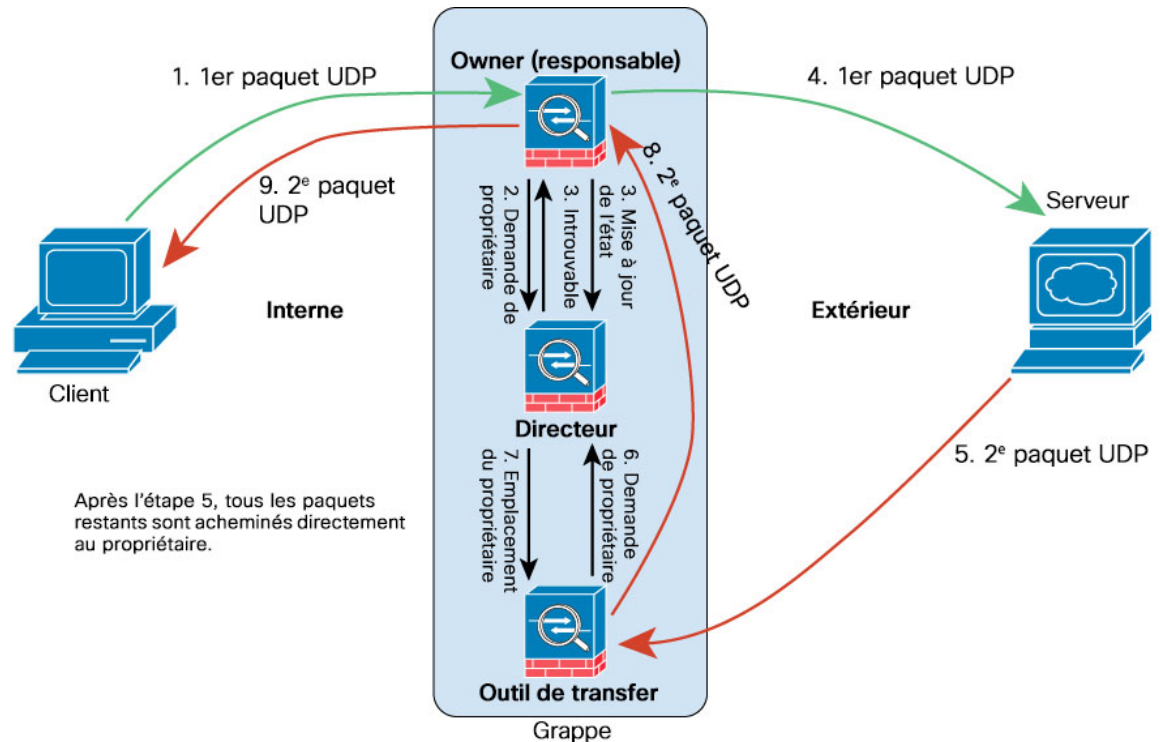
1. Le paquet SYN provient du client et est livré à un ASA (selon la méthode d'équilibrage de la charge), qui devient le propriétaire. Le propriétaire crée un flux, code les renseignements sur le propriétaire dans un témoin SYN et transfère le paquet au serveur.
2. Le paquet SYN-ACK provient du serveur et est livré à un ASA différent (selon la méthode d'équilibrage de la charge). Cet ASA est le transitaire.
3. Comme le transitaire n'est pas propriétaire de la connexion, il decode les informations sur le propriétaire à partir du témoin SYN, crée un flux de transfert vers le propriétaire et transmet le SYN-ACK au propriétaire.
4. Le propriétaire envoie une mise à jour de l'état au directeur et transmet le SYN-ACK au client.
5. Le directeur reçoit la mise à jour d'état du propriétaire, crée un flux à destination du propriétaire et enregistre les informations d'état TCP ainsi que le propriétaire. Le directeur agit en tant que propriétaire secondaire pour la connexion.
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 pour le Secure Firewall

| Nom de la caractéristique                                                                     | Version | Renseignements sur les fonctionnalités                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|-----------------------------------------------------------------------------------------------|---------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| VPN site à site distribué avec mise en grappe sur le modèle Cisco Secure Firewall 4200        | 9.23(1) | <p>Une grappe ASA sur le modèle Cisco Secure Firewall 4200 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 le nœud 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.</p> <p>Commandes ajoutées ou modifiées : <b>cluster redistribute vpn-sessiondb</b>, <b>show cluster vpn-sessiondb</b>, <b>vpn-mode</b>, <b>show cluster resource usage</b>, <b>show vpn-sessiondb</b>, <b>show conn detail</b>, <b>show crypto ikev2 stats</b></p> |
| Suppression des propos tendancieux                                                            | 9.19(1) | <p>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).</p> <p>Commandes nouvelles/modifiées : <b>cluster control-node</b>, <b>enable as-data-node</b>, <b>prompt</b>, <b>show cluster history</b>, <b>show cluster info</b></p>                                                                                                                                                                                                                                                                                                                                           |
| La prise en charge de la mise en grappe sur le pare-feu Secure Firewall 3100 a été introduite | 9.17(1) | Vous pouvez mettre en grappe jusqu'à 8 nœuds Secure Firewall 3100 en mode EtherChannel étendu.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |



## 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.