



Routage basé sur une stratégie

Ce chapitre décrit comment configurer l'ASA pour prendre en charge le routage basé sur les politiques (BPR). Les sections suivantes décrivent le routage basé sur les politiques, les lignes directrices relatives au PBR et la configuration du PBR.

- [À propos du routage basé sur les politiques, à la page 1](#)
- [Lignes directrices pour le routage basé sur des politiques, à la page 3](#)
- [Configurer le routage basé sur les politiques, à la page 5](#)
- [Exemples de routage basé sur des politiques, à la page 9](#)
- [Historique du routage basé sur des politiques, à la page 18](#)

À propos du routage basé sur les politiques

Le routage traditionnel est basé sur la destination, ce qui signifie que les paquets sont acheminés en fonction de l'adresse IP de destination. Cependant, il est difficile de modifier le routage d'un trafic spécifique dans un système de routage basé sur la destination. Avec le routage basé sur les politiques (PBR), vous pouvez définir le routage en fonction de critères autres que le réseau de destination. Le PBR vous permet d'acheminer le trafic en fonction de l'adresse source, du port source, de l'adresse de destination, du port de destination, du protocole ou d'une combinaison de ces éléments.

Routage basé sur les politiques :

- Vous permet de fournir une qualité de service (QoS) à un trafic différencié.
- Vous permet de répartir le trafic interactif et par lots sur des chemins permanents à faible bande passante et à faible coût et des chemins commutés à bande passante élevée et à coût élevé.
- Permet aux fournisseurs de services Internet et à d'autres organisations d'acheminer le trafic provenant de divers ensembles d'utilisateurs par le biais de connexions Internet bien définies.

Le routage basé sur les politiques peut mettre en œuvre la QoS en classant et en marquant le trafic à la périphérie du réseau, puis en utilisant le PBR dans tout le réseau pour acheminer le trafic marqué le long d'un chemin spécifique. Cela permet le routage de paquets provenant de différentes sources vers différents réseaux, même lorsque les destinations sont les mêmes, et cela peut être utile pour interconnecter plusieurs réseaux privés.

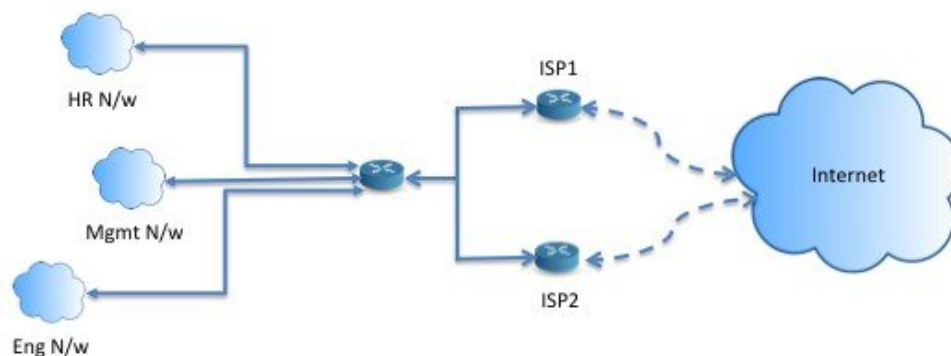
Pourquoi utiliser le routage basé sur des politiques?

Imaginez une entreprise qui dispose de deux liaisons entre des sites : l'une, une liaison onéreuse à bande passante élevée et à faible délai, et l'autre, à faible bande passante, avec un délai plus élevé et un moindre coût. Lors de l'utilisation de protocoles de routage traditionnels, la liaison à plus grande largeur de bande recevrait la majeure partie, voire la totalité, du trafic qui y est envoyé, en fonction des économies de métriques obtenues grâce aux caractéristiques de la bande passante et/ou du délai (avec EIGRP ou OSPF) de la liaison. PBR vous permet d'acheminer le trafic de priorité supérieure sur la liaison à bande passante élevée/faible délai, tout en envoyant tout le reste du trafic sur la liaison à bande passante faible/délai élevé.

Certaines applications du routage basé sur les politiques sont les suivantes :

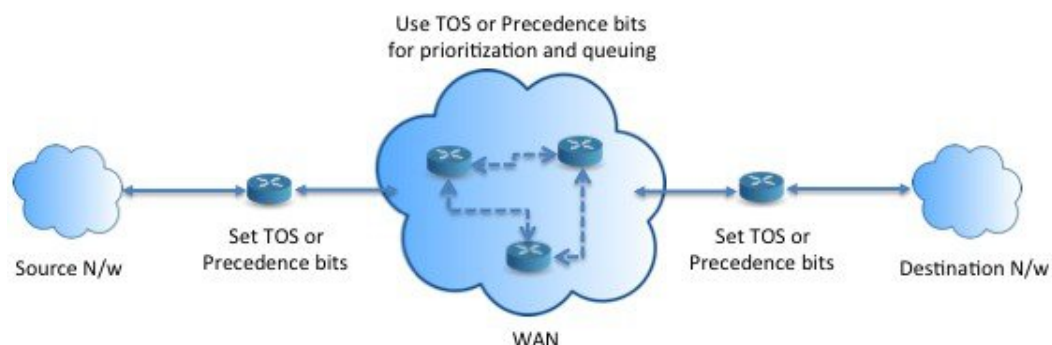
Routage à accès égal et sensible à la source

Dans cette topologie, le trafic provenant du réseau RH et du réseau de gestion peut être configuré pour passer par FAI1, et le trafic du réseau des ingénieurs peut être configuré pour passer par FAI2. Ainsi, le routage basé sur les politiques permet aux administrateurs réseau de fournir un routage à accès égal et sensible à la source, comme indiqué ici.



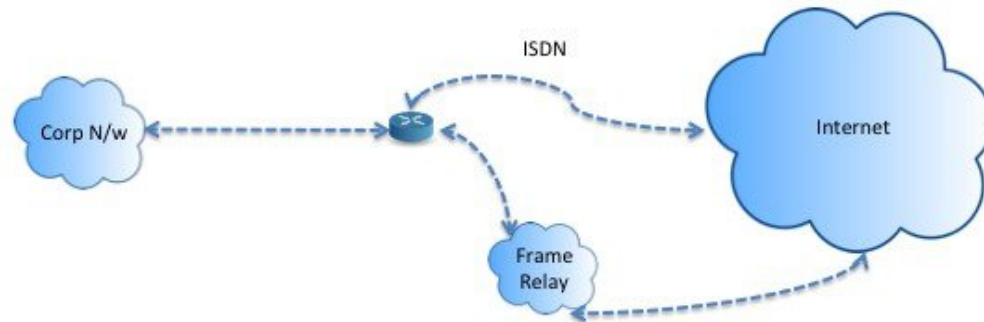
Qualité de service

En balisant les paquets avec le routage basé sur les politiques, les administrateurs réseau peuvent classer le trafic réseau au périmètre du réseau pour différentes classes de service, puis mettre en œuvre ces classes de service au cœur du réseau à l'aide de la mise en file d'attente de priorité, personnalisée ou pondérée (comme indiqué dans la figure ci-dessous). Cette configuration améliore les performances du réseau en éliminant le besoin de classer le trafic explicitement sur chaque interface WAN au cœur du réseau fédérateur.



Réduction des coûts

Une organisation peut diriger le trafic en bloc associé à une activité spécifique pour utiliser un lien à haute bande passante à coût élevé pendant une durée courte et maintenir la connectivité de base sur un lien à faible coût à bande passante inférieure pour le trafic interactif en définissant la topologie, comme indiqué ici.



Partage de la charge

En plus des fonctionnalités de partage dynamique de charge offertes par l'équilibrage de charge ECMP, les administrateurs réseau peuvent désormais mettre en œuvre des politiques pour répartir le trafic entre plusieurs chemins en fonction des caractéristiques du trafic.

Par exemple, dans la topologie décrite dans le scénario de routage à accès égalitaire sensible, un administrateur peut configurer le routage basé sur une politique pour partager la charge du trafic provenant du réseau des Ressources humaines par ISP1 et du trafic provenant du réseau Eng par ISP2.

Mise en œuvre de PBR

L'ASA utilise les ACL pour faire correspondre le trafic et effectuer des actions de routage sur le trafic. Plus précisément, vous configurez une carte de routage qui spécifie une ACL pour la correspondance, puis vous spécifiez une ou plusieurs actions pour ce trafic. Enfin, vous associez la carte de routage à une interface sur laquelle vous souhaitez appliquer le PBR à tout le trafic entrant.



Remarque

Avant de procéder à la configuration, assurez-vous que le trafic d'entrée et de sortie de chaque session traverse la même interface destinée au fournisseur de services Internet pour éviter les comportements imprévus causés par le routage asymétrique, en particulier lorsque la NAT et le VPN sont utilisés.

Lignes directrices pour le routage basé sur des politiques

Mode pare-feu

Pris en charge uniquement en mode pare-feu routé. Le mode pare-feu transparent n'est pas pris en charge.

Routage par flux

Étant donné que l'ASA effectue le routage par flux, le routage des politiques est appliqué sur le premier paquet et la décision de routage qui en résulte est stockée dans le flux créé pour le paquet. Tous les paquets suivants appartenant à la même connexion correspondent simplement à ce flux et sont acheminés de manière appropriée.

Politiques PBR non appliquées pour la recherche de route de sortie

Le routage basé sur des règles est une fonction d'entrée uniquement, c'est-à-dire qu'il n'est appliqué qu'au premier paquet d'une nouvelle connexion entrante, et c'est à ce moment-là que l'interface de sortie est sélectionnée pour le tronçon d'aller de la connexion. Notez que PBR ne sera pas déclenché si le paquet entrant appartient à une connexion existante ou si la NAT est appliquée et que cette dernière choisit l'interface de sortie.

Politiques PBR non appliquées pour le trafic amorce



Remarque

Il y a connexion amorce lorsque l'établissement de liaison nécessaire entre la source et la destination n'a pas lieu.

Lorsqu'une nouvelle interface interne est ajoutée et qu'une nouvelle politique VPN est créée à l'aide d'un groupement d'adresses unique, le PBR est appliqué à l'interface externe correspondant à la source du nouveau groupement de clients. Ainsi, PBR envoie le trafic du client au prochain saut sur la nouvelle interface. Cependant, PBR n'est pas impliqué dans le trafic de retour d'un hôte qui n'a pas encore établi de connexion avec les nouvelles routes d'interface interne vers le client. Ainsi, le trafic de retour de l'hôte vers le client VPN, en particulier la réponse du client VPN, est abandonné car il n'y a aucune voie de routage valide. Vous devez configurer une voie de routage statique pondérée avec une métrique plus élevée sur l'interface interne.

Mise en grappe

- La mise en grappe est prise en charge.
- Dans un scénario de grappe, sans routes statiques ou dynamiques, et avec la commande `ip-verify-reverse path` activée, le trafic asymétrique peut être abandonné. Il est donc conseillé de désactiver `ip-verify-reverse path`.

Prise en charge d'IPv6

IPv6 est pris en charge.

Lignes directrices relatives à la surveillance des chemins d'accès

Voici les lignes directrices pour configurer la surveillance des chemins d'accès sur les interfaces :

- Les interfaces doivent avoir un nom d'interface.
- Les interfaces de gestion uniquement ne peuvent pas être configurées avec la surveillance des chemins d'accès. Pour configurer la surveillance des chemins d'accès, vous devez décocher la case **Dedicate this interface to management only** (Dédier cette interface à la gestion uniquement).
- La surveillance des chemins d'accès n'est pas prise en charge sur les périphériques en mode de système transparent ou de contexte multiple.

- Les types de surveillance automatique (auto, auto4 et auto6) ne sont pas pris en charge pour les interfaces de tunnel.
- La surveillance des chemins d'accès ne peut pas être configurée pour les interfaces suivantes :
 - BVI
 - Boucle avec retour
 - DVTI

Directives supplémentaires

- Toutes les restrictions et limites de configuration existantes liées aux cartes de routage seront conservées.
- N'utilisez pas de cartes de routage contenant des listes de correspondance de politiques pour le routage basé sur des politiques. La liste de politiques de correspondance est uniquement utilisée pour le BGP.
- Le transfert de chemin inverse de monodiffusion (uRPF) valide l'adresse IP source des paquets reçus sur une interface par rapport à la table de routage et non par rapport à la carte de routage PBR. Lorsque uRPF est activé, les paquets reçus sur une interface par l'intermédiaire de PBR sont abandonnés tels qu'ils sont, sans l'entrée de route spécifique. Par conséquent, lorsque vous utilisez PBR, assurez-vous de désactiver uRPF.

Configurer le routage basé sur les politiques

Une carte de routage est composée d'une ou de plusieurs instructions route-map. Chaque instruction a un numéro de séquence, ainsi qu'une clause allow ou deny. Chaque instruction route-map contient des commandes match et set. La commande match indique les critères de correspondance à appliquer au paquet. La commande set indique l'action à entreprendre sur le paquet.

- Lorsqu'une carte de routage est configurée avec des clauses match/set IPv4 et IPv6 ou lorsqu'une ACL unifiée faisant correspondre le trafic IPv4 et IPv6 est utilisée, les actions définies seront appliquées en fonction de la version IP de destination.
- Lorsque plusieurs interfaces ou sauts suivants sont configurés en tant qu'action définie, toutes les options sont évaluées l'une après l'autre jusqu'à ce qu'une option utilisable valide soit trouvée. Aucun équilibrage de charge ne sera effectué parmi les multiples options configurées.
- L'option verify-availability n'est pas prise en charge en mode de contexte multiple.

Procédure

Étape 1

Définissez une liste d'accès standard ou étendue :

```
access-list name standard {permit | deny} {any4 | host ip_address | ip_address mask}
```

```
access-list name extended {permit | deny} protocol source_and_destination_arguments
```

Exemple :

```
ciscoasa(config)# access-list testacl extended permit ip
10.1.1.0 255.255.255.0 10.2.2.0 255.255.255.0
```

Si vous utilisez une liste de contrôle d'accès standard, la mise en correspondance s'effectue uniquement pour l'adresse de destination. Si vous utilisez une liste de contrôle d'accès étendue, vous pouvez mettre en correspondance la source, la destination ou les deux.

Pour la liste de contrôle d'accès étendue, vous pouvez spécifier les paramètres IPv4, IPv6, Identity Firewall ou Cisco TrustSec. Vous pouvez également inclure des objets de service réseau. Pour obtenir la syntaxe complète, consultez la référence de commande ASA.

Étape 2

Créez une entrée de carte de routage :

route-map *nom* {**permis** | **refuser**} [*numéro_séquence*]

Exemple :

```
ciscoasa(config)# route-map testmap permit 12
```

Les entrées de carte de routage sont lues dans l'ordre. Vous pouvez identifier l'ordre à l'aide de l'argument *sequence_number* ou l'ASA utilise l'ordre dans lequel vous ajoutez les entrées de la carte de routage.

La liste de contrôle d'accès comprend également ses propres instructions d'autorisation et de refus. Pour les correspondances Autoriser/Autoriser entre la carte de routage et la liste de contrôle d'accès, le traitement de routage basé sur les politiques se poursuit. Pour les correspondances Autoriser/refuser, le traitement se termine pour cette carte de routage et les autres cartes de routage sont vérifiées. Si le résultat est toujours Autoriser/refuser, la table de routage normale est utilisée. Pour les correspondances refus/refus, le traitement de routage basé sur les politiques se poursuit.

Remarque

Lorsqu'une carte de routage est configurée sans action Allow ou Deny (autoriser ou refuser) et sans numéro de séquence, elle suppose par défaut que l'action est l'autorisation et le numéro de séquence est 10.

Étape 3

Définissez les critères de correspondance à appliquer à l'aide d'une liste d'accès :

match ip address *access-list_name* [*access-list_name...*]

Exemple :

```
ciscoasa(config-route-map)# match ip address testacl
```

Remarque

Vérifiez que la liste d'accès ne contient aucune règle inactive. Vous ne pouvez pas définir une correspondance ACL avec les règles inactives sur un PBR.

Étape 4

Configurez une ou plusieurs actions définies :

- Définissez l'adresse de saut suivant :

set {ip | ipv6} next-hop *ipv4_or_ipv6_address*

Vous pouvez configurer plusieurs adresses IP de saut suivant, auquel cas elles sont évaluées dans l'ordre spécifié jusqu'à ce qu'une adresse IP de saut suivant routable soit trouvée. Les prochains sauts configurés doivent être connectés directement; Sinon, l'action définie ne sera pas appliquée.

- Définissez l'adresse de saut suivant :

set {ip | ipv6} default next-hop *ipv4_or_ipv6_address*

Si la recherche normale de route échoue pour le trafic correspondant, l'ASA transfère le trafic en utilisant l'adresse IP de saut suivant spécifiée.

- Définir une adresse IPv4 du saut suivant récursive :

set ip next-hop recursive *ip_address*

Les deux commandes **set ip next-hop** et **set ip default next-hop** exigent que le prochain saut se trouve sur un sous-réseau directement connecté. Avec **set ip next-hop recursive**, l'adresse de saut suivant n'a pas besoin d'être connectée directement. Au lieu de cela, une recherche récursive est effectuée sur l'adresse de saut suivant, et le trafic correspondant est transmis au saut suivant utilisé par cette entrée de routage en fonction du chemin de routage utilisé sur le routeur.

- Vérifiez si les prochains sauts IPv4 d'une carte de routage sont disponibles :

set ip next-hop verify-availability *next-hop-address sequence_number track object*

Vous pouvez configurer un objet de suivi de moniteur SLA pour vérifier l'accessibilité du prochain saut. Pour vérifier la disponibilité de plusieurs prochains sauts, plusieurs commandes **set ip next-hop verify-availability** peuvent être configurées avec différents numéros de séquence et différents objets de suivi.

- Définissez l'interface de sortie du paquet :

set interface *interface_name*

ou

set interface null0

Cette commande configure l'interface par laquelle le trafic correspondant est transféré. Vous pouvez configurer plusieurs interfaces, auquel cas elles sont évaluées dans l'ordre spécifié jusqu'à ce qu'une interface valide soit trouvée. Lorsque vous spécifiez **null0**, tout le trafic correspondant à la carte de routage sera abandonné. Il doit y avoir une voie de routage pour la destination qui peut être acheminée par l'interface spécifiée (statique ou dynamique).

- Définissez l'interface de sortie en fonction de son coût :

set adaptive-interface cost *interface_list*

L'interface de sortie est sélectionnée dans la liste des interfaces séparées par des espaces. Si les coûts des interfaces sont les mêmes, il s'agit d'une configuration active-active et les paquets sont équilibrés (round robin) sur les interfaces de sortie. Si les coûts sont différents, l'interface avec le coût le plus faible est sélectionnée. Les interfaces ne sont prises en compte que si elles sont actives. Par exemple :

```
set adaptive-interface cost output1 output2
```

- Définissez l'interface par défaut sur null0 :

set default interface null0

Si une recherche de route normale échoue, l'ASA transfère le trafic null0 et le trafic sera abandonné.

- Définissez la valeur de bit Ne pas fragmenter (DF) dans l'en-tête IP :

set ip df {0|1}

- Classez le trafic IP en définissant un point de code de services différentiels (DSCP) ou une valeur de priorité IP dans le paquet :

```
set {ip | ipv6} dscp new_dscp
```

Remarque

Lorsque plusieurs actions de définition sont configurées, l'ASA les évalue dans l'ordre suivant : **set ip next-hop verify-availability; set ip next-hop; set ip next-hop recursive; set interface; set adaptive-interface cost; set ip default next-hop; set default interface.**

Étape 5 Configurez une interface et passez en mode de configuration d'interface :

```
interface interface_id
```

Exemple :

```
ciscoasa(config)# interface GigabitEthernet0/0
```

Étape 6 Si vous utilisez **set adaptive-interface cost** comme critères dans la carte de routage, définissez le coût sur l'interface :

```
policy-route cost value
```

La valeur peut être 1 à 65535. La valeur par défaut est 0, et vous pouvez la réinitialiser en utilisant la version **no** de la commande. Plus le numéro de priorité est faible, plus la priorité est élevée. Par exemple, 1 a priorité sur 2.

Lorsque vous définissez le coût de routage de politique et que vous utilisez la commande **set adaptive-interface cost** dans la carte de routage, le trafic de sortie fait l'objet d'un équilibre de charge à tour de rôle sur les interfaces sélectionnées (en supposant qu'elles soient actives) qui ont le même coût d'interface. Si les coûts sont différents, les interfaces les plus coûteuses sont utilisées comme sauvegardes de l'interface la moins coûteuse.

Par exemple, en définissant le même coût pour deux liaisons WAN, vous pouvez équilibrer la charge du trafic sur ces liaisons afin d'améliorer peut-être les performances. Toutefois, si l'un des liens WAN a une bande passante plus élevée que l'autre, vous pouvez définir le coût de la liaison à bande passante supérieure à 1 et pour celui de la bande passante inférieure à 2, de sorte que le lien à bande passante inférieure ne soit utilisé que si le lien à bande passante la plus élevée est en panne.

Étape 7 Vous pouvez définir le type de surveillance pour l'homologue de l'interface afin de collecter les mesures flexibles :

```
policy-route path-monitoring {IPv4 | IPv6 | auto | auto4 | auto6}
```

Où,

- **auto** : envoie des sondes ICMP à la passerelle IPv4 par défaut de l'interface, si elle existe (identique à Auto IPv4). Sinon, envoie à la passerelle par défaut IPv6 de l'interface (identique à Auto IPv6).
- **ipv4** : envoie les sondes ICMP à l'adresse IPv4 homologue spécifiée (IP du saut suivant) pour la surveillance.
- **ipv6** : envoie les sondes ICMP à l'adresse IPv4 homologue spécifiée (IP du saut suivant) pour la surveillance.
- **auto4** : envoie des sondes ICMP à la passerelle IPv4 par défaut de l'interface.

- **auto6** : envoi des sondes ICMP à la passerelle par défaut IPv6 de l'interface.

Exemple :

```
ciscoasa(config-if)# policy-route ?

interface mode commands/options:
  cost          set interface cost
  path-monitoring Keyword for path monitoring
  route-map      Keyword for route-map
ciscoasa(config-if)# policy-route path-monitoring ?
interface mode commands/options:
  A.B.C.D      peer-ipv4
  X:X:X:X::X   peer-ipv6
  auto         Use remote peer IPv4/6 based on config
  auto4        Use only IPv4 address based on config
  auto6        Use only IPv6 address based on config
```

```
ciscoasa(config-if)# policy-route path-monitoring auto
```

Pour effacer les paramètres de surveillance de chemin d'accès sur l'interface, utilisez la commande **clear path-monitoring** :

Exemple :

```
clear path-monitoring outside1
```

Étape 8

Configurez le routage basé sur les politiques pour le trafic traversant la boîte :

```
policy-route route-map route_map_name
```

Exemple :

```
ciscoasa(config-if)# policy-route route-map testmap
```

Pour supprimer une carte de routage basée sur les politiques existante, saisissez simplement la forme **no** de cette commande.

Exemple :

```
ciscoasa(config-if)# no policy-route route-map testmap
```

Exemples de routage basé sur des politiques

Les sections suivantes présentent des exemples de configuration de carte de routage, de routage basé sur des politiques, ainsi qu'un exemple spécifique de PBR en action.

Exemples de configuration de carte de routage

Dans l'exemple suivant, comme aucune action ni aucune séquence n'est spécifiée, une action implicite d'autorisation et un numéro de séquence de 10 sont supposés :

```
ciscoasa(config)# route-map testmap
```

Dans l'exemple suivant, comme aucun critère de correspondance n'est spécifié, une correspondance implicite « any » est présumée :

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# set ip next-hop 1.1.1.10
```

Dans cet exemple, tout le trafic correspondant à <acl> sera acheminé selon la politique et transféré par l'intermédiaire d'une interface externe.

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# match ip address <acl>
ciscoasa(config-route-map)# set interface outside
```

Dans cet exemple, aucune interface ni action de saut suivant n'étant configurée, tout le trafic correspondant à <acl> verra ses champs df bit et dscp modifiés conformément à la configuration et sera transféré selon le routage normal.

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# match ip address <acl>
set ip df 1
set ip precedence afll
```

Dans l'exemple suivant, tout le trafic correspondant à <acl_1> est transféré à l'aide du saut suivant 1.1.1.10, tout le trafic correspondant <acl_2> est transféré à l'aide du saut suivant 2.1.1.10 et le reste du trafic est abandonné. L'absence de critères de « correspondance » implique une correspondance implicite « any ».

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# match ip address <acl_1>
ciscoasa(config-route-map)# set ip next-hop 1.1.1.10

ciscoasa(config)# route-map testmap permit 20
ciscoasa(config-route-map)# match ip address <acl_2>

ciscoasa(config-route-map)# set ip next-hop 2.1.1.10
ciscoasa(config)# route-map testmap permit 30
ciscoasa(config-route-map)# set interface Null0
```

Dans l'exemple suivant, l'évaluation de la carte de routage sera telle que (i) une action de carte de routage permit et une action acl permit appliqueront les actions définies (ii) une action de carte de routage deny et une action acl permit passeront à la recherche de route normale (iii) une action de carte de routage permit/deny et une action acl deny continueront avec l'entrée suivante de la carte de routage. Si aucune entrée de carte de routage suivante n'est disponible, nous reviendrons à la recherche de route normale.

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# match ip address permit_acl_1 deny_acl_2
ciscoasa(config-route-map)# set ip next-hop 1.1.1.10

ciscoasa(config)# route-map testmap deny 20
ciscoasa(config-route-map)# match ip address permit_acl_3 deny_acl_4
ciscoasa(config-route-map)# set ip next-hop 2.1.1.10

ciscoasa(config)# route-map testmap permit 30
ciscoasa(config-route-map)# match ip address deny_acl_5
```

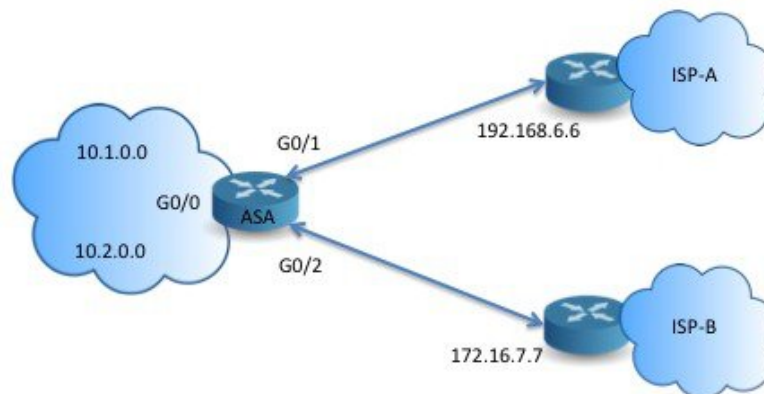
```
ciscoasa(config-route-map)# set interface outside
```

Dans l'exemple suivant, lorsque plusieurs actions définies sont configurées, elles sont évaluées dans l'ordre mentionné ci-dessus. Ce n'est que lorsque toutes les options d'une action définie sont évaluées et ne peuvent être appliquées que les actions définies suivantes seront prises en compte. Cet ordre garantira que le prochain saut le plus disponible et le moins distant sera essayé en premier, suivi du prochain saut le plus disponible et le moins distant, et ainsi de suite.

```
ciscoasa(config)# route-map testmap permit 10
ciscoasa(config-route-map)# match ip address acl_1
ciscoasa(config-route-map)# set ip next-hop verify-availability 1.1.1.10 1 track 1
ciscoasa(config-route-map)# set ip next-hop verify-availability 1.1.1.11 2 track 2
ciscoasa(config-route-map)# set ip next-hop verify-availability 1.1.1.12 3 track 3
ciscoasa(config-route-map)# set ip next-hop 2.1.1.10 2.1.1.11 2.1.1.12
ciscoasa(config-route-map)# set ip next-hop recursive 3.1.1.10
ciscoasa(config-route-map)# set interface outside-1 outside-2
ciscoasa(config-route-map)# set ip default next-hop 4.1.1.10 4.1.1.11
ciscoasa(config-route-map)# set default interface Null0
```

Exemple de configuration pour PBR

Cette section décrit l'ensemble complet de configuration requis pour configurer PBR pour le scénario suivant :



Tout d'abord, nous devons configurer les interfaces.

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# no shutdown
ciscoasa(config-if)# nameif inside
ciscoasa(config-if)# ip address 10.1.1.1 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# no shutdown
ciscoasa(config-if)# nameif outside-1
ciscoasa(config-if)# ip address 192.168.6.5 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/2
ciscoasa(config-if)# no shutdown
ciscoasa(config-if)# nameif outside-2
ciscoasa(config-if)# ip address 172.16.7.6 255.255.255.0
```

Ensuite, nous devons configurer une liste d'accès pour la correspondance du trafic.

```
ciscoasa(config)# access-list acl-1 permit ip 10.1.0.0 255.255.0.0
ciscoasa(config)# access-list acl-2 permit ip 10.2.0.0 255.255.0.0
```

Nous devons configurer une carte de routage en spécifiant la liste d'accès ci-dessus comme critères de correspondance et les actions définies requises.

```
ciscoasa(config)# route-map equal-access permit 10
ciscoasa(config-route-map)# match ip address acl-1
ciscoasa(config-route-map)# set ip next-hop 192.168.6.6

ciscoasa(config)# route-map equal-access permit 20
ciscoasa(config-route-map)# match ip address acl-2
ciscoasa(config-route-map)# set ip next-hop 172.16.7.7

ciscoasa(config)# route-map equal-access permit 30
ciscoasa(config-route-map)# set ip interface Null0
```

Maintenant, cette carte de routage doit être associée à une interface.

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# policy-route route-map equal-access
```

Pour afficher la configuration du routage des politiques.

```
ciscoasa(config)# show policy-route
Interface          Route map
GigabitEthernet0/0 equal-access
```

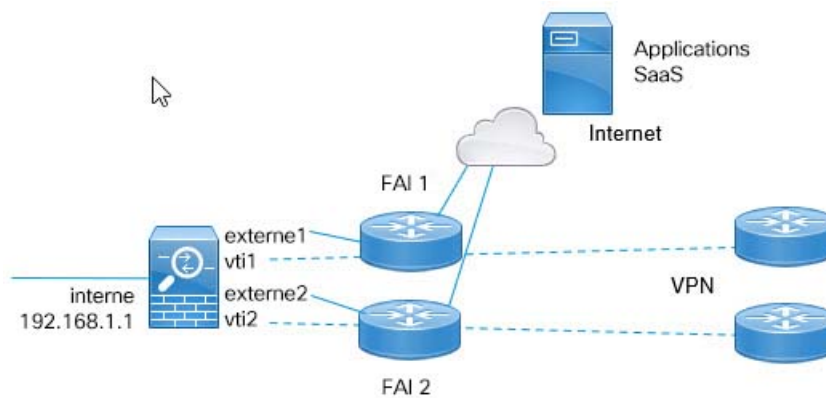
Accès Internet direct à l'aide du réseau WAN défini par logiciel

Un réseau de succursale typique utilise un VPN de site à site pour connecter la succursale à un concentrateur d'entreprise. Tout le trafic non local est ensuite acheminé vers le réseau de l'entreprise, où il est acheminé vers les services internes ou vers Internet, selon le cas.

Cette configuration crée un goulot d'étranglement au niveau du concentrateur de l'entreprise. Si du trafic de succursale est destiné à des services Internet, comme Google Search ou Gmail, il n'est pas nécessaire d'accéder d'abord au réseau d'entreprise avant d'accéder à Internet.

En utilisant le routage basé sur des politiques, vous pouvez plutôt configurer un accès Internet direct à partir de la succursale pour le trafic qui n'a pas besoin des services du réseau d'entreprise. Ainsi, le trafic vers Internet n'est pas envoyé au concentrateur de l'entreprise, et celui-ci n'a plus qu'à gérer le trafic destiné aux services internes du réseau d'entreprise. Cette configuration devrait améliorer les performances globales et le débit du réseau.

L'exemple suivant montre comment configurer l'accès Internet direct pour la configuration suivante, où deux interfaces externes se connectent à différents fournisseurs de services Internet et des interfaces de tunnel virtuel (VTI) hébergent des connexions VPN de site à site au réseau d'entreprise. L'exemple montre comment diriger le trafic destiné à certaines applications SaaS vers Internet et ainsi contourner le réseau d'entreprise.



Avant de commencer

Cet exemple suppose que vous avez déjà défini un VPN de site à site à l'aide d'interfaces de tunnel virtuel (VTI) définies sur les interfaces externes (orientées WAN) pour connecter la succursale au concentrateur de l'entreprise et que celui-ci fonctionne correctement. Le trafic acheminé vers les interfaces VTI est ainsi redirigé vers le réseau d'entreprise, tandis que le trafic acheminé directement vers les interfaces externes est acheminé vers Internet.

Cela suppose également que vous ayez configuré des serveurs DNS et activé la résolution DNS sur les interfaces de l'appareil. Utilisez la commande **show dns trusted-source detail** pour voir quels serveurs seront surveillés. Si vous souhaitez limiter les serveurs utilisés, utilisez la commande **no dns trusted-source** pour désactiver la surveillance de trafic sur les serveurs sélectionnés.

Procédure

Étape 1

Configurez les objets et les groupes de services réseau pour définir le trafic souhaité.

L'exemple suivant crée des objets pour définir Office365 et WebEx, puis crée un groupe d'objets SaaS_Applications pour les contenir. Vous devez créer un groupe d'objets, vous ne pouvez pas utiliser directement les objets dans une entrée de contrôle d'accès.

```
object network-service office365
  domain outlook.office365.com tcp eq 443
  domain onlineapps.live.com tcp eq 443
  domain skype.live.com tcp eq 443

object network-service webex
  domain webex.com tcp eq 443

object-group network-service SaaS_Applications
  network-service-member office365
  network-service-member webex
```

Étape 2

Créez une ACL étendue pour correspondre au trafic souhaité.

L'exemple suivant fait correspondre le trafic du réseau interne au groupe d'objets SaaS Applications.

```
access-list DIA_traffic extended permit ip 192.168.1.0 255.255.255.0
object-group-network-service SaaS_Applications
```

Étape 3 (Facultatif) Configurez le coût sur les interfaces de sortie.

En supposant que les interfaces output1 et output2 sont déjà configurées et fonctionnent, ajoutez simplement la commande `policy-route cost`. Cette étape est facultative si vous souhaitez configurer le système pour utiliser le traitement de circuit cyclique (round robin) pour équilibrer la charge sur les deux liaisons WAN de sortie. Cependant, vous devez définir les coûts si vous souhaitez créer une configuration active/de secours, dans laquelle une liaison est utilisée sauf si elle est hors service.

Voici un exemple de configuration active/active à coût égal.

```
interface G0/0
  nameif outside1
  policy-route cost 1

interface G0/1
  nameif outside2
  policy-route cost 1
```

Voici un exemple dans lequel output1 est la liaison préférée et output2 n'est utilisée que si output1 est hors service.

```
interface G0/0
  nameif outside1
  policy-route cost 1

interface G0/1
  nameif outside2
  policy-route cost 2
```

Étape 4 Créez une carte de routage pour faire correspondre l'ACL étendue et diriger le trafic en conséquence.

L'exemple suivant utilise la liste ACL pour faire correspondre le trafic, puis utilise le coût d'interface adaptatif pour diriger le trafic vers l'interface de sortie.

```
route-map mymap 10
  match ip address DIA_traffic
  set adaptive-interface cost outside1 outside2
```

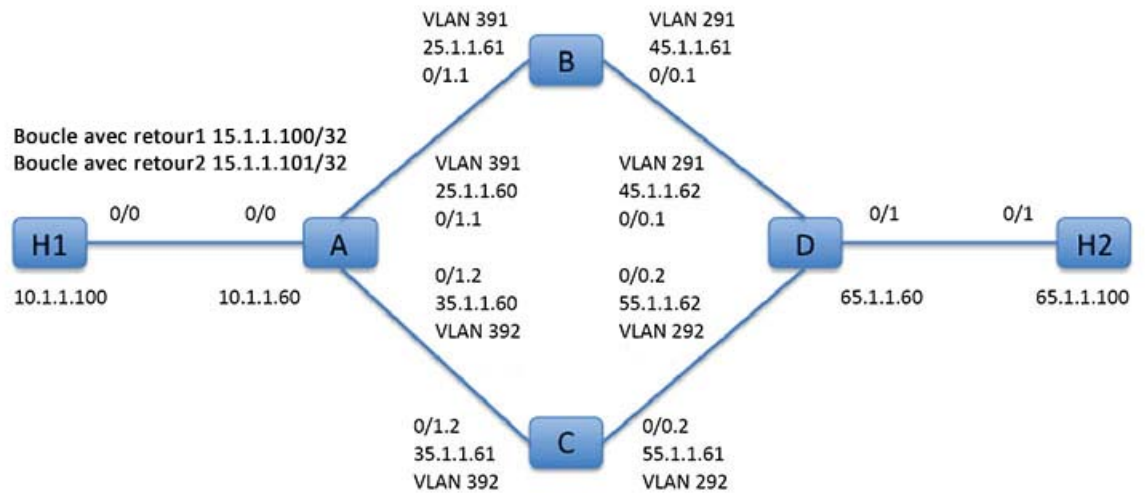
Étape 5 Configurez le routage basé sur des politiques sur l'interface d'entrée pour envoyer le trafic SaaS aux interfaces externes.

L'exemple suivant associe la carte de routage à l'interface interne pour activer le routage basé sur des politiques pour un accès Internet direct.

```
interface G1/0
  nameif inside
  policy-route route-map mymap
```

Routage basé sur des politiques en action

Nous utiliserons cette configuration de test pour configurer le routage basé sur des politiques avec différents critères de correspondance et définir des actions afin de voir comment elles sont évaluées et appliquées.



Tout d'abord, nous allons commencer par la configuration de base pour tous les périphériques impliqués dans l'installation. Ici, A, B, C et D représentent les périphériques ASA, et H1 et H2 représentent les routeurs IOS.

ASA-A :

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# nameif inside
ciscoasa(config-if)# security-level 100
ciscoasa(config-if)# ip address 10.1.1.60 255.255.255.0
ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# no shut

ciscoasa(config)# interface GigabitEthernet0/1.1
ciscoasa(config-if)# vlan 391
ciscoasa(config-if)# nameif outside
ciscoasa(config-if)# security-level 0
ciscoasa(config-if)# ip address 25.1.1.60 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/1.2
ciscoasa(config-if)# vlan 392
ciscoasa(config-if)# nameif dmz
ciscoasa(config-if)# security-level 50
ciscoasa(config-if)# ip address 35.1.1.60 255.255.255.0
```

ASA-B :

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# no shut

ciscoasa(config)# interface GigabitEthernet0/0.1
ciscoasa(config-if)# vlan 291
ciscoasa(config-if)# nameif outside
ciscoasa(config-if)# security-level 0
ciscoasa(config-if)# ip address 45.1.1.61 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# no shut
```

```
ciscoasa(config)# interface GigabitEthernet0/1.1
ciscoasa(config-if)# vlan 391
ciscoasa(config-if)# nameif inside
ciscoasa(config-if)# security-level 100
ciscoasa(config-if)# ip address 25.1.1.61 255.255.255.0
```

ASA-C :

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# no shut

ciscoasa(config)# interface GigabitEthernet0/0.2
ciscoasa(config-if)# vlan 292
ciscoasa(config-if)# nameif outside
ciscoasa(config-if)# security-level 0
ciscoasa(config-if)# ip address 55.1.1.61 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# no shut

ciscoasa(config)# interface GigabitEthernet0/1.2
ciscoasa(config-if)# vlan 392
ciscoasa(config-if)# nameif inside
ciscoasa(config-if)# security-level 0
ciscoasa(config-if)# ip address 35.1.1.61 255.255.255.0
```

ASA-D :

```
ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# no shut

ciscoasa(config) #interface GigabitEthernet0/0.1
ciscoasa(config-if)# vlan 291
ciscoasa(config-if)# nameif inside-1
ciscoasa(config-if)# security-level 100
ciscoasa(config-if)# ip address 45.1.1.62 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/0.2
ciscoasa(config-if)# vlan 292
ciscoasa(config-if)# nameif inside-2
ciscoasa(config-if)# security-level 100
ciscoasa(config-if)# ip address 55.1.1.62 255.255.255.0

ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# nameif outside
ciscoasa(config-if)# security-level 0
ciscoasa(config-if)# ip address 65.1.1.60 255.255.255.0
```

H1 :

```
ciscoasa(config)# interface Loopback1
ciscoasa(config-if)# ip address 15.1.1.100 255.255.255.255

ciscoasa(config-if)# interface Loopback2
ciscoasa(config-if)# ip address 15.1.1.101 255.255.255.255

ciscoasa(config)# ip route 0.0.0.0 0.0.0.0 10.1.1.60
```

H2 :

```
ciscoasa(config)# interface GigabitEthernet0/1
ciscoasa(config-if)# ip address 65.1.1.100 255.255.255.0

ciscoasa(config-if)# ip route 15.1.1.0 255.255.255.0 65.1.1.60
```

Nous configurerons PBR sur ASA-A pour acheminer le trafic provenant de H1.

ASA-A :

```
ciscoasa(config-if)# access-list pbracl_1 extended permit ip host 15.1.1.100 any

ciscoasa(config-if)# route-map testmap permit 10
ciscoasa(config-if)# match ip address pbracl_1
ciscoasa(config-if)# set ip next-hop 25.1.1.61

ciscoasa(config)# interface GigabitEthernet0/0
ciscoasa(config-if)# policy-route route-map testmap

ciscoasa(config-if)# debug policy-route
```

H1 : envoi d'un message Ping 65.1.1.100 répétition 1 boucle avec retour à la source 1

```
pbr: policy based route lookup called for 15.1.1.100/44397 to 65.1.1.100/0 proto 1 sub_proto
 8 received on interface inside
pbr: First matching rule from ACL(2)
pbr: route map testmap, sequence 10, permit; proceed with policy routing
pbr: evaluating next-hop 25.1.1.61
pbr: policy based routing applied; egress_ifc = outside : next_hop = 25.1.1.61
```

Le paquet est transféré comme prévu en utilisant l'adresse de saut suivant dans la carte de routage.

Lorsqu'un prochain saut est configuré, nous effectuons une recherche dans la table de routage d'entrée pour identifier une route connectée vers le prochain saut configuré et utiliser l'interface correspondante. La table de routage d'entrée pour cet exemple est affichée ici (avec l'entrée de route correspondante mise en surbrillance).

```
in  255.255.255.255 255.255.255.255 identity
in  10.1.1.60      255.255.255.255 identity
in  25.1.1.60      255.255.255.255 identity
in  35.1.1.60      255.255.255.255 identity
in  10.127.46.17   255.255.255.255 identity
in  10.1.1.0       255.255.255.0   inside
in  25.1.1.0       255.255.255.0   outside
in  35.1.1.0       255.255.255.0   dmz
```

Ensuite, configurons ASA-A pour acheminer les paquets H1 de boucle avec retour à la source 2 hors de l'interface DMZ d'ASA-A.

```
ciscoasa(config)# access-list pbracl_2 extended permit ip host 15.1.1.101 any

ciscoasa(config)# route-map testmap permit 20
ciscoasa(config-route-map)# match ip address pbracl_2
ciscoasa(config-route-map)# set ip next-hop 35.1.1.61

ciscoasa(config)# show run route-map
!
```

```

route-map testmap permit 10
  match ip address pbracl_1
  set ip next-hop 25.1.1.61
!
route-map testmap permit 20
  match ip address pbracl_2
  set ip next-hop 35.1.1.61
!

```

H1 : envoi d'un message Ping 65.1.1.100 répétition 1 boucle avec retour à la source 2

Les débogages sont affichés ici :

```

pbr: policy based route lookup called for 15.1.1.101/1234 to 65.1.1.100/1234 proto 6 sub_proto
0 received on interface inside
pbr: First matching rule from ACL(3)
pbr: route map testmap, sequence 20, permit; proceed with policy routing
pbr: evaluating next-hop 35.1.1.61
pbr: policy based routing applied; egress_ifc = dmz : next_hop = 35.1.1.61

```

et l'entrée de route choisie dans la table de routage d'entrée est affichée ici :

```

in  255.255.255.255 255.255.255.255 identity
in  10.1.1.60      255.255.255.255 identity
in  25.1.1.60      255.255.255.255 identity
in  35.1.1.60      255.255.255.255 identity
in  10.127.46.17   255.255.255.255 identity
in  10.1.1.0       255.255.255.0    inside
in  25.1.1.0       255.255.255.0    outside
in  35.1.1.0       255.255.255.0    dmz

```

Historique du routage basé sur des politiques

Tableau 1 : Historique des cartes de routage

Nom de la caractéristique	Versions de plateforme	Renseignements sur les fonctionnalités
Indicateurs de supervision des chemins d'accès dans PBR.	9.18(1)	PBR utilise les mesures pour déterminer le meilleur chemin (interface de sortie) pour transférer le trafic. La supervision des chemins informe périodiquement PBR de l'interface surveillée dont la métrique a été modifiée. PBR récupère les dernières valeurs de métrique pour les interfaces surveillées à partir de la base de données de surveillance des chemins et met à jour le chemin d'accès des données. Commandes nouvelles/modifiées : clear path-monitoring, policy-route, show path-monitoring

Nom de la caractéristique	Versions de plateforme	Renseignements sur les fonctionnalités
Routage basé sur les politiques	9.4(1)	Le routage basé sur les politiques (PBR) est un mécanisme par lequel le trafic est acheminé par des chemins spécifiques avec une QoS spécifiée à l'aide d'ACL. Les ACL permettent de classer le trafic en fonction du contenu des en-têtes de couche 3 et de couche 4 du paquet. Cette solution permet aux administrateurs de fournir une QoS à un trafic différencié, de répartir le trafic interactif et par lots entre des chemins permanents à faible bande passante et à faible coût et des chemins commutés à bande passante élevée et à coût élevé, et permet aux fournisseurs de services Internet et à d'autres organisations d'acheminer le trafic provenant de divers ensembles d'utilisateurs par le biais de connexions Internet bien définies. Nous avons introduit les commandes suivantes : set ip next-hop verify-availability, set ip next-hop, set ip next-hop recursive, set interface, set ip default next-hop, set default interface, set ip df, set ip dscp, policy-route route-map, show policy-route, debug policy-route
Prise en charge d'IPv6 pour le routage basé sur les politiques	9.5(1)	Les adresses IPv6 sont désormais prises en charge pour le routage basé sur les politiques. Nous avons introduit les commandes suivantes : set ipv6 next-hop, set default ipv6-next hop, set ipv6 dscp
Prise en charge de VXLAN pour le routage basé sur les politiques	9.5(1)	Vous pouvez maintenant activer le routage basé sur les politiques sur une interface VNI. Nous n'avons pas modifié de commandes.
Prise en charge du routage basé sur les politiques pour le pare-feu d'identité et Cisco TrustSec	9.5(1)	Vous pouvez configurer le pare-feu d'identité et Cisco TrustSec, puis utiliser les ACL du pare-feu d'identité et Cisco TrustSec dans les cartes de routage basées sur les politiques. Nous n'avons pas modifié de commandes.

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.