



Grappe ASA pour l'ASA virtuel dans un nuage public

La mise en grappe vous permet de regrouper plusieurs ASA virtuel 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. Vous pouvez déployer des grappes ASA virtuel dans un nuage public en utilisant les éléments suivants :

- Amazon Web Services (AWS)

Actuellement, seul le mode pare-feu routé est pris en charge.



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 39](#).

- [À propos de la mise en grappe d'ASA virtuel dans le nuage public, à la page 1](#)
- [Licences pour la mise en grappe d'ASA virtuel, à la page 6](#)
- [Exigences et conditions préalables à la mise en grappe d'ASA virtuel, à la page 6](#)
- [Lignes directrices relatives à la mise en grappe d'ASA virtuel, à la page 8](#)
- [Déployer la grappe dans AWS, à la page 9](#)
- [Personnaliser l'opération de mise en grappe, à la page 19](#)
- [Gérer les nœuds de la grappe, à la page 23](#)
- [Surveillance de la grappe, à la page 27](#)
- [Référence pour la mise en grappe, à la page 38](#)
- [Historique de la mise en grappe d'ASA virtuel dans le nuage public, à la page 54](#)

À propos de la mise en grappe d'ASA virtuel dans le nuage public

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 un seul périphérique. Pour agir comme une grappe, les pare-feu ont besoin de l'infrastructure suivante :

- Réseau isolé pour la communication intra-grappe, appelé *liaison de commande de grappe*, qui utilise des interfaces VXLAN. Les VXLAN, qui agissent comme des réseaux virtuels de couche 2 sur des réseaux physiques de couche 3, permettent à la ASA virtuel d'envoyer des messages en diffusion ou en multidiffusion sur la liaison de commande de grappe.
- Équilibreur(s) de charges : pour l'équilibrage de charges externe, vous avez les options suivantes :

- Équilibreur de charge de passerelle AWS

L'équilibreur de charge de passerelle AWS combine une passerelle de réseau transparente et un équilibreur de charge qui répartit le trafic et fait évoluer les périphériques virtuels à la demande. Le ASA virtuel prend en charge le plan de contrôle centralisé de l'équilibreur de charge de passerelle avec un plan de données distribué (point terminal de l'équilibreur de charge de passerelle) à l'aide d'un serveur mandataire à un seul bras d'interface de Geneve.

- Routage à chemins multiples à coût égal (ECMP) utilisant des routeurs internes et externes comme le routeur des services en nuage de Cisco

Le routage ECMP peut transférer des paquets sur plusieurs « meilleurs chemins » qui se partagent la première place dans la mesure du routage. Comme pour l'EtherChannel, un hachage des adresses IP source et de destination ou des ports source et de destination peut être utilisé pour envoyer un paquet vers l'un des sauts suivants. Si vous utilisez des routes statiques pour le routage ECMP, la défaillance de ASA virtuel peut provoquer des problèmes. le routage continue d'être utilisé et le trafic vers le ASA virtuel défaillant sera perdu. Si vous utilisez des routes statiques, veillez à utiliser une fonctionnalité de surveillance de routage statique telle que le suivi d'objets. Nous recommandons d'utiliser des protocoles de routage dynamique pour ajouter et supprimer des routes, auquel cas vous devez configurer chaque ASA virtuel pour qu'il participe au routage dynamique.



Remarque

Les canaux EtherChannels étendus de couche 2 ne sont pas pris en charge pour l'équilibrage de la charge.

Nœuds de la grappe

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

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 individuelles

Vous pouvez configurer les interfaces de grappe en tant qu'*interfaces individuelles*.

Les interfaces individuelles sont des interfaces de routage normales, chacune avec sa propre adresse IP locale. La configuration d'interface doit être configurée uniquement sur le nœud de contrôle et chaque interface utilise DHCP.

**Remarque**

Les canaux EtherChannels étendus de couche 2 ne sont pas pris en charge.

Liaison de commande de grappe

Chaque nœud doit dédier une interface en tant qu'interface VXLAN (VTEP) pour la liaison de commande de grappe. Pour en savoir plus sur VXLAN, consultez [Interfaces VXLAN](#).

Point terminal du tunnel VXLAN

Les périphériques de point terminal de tunnel VXLAN (VTEP) effectuent l'encapsulation et la désencapsulation VXLAN. Chaque VTEP comporte deux types d'interface : une ou plusieurs interfaces virtuelles appelées interfaces VNI (VXLAN Network Identifier), et une interface normale appelée interface source du VTEP qui canalise les interfaces VNI entre les VTEP. L'interface source du VTEP est connectée au réseau IP de transport pour la communication de VTEP à VTEP.

Interface de la source VTEP

L'interface source du VTEP est une interface ASA virtuel classique à laquelle vous prévoyez associer l'interface VNI. Vous pouvez configurer une interface source de VTEP pour qu'elle agisse en tant que liaison de commande de grappe. L'interface source est réservée à une utilisation avec la liaison de commande de grappe uniquement. Chaque interface source de VTEP 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.

Interface VNI

Une interface VNI est semblable à une interface VLAN : il s'agit d'une interface virtuelle qui sépare le trafic réseau sur une interface physique donnée au moyen de balisage. Vous ne pouvez configurer qu'une seule interface VNI. Chaque interface VNI possède une adresse IP sur le même sous-réseau.

VTEP homologues

Contrairement au VXLAN habituel pour les interfaces de données, qui autorise un seul homologue VTEP, la mise en grappe ASA virtuel vous permet de configurer plusieurs homologues.

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.

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 virtuel 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 du DHCP ou de l'ensemble d'adresses IP de la grappe. Si vous utilisez un ensemble d'adresses IP de grappe, 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 le nœud de contrôle). Vous devez utiliser le port de console (si disponible) pour toute autre configuration.

Réplication de la configuration

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

Gestion des grappes ASA virtuel

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

Réseau de gestion

Nous vous conseillons de connecter tous les nœuds à un seul réseau de gestion. Ce réseau est distinct de la liaison de commande de grappe.

Interface de gestion

Utilisez l'interface Management 0/0 pour la gestion.



Remarque

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

Vous pouvez utiliser l'adressage statique ou DHCP pour l'adresse IP de gestion.

Si vous utilisez l'adressage statique, vous pouvez utiliser l'adresse IP de la grappe principale qui est une adresse fixe pour la grappe qui appartient toujours au nœud de contrôle actuel. Pour chaque interface, vous configurez également une plage d'adresses de sorte que chaque nœud, y compris le nœud de contrôle actuel, 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'un nœud de contrôle change, l'adresse IP de la grappe principale est déplacée vers le nouveau nœud 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 au nœud de contrôle actuel. Pour gérer un membre, vous pouvez vous connecter à l'adresse IP locale. Pour le trafic de gestion sortant tel que TFTP ou le journal système, chaque nœud, y compris le nœud de contrôle, utilise l'adresse IP locale pour se connecter au serveur.

Si vous utilisez DHCP, vous n'utilisez pas un ensemble d'adresses locales ou n'avez pas d'adresse IP de grappe principale.



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.

Gestion de nœud de contrôle vs Gestion de nœud 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.

Licences pour la mise en grappe d'ASA virtuel

Chaque nœud de grappe nécessite la même licence de modèle. Nous vous conseillons d'utiliser le même nombre de CPU et de mémoire pour tous les nœuds, sinon les performances seront limitées sur tous les nœuds pour correspondre au membre le moins performant. Le niveau de débit sera répliqué du nœud de contrôle à chaque nœud de données afin qu'ils correspondent.



Remarque

Si vous annulez l'enregistrement de l'ASA virtuel pour qu'il ne soit pas sous licence, il reviendra à un état de débit fortement limité si vous rechargez l'ASA virtuel. Un nœud de grappe sans licence et à faible performance aura une incidence négative sur les performances de l'ensemble de la grappe. Assurez-vous de conserver tous les nœuds de la grappe sous licence ou de supprimer les nœuds sans licence.

Exigences et conditions préalables à la mise en grappe d'ASA virtuel

Exigences du modèle

- ASAv30, ASAv50, ASAv100

- Les services de nuage public suivants :
 - Amazon Web Services (AWS)
- Maximum de 16 nœuds

Consultez également les exigences générales pour l'ASA virtuel dans le [Guide de démarrage de l'ASA virtuel](#).

Configuration matérielle et logicielle requise

Tous les nœuds d'une grappe :

- Il doit s'agir du même niveau de performance. Nous vous recommandons d'utiliser le même nombre de CPU et de mémoire pour tous les nœuds, sinon les performances seront limitées sur tous les nœuds pour correspondre au nœud le moins performant.
- 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.
- Prise en charge du déploiement de zone de disponibilité unique.
- Les interfaces de liaison de commande de grappe doivent se trouver dans le même sous-réseau, de sorte que la grappe doit être déployée dans le même sous-réseau.

MTU

Assurez-vous que les ports connectés à la liaison de commande de grappe ont une MTU correcte (plus élevée) configurée. En cas de non-concordance MTU, la formation de la grappe échouera.

La MTU de la liaison de commande de grappe doit être 154 octets supérieure aux interfaces de données. Étant donné que le trafic de la liaison de commande de grappe comprend la transmission de paquets de données, celle-ci doit prendre en charge la taille totale d'un paquet de données, plus les surcharges de trafic de la grappe (100 octets) et les surcharges VXLAN (54 octets).

Pour AWS avec GWLB, l'interface de données utilise l'encapsulation de Geneve. Dans ce cas, le datagramme Ethernet entier est encapsulé, de sorte que le nouveau paquet est plus volumineux et nécessite une MTU plus grande. Vous devez définir la MTU de l'interface source comme étant la MTU du réseau + 306 octets. Ainsi, pour le chemin réseau standard de 1 500 MTU, la MTU de l'interface source doit être de 1 806 et la MTU de la liaison de commande de grappe doit être de +154, 1 960.

Le tableau suivant présente la MTU de la liaison de commande de grappe et la MTU de l'interface de données suggérées.



Remarque

Nous ne recommandons pas de définir la MTU de la liaison de commande de grappe entre 2561 et 8362. En raison de la gestion du groupe de blocs, cette taille MTU n'est pas optimale pour le fonctionnement du système.

Tableau 1 : MTU suggérée

Nuage public	Liaison de commande de grappe	MTU de l'interface de données
AWS avec GWLB	1960	1806

Nuage public	Liaison de commande de grappe	MTU de l'interface de données
AWS	1654	1 500

Lignes directrices relatives à la mise en grappe d'ASA virtuel

Haute disponibilité

La haute disponibilité n'est pas prise en charge par la mise en grappe.

IPv6

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

Directives supplémentaires

- Lorsque des modifications de topologie surviennent (par exemple, ajout ou suppression d'une interface de données, activation ou désactivation d'une interface sur l'ASA virtuel ou le commutateur, ajout d'un commutateur supplémentaire pour former un système de commutation redondant), 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'un nœud à une grappe existante ou lors du rechargement d'un nœud, il se produit 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.
- Pour les connexions TLS/SSL déchiffrées, les états de déchiffrement ne sont pas synchronisés, et si le propriétaire de la connexion échoue, les connexions déchiffrées sont réinitialisées. De nouvelles connexions devront être établies vers un nouveau nœud. Les connexions qui ne sont pas déchiffrées (elles correspondent à une règle « ne pas déchiffrer ») ne sont pas affectées et sont répliquées correctement.
- L'évolutivité dynamique n'est pas prise en charge.
- Le basculement avec état de la cible n'est pas pris en charge lorsque vous déployez la grappe sur AWS.
- Effectuer un déploiement global à la fin de chaque fenêtre de maintenance.
- Assurez-vous de ne pas supprimer plusieurs périphériques à la fois du groupe d'évolutivité automatique (AWS) ou de l'ensemble d'évolutivité (Azure). Nous vous conseillons également d'exécuter la commande **cluster disable** sur le périphérique avant de retirer ce périphérique du groupe d'évolutivité (AWS) ou de l'ensemble d'évolutivité (Azure).
- Si vous souhaitez désactiver les nœuds de données et le nœud de contrôle dans une grappe, nous vous recommandons de désactiver les nœuds de données avant de désactiver le nœud de contrôle. Si un nœud de contrôle est désactivé alors qu'il y a d'autres nœuds de données dans la grappe, l'un d'eux doit être promu au rang de nœud de contrôle. Notez que le changement de rôle pourrait perturber la grappe.

- Dans les scripts de configuration Day0 présentés dans ce guide, vous pouvez modifier les adresses IP selon vos besoins, fournir des noms d'interface personnalisés et modifier la séquence de l'interface CCL-Link.

Valeurs par défaut pour la mise en grappe

- L'ID du 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.
- Un délai de duplication de connexion de 5 secondes est activé par défaut pour le trafic HTTP.

Déployer la grappe dans AWS

Pour déployer une grappe dans AWS, vous pouvez soit la déployer manuellement, soit utiliser des modèles CloudFormation pour déployer une pile. Vous pouvez utiliser la grappe avec l'équilibreur de charge de passerelle AWS ou avec un équilibreur de charge non natif comme le routeur des services en nuage de Cisco.

Équilibreur de charge de passerelle AWS et serveur mandataire à un seul volet de Geneve

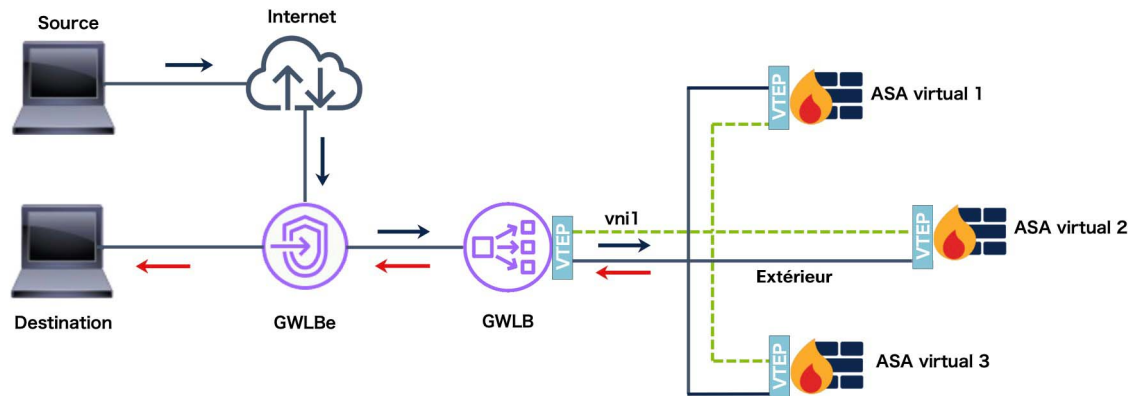


Remarque

Ce scénario est le seul actuellement pris en charge pour les interfaces de Geneve.

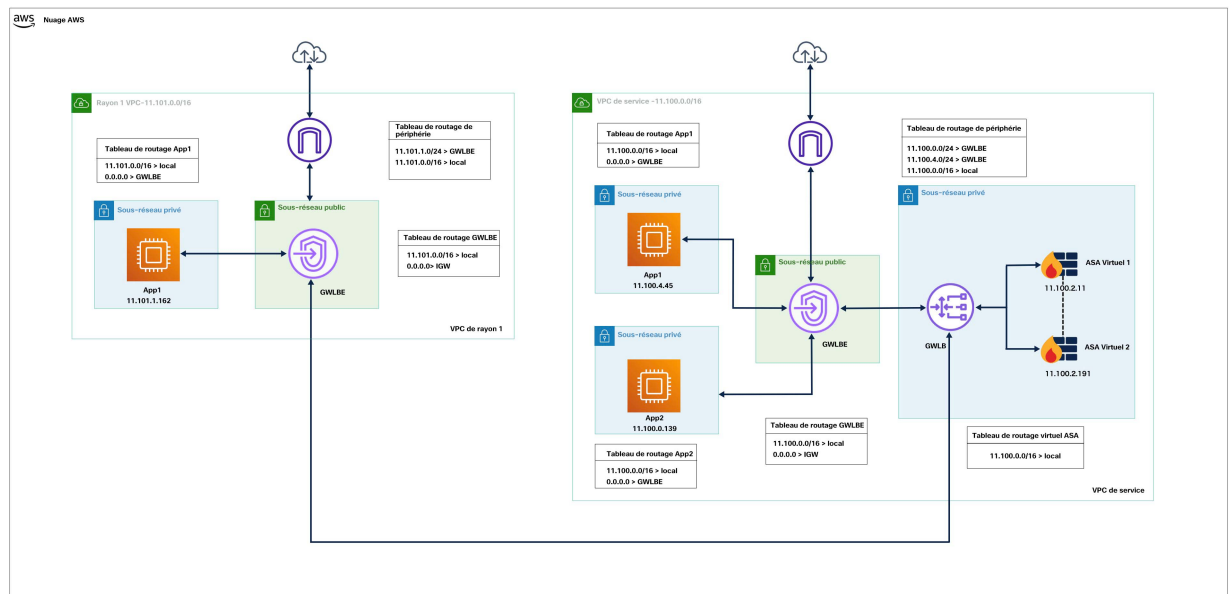
L'équilibreur de charge de passerelle AWS combine une passerelle de réseau transparente et un équilibreur de charge qui répartit le trafic et fait évoluer les périphériques virtuels à la demande. L'ASA virtuel prend en charge le plan de commande centralisé de l'équilibreur de charges de passerelle avec un plan de données distribué (point terminal de l'équilibreur de charges de passerelle). La figure suivante montre le trafic acheminé vers l'équilibreur de charge de passerelle à partir du point terminal de l'équilibreur de charge de passerelle. L'équilibreur de charges de passerelle équilibre le trafic entre plusieurs ASA virtuels, qui inspectent le trafic avant de l'abandonner ou de le renvoyer à l'équilibreur de charges de passerelle (trafic à demi-tour). L'équilibreur de charge de passerelle renvoie ensuite le trafic au point terminal de l'équilibreur de charge de passerelle et à la destination.

Illustration 1 : Serveur mandataire à un seul volet Geneve



Exemple de topologie

La topologie indiquée ci-dessous décrit le flux de trafic entrant et sortant. Il y a deux instances ASA virtuel dans la grappe connectée à une GWLB.



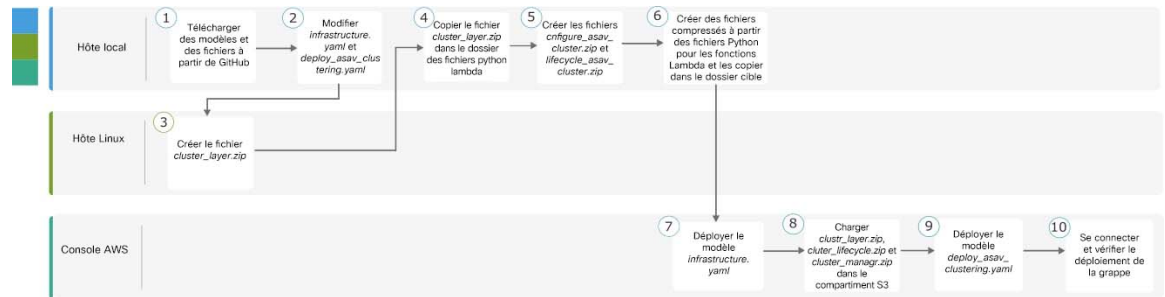
Le trafic entrant provenant d'Internet est dirigé vers le point terminal de la GWLB, qui le transmet ensuite à la GWLB. Le trafic est ensuite acheminé vers la grappe ASA virtuel. Une fois que le trafic a été inspecté par une instance ASA virtuel dans la grappe, il est transféré à la VM de l'application, App1.

Le trafic sortant d'App1 est transmis au point terminal de la GWLB, qui l'envoie ensuite vers Internet.

Processus de bout en bout pour le déploiement de la grappe ASA virtuel sur AWS

Déploiement basé sur un modèle

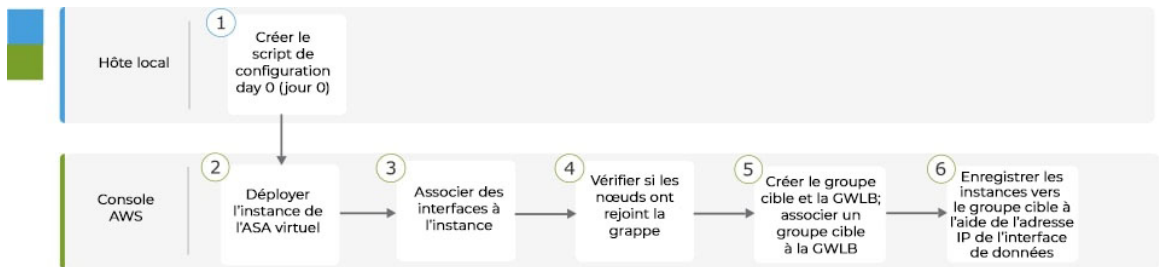
Le diagramme suivant illustre le flux de travail pour le déploiement basé sur le modèle de la grappe ASA virtuel sur AWS.



	Espace de travail	Étapes
1	Hôte local	Téléchargez des modèles et des fichiers à partir de GitHub.
2	Hôte local	Modifiez les modèles <i>infrastructure.yaml</i> et <i>deploying_asav_clustering.yaml</i> .
3	Hôte Linux	Créez le fichier <i>cluster_layer.zip</i> .
4	Hôte local	Copiez le fichier <i>cluster_layer.zip</i> dans le dossier des fichiers Python Lambda.
5	Hôte local	Créez les fichiers <i>configure_asav_cluster.zip</i> et <i>lifecycle_asav_cluster.zip</i> .
6	Hôte local	Créez des fichiers compressés à partir des fichiers Python pour les fonctions Lambda et les copier dans le dossier cible.
7	Console AWS	Déployez le modèle <i>infrastructure.yaml</i> .
8	Console AWS	Chargez <i>cluster_layer.zip</i> , <i>cluster_lifecycle.zip</i> et <i>cluster_manager.zip</i> dans le compartiment S3.
9	Console AWS	Déployez le modèle <i>deploy_asav_clustering.yaml</i> .
10	Console AWS	Connectez-vous et vérifiez le déploiement de la grappe.

Déploiement manuel

Le diagramme suivant illustre le flux de travail pour le déploiement manuel de la grappe ASA virtuel sur AWS.



	Espace de travail	��tapes
1	H��te local	Cr��ez le script de configuration de jour 0.
2	Console AWS	D��ployez l'instance de l'ASA virtuel.
3	Console AWS	Associez des interfaces �� l'instance.
4	Console AWS	V��rifiez si les n��uds ont rejoint la grappe.
5	Console AWS	Cr��ez le groupe cible et la GWLB; associer un groupe cible �� la GWLB.
6	Console AWS	Enregistrez les instances avec le groupe cible �� l'aide de l'adresse IP de l'interface de donn��es.

Mod  les

Les mod  les fournis ci-dessous sont disponibles dans GitHub. Les valeurs des param  tres sont explicites et les noms des param  tres, les valeurs par d  faut, les valeurs autoris  es et la description sont donn  es dans le mod  le.

- [infrastructure.yaml](#) : mod  le pour le d  ploiement de l'infrastructure
- [deploy_asav_clustering.yaml](#) : mod  le pour le d  ploiement en grappe.



Remarque

Assurez-vous de consulter la liste des types d'instances AWS pris en charge avant de d  ployer les n  uds de la grappe. Cette liste se trouve dans le mod  le *deploy_asav_clustering.yaml*, sous les valeurs autoris  es pour le param  tre InstanceType (Type d'instance).

D  ployer la pile dans AWS    l'aide d'un mod  le CloudFormation

D  ployez la pile dans AWS    l'aide du mod  le personnalis   Cloud Formation.

Avant de commencer

- Vous avez besoin d'un ordinateur Linux avec Python 3.

Procédure

Étape 1

Préparez le modèle.

- a) Copiez le référentiel github dans votre dossier local. Consultez <https://github.com/CiscoDevNet/cisco-asav/tree/master/cluster/aws>.
- b) Modifiez **infrastructure.yaml** et **deploy_asav_clustering.yaml** avec les paramètres requis.
- c) Créez un fichier nommé **cluster_layer.zip** pour fournir les bibliothèques Python essentielles aux fonctions Lambda.

Nous vous recommandons d'utiliser Amazon Linux avec Python 3.9 installé pour créer le fichier **cluster_layer.zip**.

Remarque

Si vous avez besoin d'un environnement Amazon Linux, vous pouvez créer une instance EC2 à l'aide de l'AMI d'Amazon Linux 2023 ou utiliser AWS Cloudshell, qui exécute la dernière version d'Amazon Linux.

Pour créer le fichier **cluster-layer.zip**, vous devez d'abord créer le fichier **requirements.txt** contenant les détails du paquet de la bibliothèque python, puis exécuter le script d'interface Shell.

1. Créez le fichier **requirements.txt** en précisant les détails du paquet Python.

Voici un exemple de détails de paquet que vous fournissez dans le fichier **requirements.txt** :

```
$ cat requirements.txt
pycryptodome
paramiko
requests
scp
jsonschema
cffi
zip
importlib-metadata
```

2. Exécutez le script Shell suivant pour créer le fichier **cluster_layer.zip**.

```
$ pip3 install --platform manylinux2014_x86_64
--target=./python/lib/python3.9/site-packages
--implementation cp --python-version 3.9 --only-binary=:all:
--upgrade -r requirements.txt
$ zip -r cluster_layer.zip ./python
```

Remarque

Si vous rencontrez une erreur de conflit de dépendances lors de l'installation, comme une erreur liée à urllib3 ou une erreur de cryptographie, il est recommandé d'inclure les paquets en conflit ainsi que leurs versions recommandées dans le fichier **requirements.txt**. Après cela, vous pouvez exécuter à nouveau l'installation pour résoudre le conflit.

- d) Copiez le fichier **cluster_layer.zip** résultant dans le dossier des fichiers lambda python.
- e) Créez les fichiers **configure_asav_cluster.zip** et **lifecycle_asav_cluster.zip**.

Un fichier **make.py** se trouve dans le répertoire supérieur du référentiel cloné. Cela compressera les fichiers python dans un fichier compressé et les copiera dans un dossier cible.

python3 make.py build

Étape 2

Déployez **infrastructure.yaml** et notez les valeurs de sortie pour le déploiement en grappe.

- Sur la console AWS, accédez à **CloudFormation** et cliquez sur **Create stack** (Créer une pile); sélectionnez **With new resources (standard)** (Avec de nouvelles ressources (standard)).
- Sélectionnez **Upload a template file** (Charger un fichier modèle), cliquez sur **Choose file** (Choisir un fichier) et sélectionnez **infrastructure.yaml** dans le dossier cible.
- Cliquez sur **Next** (suivant) et fournissez les informations requises.
- Cliquez sur **Next** (Suivant), puis sur **Create stack** (Créer une pile).
- Une fois le déploiement terminé, accédez à **Outputs** (Résultats) et notez le nom de **compartiment S3**.

Illustration 2 : Sortie de infrastructure.yaml

Outputs (13)				
Search outputs				
Key	Value	Description	Export name	
AZ	sa-east-1a	Availability zone	-	
BucketName	ran-cls-infra-s3bucketcluster-kckr7518u00l	Name of the Amazon S3 bucket	-	
BucketUrl	http://ran-cls-infra-s3bucketcluster-kckr7518u00l.s3-website-sa-east-1.amazonaws.com	URL of S3 Bucket Static Website	-	
CCLSubnetId	subnet-050feb347e57eba99	CCL subnet ID	-	
EIPforNATgw	52.67.246.95	EIP reserved for NAT GW	-	
InInterfaceSGId	sg-0333e92f36b2aa0bf	Security Group ID for Instances Inside Interface	-	
InsideSubnetIds	subnet-047c0a2beffb5a70f	Inside subnet ID	-	
InstanceSGId	sg-0c0c6bfb5ba5f1c10	Security Group ID for Instances Management Interface	-	
LambdaSecurityGroupId	sg-01771b0d3012a40c5	Security Group ID for Lambda Functions	-	
LambdaSubnetIds	subnet-0fb24785c687d50e4,subnet-0f1996a02ffaa2e62	List of lambda subnet IDs (comma seperated)	-	
MgmtSubnetIds	subnet-02d4a757b95a9a5b9	Mangement subnet ID	-	
UseGWLB	Yes	Use Gateway Load Balancer	-	
VpcName	vpc-003b592ad2518d03d	Name of the VPC created	-	

Étape 3

Chargez **cluster_layer.zip**, **cluster_lifecycle.zip** et **cluster_manager.zip** dans le compartiment S3 créé par **infrastructure.yaml**.

Étape 4

Déployez **deploy_asav_clustering.yaml** :

- Allez sur **CloudFormation** et cliquez sur **Create stack** (Créer une pile); sélectionnez **With new resources (standard)** (Avec de nouvelles ressources (standard)).
- Cliquez sur **Upload a template file** (Charger un fichier modèle), cliquez sur **Choose file** (Choisir un fichier) et sélectionnez **deploy_asav_clustering.yaml** dans le dossier cible.
- Cliquez sur **Next** (suivant) et fournissez les informations requises.
- Cliquez sur **Next** (Suivant), puis sur **Create stack** (Créer une pile).

Illustration 3 : Ressources déployées

Resources (21)					
Search resources					
Logical ID ▲	Physical ID ▼	Type ▼	Status ▼	Status reason	
ASAvGroup	ran-clt-1 🔗	AWS::AutoScaling::AutoScalingGroup	✔️ CREATE_COMPLETE	-	
ASAvLaunchTemplate	lt-056fd20764270c893 🔗	AWS::EC2::LaunchTemplate	✔️ CREATE_COMPLETE	-	
CLSmanagerTopic	arn:aws:sns:sa-east-1:797661843114:ran-clt-1-cluster-manager-topic 🔗	AWS::SNS::Topic	✔️ CREATE_COMPLETE	-	
ClusterManager	ran-clt-1-manager-lambda 🔗	AWS::Lambda::Function	✔️ CREATE_COMPLETE	-	
ClusterManagerLogGrp	/aws/lambda/ran-clt-1-manager-lambda 🔗	AWS::Logs::LogGroup	✔️ CREATE_COMPLETE	-	
ClusterManagerSNS1	arn:aws:sns:sa-east-1:797661843114:ran-clt-1-cluster-manager-topic:e13bfc0-d698-4215-88a5-278474e22c32	AWS::SNS::Subscription	✔️ CREATE_COMPLETE	-	
ClusterManagerSNS1Permission	ran-clt-ClusterManagerSNS1Permission-S6BQAE05OG6U	AWS::Lambda::Permission	✔️ CREATE_COMPLETE	-	
InstanceEvent	ran-clt-1-notify-instance-event 🔗	AWS::Events::Rule	✔️ CREATE_COMPLETE	-	
InstanceEventInvokeLambdaPermission	ran-clt-InstanceEventInvokeLambdaPermission-1XPS21Q4G2DY6	AWS::Lambda::Permission	✔️ CREATE_COMPLETE	-	

L'état passe de CREATE_IN_PROGRESS à CREATE COMPLETE indiquant le déploiement réussi.

Étape 5

Vérifiez le déploiement de la grappe en vous connectant à l'un des nœuds et en saisissant la commande **show cluster info**.

show cluster info

```
Cluster oneclicktest-cluster: On
Interface mode: individual
Cluster Member Limit : 16
This is "200" in state CONTROL_NODE
ID : 0
Version : 9.19.1
Serial No.: 9AU42EN5D1E
CCL IP : 1.1.1.200
CCL MAC : 4201.0a0a.0fc7
Module : ASAv
Resource : 4 cores / 8192 MB RAM
Last join : 15:26:22 UTC Jul 17 2022
Last leave: N/A
Other members in the cluster:
Unit "204" in state DATA_NODE
ID : 1
Version : 9.19.1
Serial No.: 9AJ9N46947R
CCL IP : 1.1.1.204
CCL MAC : 4201.0a0a.0fcb
Module : ASAv
Resource : 4 cores / 8192 MB RAM
```

Last join : 16:57:42 UTC Jul 17 2022
 Last leave: 16:03:25 UTC Jul 17 2022

Déployer manuellement la grappe dans AWS

Pour déployer la grappe manuellement, préparez la configuration Day0 et déployez chaque nœud.

Créer la configuration Day0 pour AWS

Fournissez la configuration de démarrage pour chaque nœud de grappe à l'aide des commandes suivantes :

Exemple d'équilibreur de charge de passerelle

Dans l'exemple suivant, une configuration en cours d'exécution est créée pour un équilibreur de charges de passerelle avec une interface Geneve pour le trafic en demi-tour et une interface VXLAN pour la liaison de commande de grappe.

```
cluster interface-mode individual force
policy-map global_policy
class inspection_default
no inspect h323 h225
no inspect h323 ras
no inspect rtsp
no inspect skinny

int m0/0
management-only
nameif management
security-level 100
ip address dhcp setroute
no shut

interface TenGigabitEthernet0/0
nameif geneve-vtep-ifc
security-level 0
ip address dhcp
no shutdown

interface TenGigabitEthernet0/1
nve-only cluster
nameif ccl_link
security-level 0
ip address dhcp
no shutdown

interface vni1
description Clustering Interface
segment-id 1
vtep-nve 1

interface vni2
proxy single-arm
nameif ge
security-level 0
vtep-nve 2

object network ccl_link
```

```

range 10.1.90.4 10.1.90.254 //Mandatory user input, use same range on all nodes
object-group network cluster_group
network-object object ccl_link
nve 2
encapsulation geneve
source-interface geneve-vtep-ifc
nve 1
encapsulation vxlan
source-interface ccl_link
peer-group cluster_group

cluster group asav-cluster // Mandatory user input, use same cluster name on all nodes
local-unit 1 //Value in bold here must be unique to each node
cluster-interface vnil ip 1.1.1.1 255.255.255.0 //Value in bold here must be unique to each
node
priority 1
enable noconfirm

mtu geneve-vtep-ifc 1806
mtu ccl link 1960
aaa authentication listener http geneve-vtep-ifc port 7575 //Use same port number on all
nodes
jumbo-frame reservation
wr mem

```

**Remarque**

Pour les paramètres de vérification de l'intégrité d'AWS, assurez-vous de spécifier le port **aaa authentication listener http** que vous définissez ici.

Exemple d'équilibreur de charge non natif

Dans l'exemple suivant, une configuration à utiliser avec des équilibreurs de charges non natifs avec des interfaces de gestion, interne et externe est créée, ainsi qu'une interface VXLAN pour la liaison de commande de grappe.

```

cluster interface-mode individual force
interface Management0/0
management-only
nameif management
ip address dhcp

interface GigabitEthernet0/0
no shutdown
nameif outside
ip address dhcp

interface GigabitEthernet0/1
no shutdown
nameif inside
ip address dhcp

interface GigabitEthernet0/2
nve-only cluster
nameif ccl_link
ip address dhcp
no shutdown

interface vnil
description Clustering Interface
segment-id 1

```

```

vtep-nve 1

jumbo-frame reservation
mtu ccl_link 1654
object network ccl_link
range 10.1.90.4 10.1.90.254 //mandatory user input
object-group network cluster_group
network-object object ccl_link

nve 1
encapsulation vxlan
source-interface ccl_link
peer-group cluster_group

cluster group asav-cluster //mandatory user input
local-unit 1 //mandatory user input
cluster-interface vni1 ip 10.1.1.1 255.255.255.0 //mandatory user input
priority 1
enable

```

**Remarque**

Si vous copiez et collez la configuration donnée ci-dessus, veuillez à supprimer **//entrée utilisateur obligatoire** de la configuration.

Déployer les nœuds de la grappe

Déployez les nœuds de la grappe pour former une grappe.

Procédure

Étape 1 Déployez l'instance ASA virtuel en utilisant la configuration de jour 0 de la grappe avec le nombre d'interfaces requis (trois interfaces si vous utilisez l'équilibreur de charges de passerelle (GWLB) ou quatre interfaces si vous utilisez un équilibreur de charges non natif). Pour ce faire, dans la section **Configure Instance Details** (Configurer les détails de l'instance) > Advanced Details (détails avancés), collez la configuration du jour 0 de la grappe.

Remarque

Assurez-vous d'associer des interfaces aux instances dans l'ordre indiqué ci-dessous.

- Équilibreur de charges de passerelle AWS : trois interfaces : liaison de gestion, externe et de commande de grappe.
- Équilibreurs de charges non natifs : quatre interfaces : liaison de gestion, interne, externe et de commande de grappe.

Pour en savoir plus sur le déploiement d'ASA virtuel sur AWS, consultez [Déployer l'ASA virtuel sur AWS](#).

Étape 2 Répétez l'étape 1 pour déployer le nombre requis de nœuds supplémentaires.

Étape 3 Utilisez la commande **show cluster info** de la console ASA virtuel pour vérifier si tous les nœuds ont bien rejoint la grappe.

Étape 4 Configurez l'équilibreur de charges de passerelle AWS

- Créez un groupe cible et un GWLB.

- b) Associez le groupe cible au GWLB.

Remarque

Assurez-vous de configurer le GWLB pour utiliser les paramètres de groupe de sécurité, de configuration d'écouteur et de vérification de l'intégrité adéquats.

- c) Enregistrez l'interface de données (interface interne) avec le groupe cible à l'aide des adresses IP. Pour en savoir plus, consultez [Créer un équilibreur de charges de passerelle](#).

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, dans le cadre de la configuration day0 ou après le déploiement de la grappe.

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.

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. La surveillance de l'intégrité n'est pas effectuée sur les sous-interfaces VLAN. 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*]

Pour déterminer l'intégrité des nœuds, les nœuds de la grappe ASA envoient aux autres nœuds des messages de pulsation sur la liaison de commande de la grappe. Si un nœud ne reçoit aucun message de pulsation d'un nœud homologue au cours de la période de rétention, le nœud homologue est considéré comme ne répondant pas ou comme étant inactif.

- **holdtime** *timeout* : détermine l'intervalle de temps entre les messages d'état de pulsation des nœuds, entre 0,3 et 45 secondes; la valeur par défaut est de 3 secondes.

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), vous devez désactiver la fonctionnalité de 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. Le délai avant la suppression par l'ASA d'un membre de la grappe dépend du fait que le nœud est un membre établi ou qu'il rejoint 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 supervision d'une interface. La surveillance de l'intégrité n'est pas effectuée sur les sous-interfaces VLAN. 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), vous devez désactiver la fonctionnalité de 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.

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. 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; désactive la surveillance sur l'interface Management 0/0; 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
ciscoasa(cfg-cluster)# no health-check monitor-interface management0/0
ciscoasa(cfg-cluster)# health-check data-interface auto-rejoin 4 2 3
ciscoasa(cfg-cluster)# health-check cluster-interface auto-rejoin 6 2 1
```

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.

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 un nœud de données à partir du nœud de contrôle

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.

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.

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.

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.

Procédure

Définissez un nouveau nœud comme nœud de contrôle :

cluster control-node unit*node_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.

Surveillance de la grappe

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.: P30000000025
    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 unicast 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 une 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) 0xclce5c49
access-list 101 line 8 extended permit tcp host 192.168.1.177 host 192.168.43.192
(hitcnt=3, 0, 1, 1, 1) 0xb6f59512
access-list 101 line 9 extended permit tcp host 192.168.1.177 host 192.168.43.44
(hitcnt=0, 0, 0, 0, 0) 0xdc104200
access-list 101 line 10 extended permit tcp host 192.168.1.112 host 192.168.43.44
(hitcnt=429, 109, 107, 109, 104)
0xce4f281d
access-list 101 line 11 extended permit tcp host 192.168.1.170 host 192.168.43.238
(hitcnt=3, 1, 0, 0, 2) 0x4143a818
access-list 101 line 12 extended permit tcp host 192.168.1.170 host 192.168.43.169
(hitcnt=2, 0, 1, 0, 1) 0xb18dfea4
access-list 101 line 13 extended permit tcp host 192.168.1.170 host 192.168.43.229
(hitcnt=1, 1, 0, 0, 0) 0x21557d71
access-list 101 line 14 extended permit tcp host 192.168.1.170 host 192.168.43.106
(hitcnt=0, 0, 0, 0, 0) 0x7316e016
access-list 101 line 15 extended permit tcp host 192.168.1.170 host 192.168.43.196
(hitcnt=0, 0, 0, 0, 0) 0x013fd5b8
access-list 101 line 16 extended permit tcp host 192.168.1.170 host 192.168.43.75
(hitcnt=0, 0, 0, 0, 0) 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.

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

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
- Routage multidiffusion

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.

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.

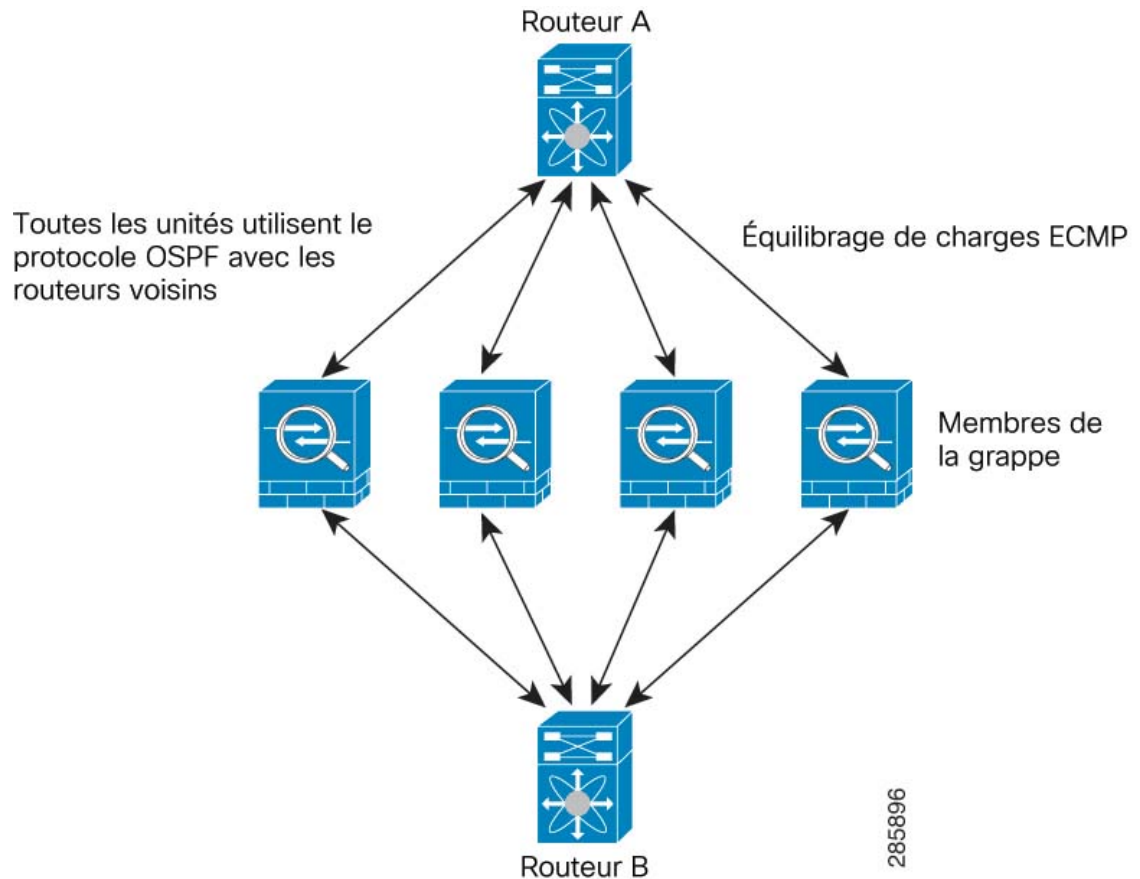
Paramètres de connexion et mise en grappe

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.

Routage et mise en grappe dynamiques

En mode d'interface individuel, chaque nœud exécute le protocole de routage en tant que routeur autonome, et les routes sont apprises par chaque nœud indépendamment.

Illustration 4 : Routage dynamique en mode d'interface individuelle



Dans le diagramme ci-dessus, le routeur A détecte qu'il existe quatre chemins à coûts égaux vers le routeur B, chacun passant par un nœud. ECMP est utilisé pour équilibrer la charge du trafic entre les quatre chemins. Chaque nœud choisit un ID de routeur différent lorsqu'il communique avec des routeurs externes.

Vous devez configurer un groupement de grappes pour l'ID de routeur afin que chaque nœud ait un ID de routeur distinct.

Le protocole EIGRP ne forme pas de relations de voisinage avec les homologues de la grappe en mode d'interface individuelle.



Remarque

Si la grappe comporte plusieurs contiguïtés avec le même routeur à des fins de redondance, le routage dissymétrique peut entraîner une perte de trafic inacceptable. Pour éviter le routage dissymétrique, regroupez toutes ces interfaces de nœud dans la même zone de trafic. Voir [Configurer une zone de trafic](#).

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

En mode d'interface individuel, les unités n'agissent pas indépendamment avec la multidiffusion. Toutes les données et les paquets de routage sont traités et transmis par l'unité de contrôle, évitant ainsi la duplication des paquets.

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 :

- No Proxy ARP (Pas de serveur mandataire ARP) : pour les interfaces individuelles, une réponse de serveur mandataire ARP n'est jamais envoyée pour les adresses mappées. Cela empêche le routeur adjacent de maintenir une relation d'homologue avec un ASA qui ne fait plus partie de la grappe. Le routeur en amont a besoin d'une route statique ou d'un PBR avec suivi d'objets pour les adresses mappées qui pointe vers l'adresse IP de la grappe principale. Ce n'est pas un problème pour un EtherChannel étendu, car il n'y a qu'une seule adresse IP associée à l'interface de grappe.
- 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.

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 et 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 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.

Pour les connexions à une interface individuelle lors de l'utilisation de PBR ou d'ECMP, vous devez toujours vous connecter à l'adresse IP de la grappe principale, et non à une adresse locale.

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 élient leurs propres nœuds de contrôle. Par exemple, si un routeur tombe en panne entre deux emplacements de grappe, le nœud de contrôle d'origine à l'emplacement 1 supprimera les nœuds de données de l'emplacement 2 de la grappe. Pendant ce temps, les nœuds de l'emplacement 2 élient leur propre nœud de contrôle et formeront leur propre grappe. Notez que le trafic symétrique peut échouer dans ce scénario. Une fois la liaison de commande de grappe restaurée, le nœud de contrôle qui a la priorité la plus élevée conservera le rôle de nœud de contrôle.

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.

Lorsque vous activez la surveillance de l'intégrité, toutes les interfaces physiques sont supervisées par défaut; vous pouvez éventuellement désactiver la surveillance par interface. Seules les interfaces nommées peuvent être surveillées.

Un nœud est supprimé de la grappe en cas de défaillance de ses interfaces surveillées. Le délai avant la suppression par l'ASA d'un membre de la grappe dépend du fait que le nœud est un membre établi ou qu'il rejoint la grappe. L'ASA ne supervise 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 l'ASA de la grappe. Le nœud est supprimé après 500 ms, quel que soit son état.

État après l'échec

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 de l'interface de ligne de commande 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 de l'interface de ligne de commande 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.

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 2 : Fonctionnalités répliquées dans la grappe

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

Gestion des connexions par la grappe

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

Rôles de connexion

Consultez les rôles suivants, définis pour chaque connexion :

- Propriétaire : généralement, le nœud qui reçoit initialement la connexion. Le propriétaire gère l'état TCP et traite les paquets. Une connexion n'a qu'un seul propriétaire. Si le propriétaire d'origine échoue, lorsque les nouveaux nœuds reçoivent des paquets de la connexion, le directeur choisit un nouveau propriétaire dans ces nœuds.

- Propriétaire du sauvegarde : nœud qui stocke les informations d'état TCP/UDP reçues du propriétaire, de sorte que la connexion puisse être transférée en toute transparence à un nouveau propriétaire en cas de défaillance. Le propriétaire de secours ne prend pas en charge la connexion en cas d'échec. Si le propriétaire devient indisponible, le premier nœud à recevoir les paquets de la connexion (selon l'équilibrage de la charge) contacte le propriétaire de secours pour obtenir les informations d'état pertinentes, afin qu'il puisse devenir le nouveau propriétaire.

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

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

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

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

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

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

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

Détails du hachage ICMP/ICMPv6 :

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

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

**Remarque**

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

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

Connexions par traduction d'adresse de port

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

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

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

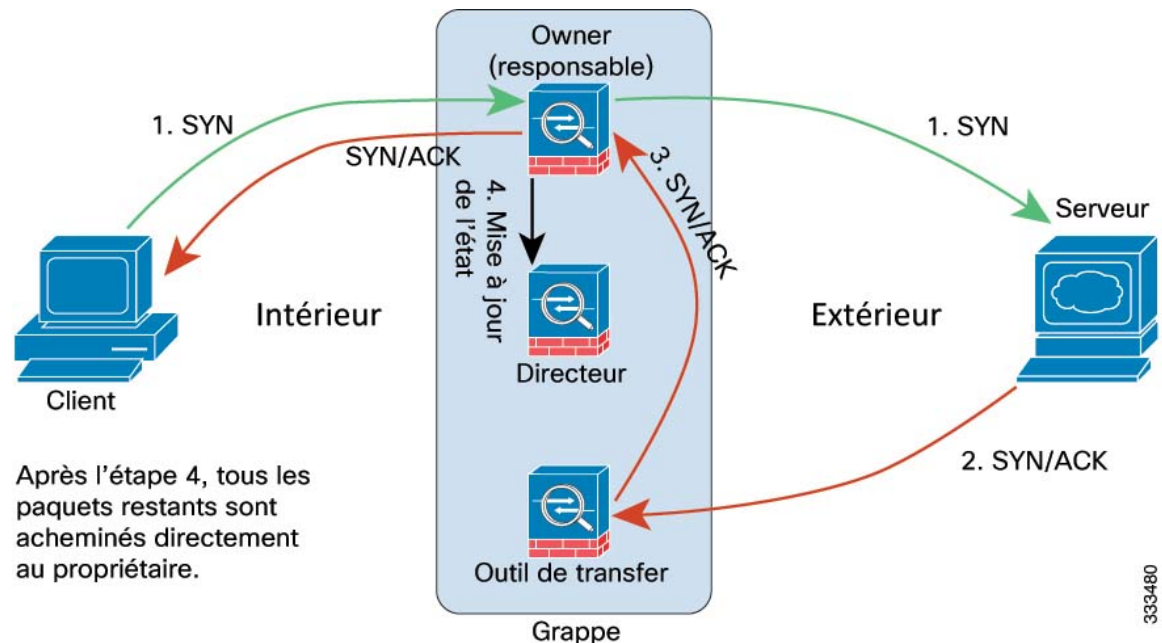
Vous pouvez modifier les valeurs par défaut de PAT par session pour les protocoles TCP et UDP afin que les connexions à ces protocoles soient gérées par session ou multisession selon la configuration. Pour ICMP, vous ne pouvez pas modifier la PAT multisession par défaut. Référez-vous au guide de configuration 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.



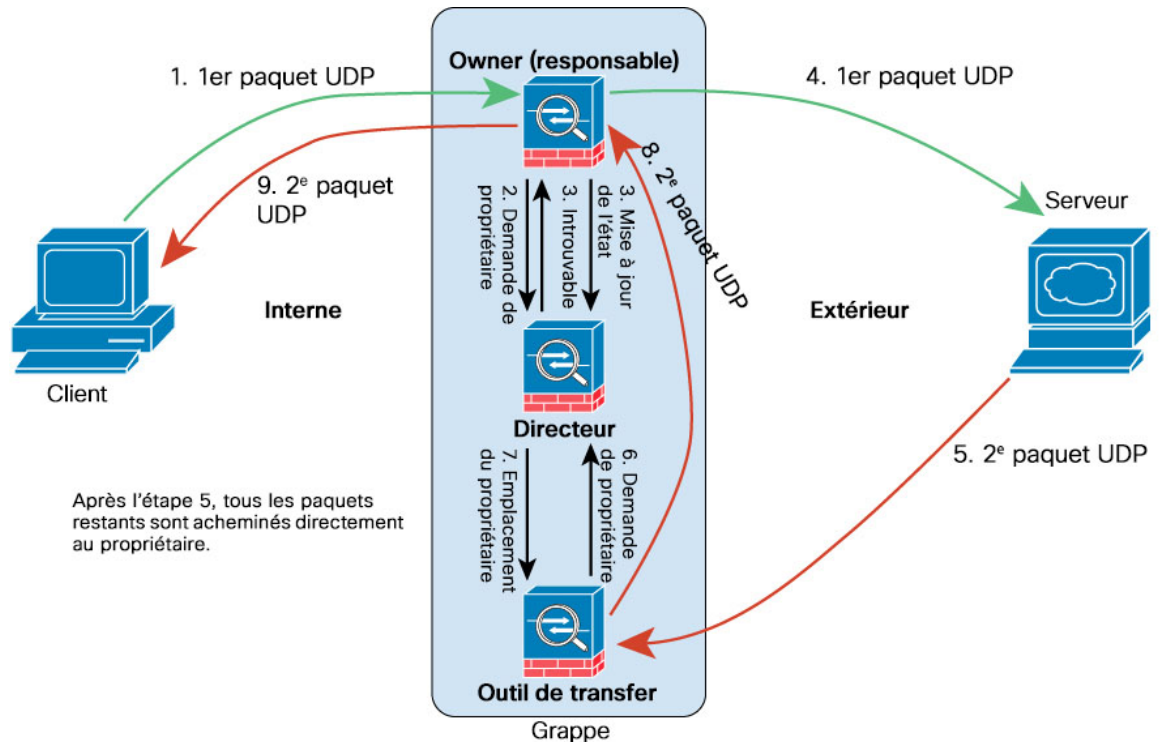
1. Le paquet SYN provient du client et est livré à un ASA (selon la méthode d'équilibrage de la charge), qui devient le propriétaire. Le propriétaire crée un flux, code les renseignements sur le propriétaire dans un témoin SYN et transfère le paquet au serveur.
2. Le paquet SYN-ACK provient du serveur 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 5 : 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.

Historique de la mise en grappe d'ASA virtuel dans le nuage public

Fonctionnalités	Version	Détails
Mise en grappe d'ASA virtuel sur Amazon Web Services (AWS)	9.19(1)	L'ASA virtuel prend en charge la mise en grappe d'interfaces individuelles pour un maximum de 16 nœuds sur AWS. Vous pouvez utiliser la mise en grappe avec ou sans équilibreur de charge de la passerelle AWS.

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.