

Protection de la gestion du réseau

Par une bonne maîtrise de la gestion du réseau, il est possible de se prémunir de la plupart des problèmes de sécurité réseau. Cela recouvre les services qui gravitent autour de la gestion du réseau, tels les services (routage, supervision, etc.) de résolution de noms de domaines, de synchronisation des horloges des équipements réseau, etc.

Il est primordial de considérer comme critique l'architecture et les systèmes en charge de la gestion et de la supervision du réseau. Une bonne gestion du réseau permet d'apporter un premier niveau de sécurité face aux attaques suivantes :

- Attaques par injection de routes, qui consistent à injecter un nombre important de fausses routes ou de routes dupliquées afin de rendre instable le processus de routage du réseau.
- Attaques permettant de générer une instabilité des routes, qui consistent à injecter des mises à jour importantes, par exemple sur une route légitime, afin d'impacter le processus de routage ou de bloquer certaines routes.
- Attaques par déni de service sur les services de noms de domaines, qui peuvent bloquer l'établissement de sessions IP sur un réseau si le service DNS ne répond plus.
- Attaques sur le protocole de gestion du réseau SNMP (Simple Network Management Protocol), qui peuvent impacter le réseau et ses services.

Le tableau 11.1 récapitule les mesures à implémenter afin de garantir un niveau de sécurité acceptable pour l'administration réseau.

Règles de sécurité pour la gestion de réseau

Les règles de sécurité à considérer pour la gestion de réseau sont les suivantes :

- Les accès logiques d'administration des équipements réseau ne sont possibles que depuis une zone dédiée à la gestion du réseau.
- La zone dédiée à la gestion du réseau fait l'objet d'une politique de sécurité spécifique et implémente un niveau de sécurité maximal.
- Les protocoles mis en œuvre pour la gestion du réseau implémentent au maximum les options de sécurité.
- Tous les services et systèmes associés à l'activité réseau sont partie prenante de la politique de sécurité réseau.

Tableau 11.1 Mesures de protection de l'administration réseau

Domaine	Mesure de sécurité
Protection des équipements	Définir des règles de configuration des équipements par type d'équipement et par fonction
Routage réseau (IS-IS, OSPF, BGP, etc.)	– Définir des règles de configuration du protocole IGP permettant d'assurer un périmètre de sécurité du processus de routage – Définir des règles de configuration du protocole EGP permettant d'assurer un périmètre de sécurité du processus de routage – Mettre en place une supervision des indicateurs relatifs aux tables de routage du réseau – Définir des procédures d'intervention en cas de perturbation du routage du réseau
Supervision réseau (SNMP)	– Dédier des serveurs pour la supervision SNMP – Renforcer au maximum la sécurité des systèmes d'exploitation des serveurs – Localiser les serveurs SNMP dans la zone d'administration – Implémenter au maximum les options de sécurité disponibles (authentification) – Suivre et appliquer tous les patchs de sécurité relatifs aux serveurs et au service SNMP – Migrer vers une administration reposant sur le protocole IPsec
Service de noms de domaines (DNS)	– Dédier des serveurs à la résolution de noms de domaines – Renforcer au maximum la sécurité des systèmes d'exploitation des serveurs DNS – Localiser les serveurs DNS dans la zone d'administration – Implémenter au maximum les options de sécurité disponibles (authentification) – Suivre et appliquer tous les patchs de sécurité relatifs aux serveurs et au service DNS
Service de mise à l'heure (NTP)	– Dédier des serveurs à la mise à jour des horloges des équipements réseau – Renforcer au maximum la sécurité des systèmes d'exploitation des serveurs NTP – Localiser les serveurs NTP dans la zone d'administration – Implémenter au maximum les options de sécurité disponibles (authentification) – Suivre et appliquer tous les patchs de sécurité relatifs aux serveurs et au service NTP
Zone d'administration	– Dédier une zone d'administration pour le réseau – Renforcer la sécurité du périmètre de sécurité par des contrôles de filtrage très stricts – Authentifier tous les accès à la zone d'administration – Générer des traces des accès et des commandes passées à des fins d'investigation de sécurité – Installer des systèmes de détection d'intrusion au sein de la zone d'administration – Dédier un plan d'adressage spécifique

Le routage réseau

D'une manière générale, tous les protocoles de routage ont pour objectif de maintenir les tables de routage du réseau dans un état intègre et cohérent. Pour y parvenir, les protocoles diffusent des informations de routage aux autres systèmes du réseau afin de transmettre les modifications des tables de routage. Ces protocoles réceptionnent en contrepartie les informations de routage d'autres systèmes du réseau afin de mettre à jour les tables de routage.

Le déploiement de réseaux IP de grande taille a rapidement nécessité la mise au point de protocoles de routage dynamique chargés de déterminer le plus efficacement possible la meilleure route pour atteindre une destination donnée. Il a aussi été nécessaire de découper le réseau en différents systèmes autonomes, ou AS (Autonomous System), afin de réduire cette complexité. Les systèmes autonomes du cœur de réseau Internet sont gérés par les opérateurs de télécommunications.

Ces considérations ont donné lieu à une classification des protocoles de routage dynamique en deux grandes familles : les protocoles IGP (Interior Gateway Protocol), qui permettent d'échanger des informations d'accessibilité au sein d'un système autonome, et les protocoles EGP (Exterior Gateway Protocol), qui permettent d'échanger des informations d'accessibilité entre systèmes autonomes.

Voici une liste non exhaustive des algorithmes qui peuvent être mis en œuvre lors du processus de routage (ces algorithmes doivent être le plus simple possible afin d'être efficaces pour le calcul et la propagation des mises à jour des tables de routage) :

- **Algorithmes de routage hiérarchique.** Définissent des groupes logiques de nœuds, appelés domaines, systèmes autonomes ou zones. Certains routeurs peuvent communiquer avec les routeurs d'autres domaines, alors que d'autres routeurs ne peuvent communiquer qu'à l'intérieur de leur propre domaine, simplifiant ainsi les algorithmes en fonction des exigences de routage des routeurs appartenant à la hiérarchie.
- **Algorithmes de routage intradomaine.** Ne fonctionnent que dans les limites d'un domaine.
- **Algorithmes de routage interdomaine.** Fonctionnent tant au sein d'un domaine qu'entre divers domaines.
- **Algorithmes de routage d'état des liens.** Testent régulièrement l'état des liens avec leurs voisins et diffusent périodiquement ces états à tous les autres routeurs du domaine. L'algorithme du plus court chemin est généralement fondé sur l'algorithme de Dijkstra, qui calcule le plus court chemin vers chaque destination. Les avantages de tels algorithmes sont d'offrir une convergence rapide sans boucle et à chemins multiples. De plus, chaque passerelle calcule ses propres routes indépendamment des autres. Les métriques sont généralement précises et couvrent plusieurs besoins. En revanche, ces algorithmes sont souvent plus complexes à mettre en œuvre et consommateurs de ressources.
- **Algorithmes de routage à vecteur de distance.** Diffusent régulièrement aux voisins l'état des routes. En se fondant sur les routes reçues, chaque voisin met à jour sa propre

table en fonction de l'adresse du réseau destination, de celle du routeur permettant d'atteindre le réseau destination et du nombre de sauts nécessaire pour l'atteindre. Le calcul de routes distribuées s'appuie le plus souvent sur l'algorithme de Bellman-Ford. Les avantages d'un tel algorithme sont une forte interopérabilité entre systèmes réseau et de faibles impacts sur les ressources système. La convergence des tables de routage se montre en revanche lente lorsque les réseaux deviennent importants, la taille des informations de routage étant proportionnelle au nombre de réseau. De plus, des phénomènes de bouclage peuvent intervenir.

Suivant l'algorithme utilisé, plusieurs paramètres peuvent intervenir pour une décision de routage. Les critères de routage s'appuient généralement sur les éléments suivants :

- **Longueur du trajet.** Définit un critère de décision à partir du nombre de liens qu'un paquet doit traverser pour se rendre du point d'origine au point de destination.
- **Fiabilité.** Définit un critère de décision fondé sur la fiabilité de chaque lien du réseau.
- **Délai de transmission.** Définit un critère de décision fondé sur le temps requis afin d'acheminer un paquet du point d'origine au point de destination.
- **Largeur de bande.** Définit un critère de décision fondé sur la capacité de transmission d'un lien.
- **Charge.** Définit un critère de décision fondé sur les ressources d'un routeur (nombre de paquets traités par seconde, ressource mémoire, etc.).
- **Coût de la communication.** Définit un critère de décision fondé sur un coût appliqué à un lien.

Comme expliqué précédemment, un réseau de routage est découpé en systèmes autonomes (AS), ou zones de responsabilité. Dans un système autonome, le protocole de routage utilisé est de type IGP. Pour les échanges de routage entre systèmes autonomes différents, le protocole de routage utilisé est de type EGP. Pour le domaine du multicast, on distingue les protocoles d'accès, les protocoles de routage intradomaine et les protocoles de routage interdomaine.

Comme l'illustre la figure 11.1 (exemples de protocoles), les protocoles de routage IP interdomaine se situent au-dessus de la couche transport du modèle OSI afin d'assurer une fiabilité dans les sessions pour la mise à jour des routes par l'utilisation du protocole TCP orienté session.

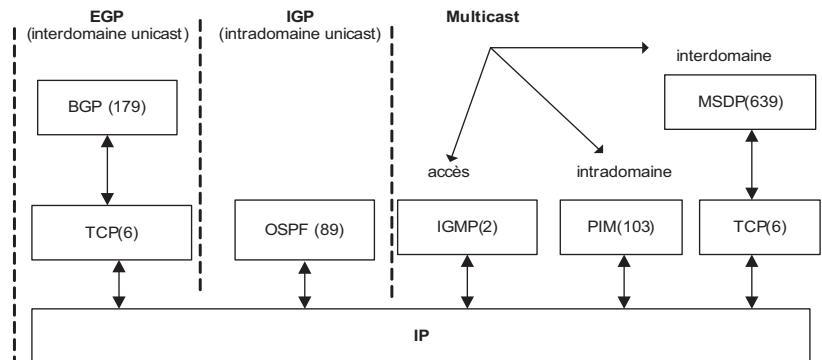
Le routage réseau est donc une pièce maîtresse dans la gestion du réseau. De ce fait, il peut entraîner une faiblesse capitale si des attaques parviennent à perturber directement le processus de routage. Un réseau sans routage perd une fonction fondamentale de sa sécurité, qui est sa disponibilité.

Les protocoles de routage IGP

Les protocoles IGP sont conçus pour gérer le routage interne d'un réseau avec des objectifs de forte convergence des nouvelles routes injectées dans les tables de routage. Les

Figure 11.1

Représentation en couches
des protocoles de routage



décisions de routage s'appuient sur une unique métrique afin de favoriser la fonction de convergence. Le nombre d'entrée dans les tables de routage doit aussi être limité afin de renforcer la fonction de convergence.

Le routage IGP repose généralement sur l'algorithme de Dijkstra. Il s'agit d'un algorithme permettant de trouver, à partir d'un sommet origine unique, le plus court chemin dans un graphe $G = (S, A)$ pondéré, où les arêtes ont des coûts positifs ou nuls. Il s'agit donc d'un algorithme à fixation d'étiquettes (*label setting algorithm*) traitant définitivement un sommet et son étiquette (ou distance) à chaque itération. Il exploite en outre la propriété que les sous-chemins de plus courts chemins sont de plus courts chemins. Le temps processeur nécessaire pour le calcul des tables de routage par un équipement réseau peut être non négligeable.

Par exemple, avec un graphe dense, l'algorithme de Dijkstra a une complexité en temps de l'ordre de $O(S^2)$. Si l'on modifie 10 préfixes par seconde dans un graphe comportant 900 sommets, le temps nécessaire pour mettre à jour les tables de routage nécessite de l'ordre de plusieurs millions d'opérations par seconde dans le pire des cas.

Si un équipement réseau consomme trop de ressources pour le calcul des tables de routage, il impacte le routage proprement dit et par conséquent le trafic réseau.

Les protocoles parmi les plus utilisés de nos jours sont les suivants :

- **OSPF (Open Shortest Path First).** Protocole de routage à état des liens fondé sur le calcul des chemins les plus courts et sur une architecture hiérarchique constituée de zones OSPF.
- **IS-IS (Intermediate System to Intermediate Systems).** Protocole de routage à état des liens fondé sur le calcul des chemins les plus courts et sur une architecture hiérarchique constituée de domaines IS-IS.

D'autres protocoles, tels RIP (Routing Information Protocol) ou IGRP (Interior Gateway Routing Protocol), sont des protocoles de routage à vecteur de distance.

Le protocole de routage IS-IS

IS-IS est un protocole interne de routage. Issu de l'ensemble des protocoles OSI, il fournit un support pour la mise à jour d'informations de routage entre de multiples protocoles. Il s'agit d'un protocole par état des liaisons de type SPF (Shortest Path First), ou chemin le plus court, dont la dernière version est conforme à la norme ISO 10589.

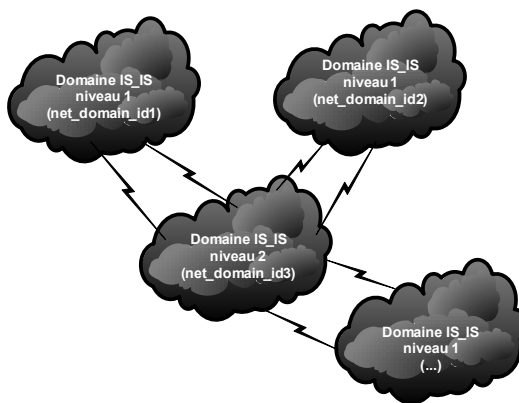
Le routage IS-IS utilise deux niveaux hiérarchiques de routage. La topologie de routage IS-IS est donc partitionnée en domaines de routage de niveaux 1 ou 2. Les routeurs de niveau 1 connaissent la topologie dans leur domaine, incluant tous les routeurs de ce domaine. Cependant, ces routeurs de niveau 1 ne connaissent ni l'identité des routeurs ni les destinations à l'extérieur de leur domaine. Ils routent tout le trafic vers les routeurs interconnectés au niveau 2 dans leur domaine.

Les routeurs de niveau 2 connaissent la topologie réseau du niveau 2 et savent quelles adresses sont atteignables pour chaque routeur. Les routeurs de niveau 2 n'ont pas besoin de connaître la topologie à l'intérieur d'un domaine de niveau 1. Seuls les routeurs de niveau 2 peuvent échanger les paquets de données ou les informations de routage direct avec les routeurs externes situés en dehors de leur domaine de routage.

Le domaine de niveau 2 agit comme domaine d'échange entre les domaines de niveau 1. La figure 11.2 illustre une topologie type d'interconnexion entre les domaines IS-IS de niveaux 1 et 2.

Figure 11.2

Topologie des interconnexions des domaines IS-IS

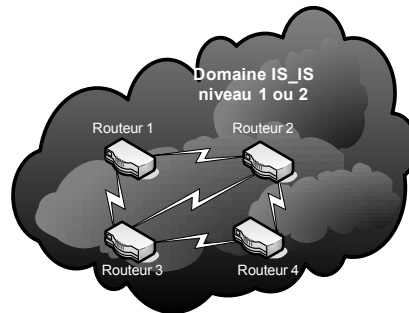


Pour des raisons de disponibilité des connexions entre les domaines, le nombre de sessions d'interconnexion entre les domaines 1 et 2 doit être supérieur au minimum à deux sessions IS-IS.

Le réseau de routage formé par les routeurs des domaines 1 ou 2 doit définir un graphe connexe (il existe un chemin entre tous les nœuds du graphe). La figure 11.3 illustre une topologie type d'interconnexion des équipements réseau dans un domaine IS-IS de niveau 1 ou 2.

Figure 11.3

Topologie des interconnexions des équipements réseau dans un domaine



Le réseau de routage formé par les routeurs au sein d'un domaine doit être un graphe connexe (il existe un chemin entre tous les nœuds du graphe) et sans point d'articulation.

Règles de sécurité pour l'architecture de routage IS-IS

Les règles de sécurité à considérer pour l'architecture de routage IS-IS sont les suivantes (pour une configuration de routeur Cisco) :

- L'architecture de routage IS-IS et le découpage en domaines IS-IS sont clairement explicités et documentés. La topologie de routage est décrite dans les documents de l'ingénierie.
- Les échanges des tables de routage sont authentifiés lors d'une session IS-IS.

La commande suivante :

```
isis password password {level-isis}
```

permet de définir un mot de passe par session IS-IS.

- Les tables de routage sont limitées aux classes d'adresses IP autorisées.

La commande suivante :

```
distance weight {ip-address {ip-address mask}} [ip access-list]
access-list access-list-number [dynamic list-name [timeout value]] {deny
| permit} protocol source source-wildcard destination destination-
wildcard [precedence precedence] [tos tos] [log| log-input]
```

permet de définir une ACL fondée sur les adresses IP à filtrer.

Les protocoles de routage EGP

Nous ne décrivons dans cette section que le protocole BGP (Border Gateway Protocol), qui s'est imposé comme le protocole EGP du réseau Internet.

BGP s'appuie sur la couche TCP (port 179) pour établir une connexion TCP entre deux routeurs et échanger d'une manière dynamique les annonces de routes.

Le routage BGP repose généralement sur l'algorithme de Bellman-Ford distribué. Il s'agit d'un algorithme réparti et autostabilisant, dans lequel chaque sommet x maintient une table des distances donnant le voisin z à utiliser pour joindre la destination y . On le note $D^x(y,z)$.

L'algorithme se fonde sur le calcul de l'invariant suivant pour chaque sommet et pour chacune de ses destinations :

$$D^x(y,z) = c(x,y) + \min_w D^c(y,w)$$

où w désigne les voisins de z .

L'algorithme exploite la propriété que les sous-chemins de plus courts chemins sont des plus courts chemins. Lorsqu'un nœud calcule un nouveau coût minimal pour une destination lors d'une mise à jour de routage provenant de ses voisins ou lorsque le coût d'une de ses adjacences a changé, il informe ses voisins de cette nouvelle valeur. Il s'agit donc d'un algorithme à correction d'étiquettes (*label correcting algorithm*) pouvant affiner à chaque itération l'étiquette (ou distance) de chaque sommet.

Si l'on considère qu'un équipement réseau a de l'ordre de 20 voisins en moyenne et que chaque voisin envoie une mise à jour de routage modifiant 5 000 préfixes par seconde, il faut de l'ordre de plusieurs millions d'opérations par seconde dans le pire des cas pour mettre à jour les tables de routage. De plus, le nombre de messages de mises à jour de routage envoyé par cet équipement réseau peut être important et générer une attaque par déni de service sur le réseau considéré (la taille d'un message BGP peut aller jusqu'à 4 096 octets).

Si un équipement réseau consomme trop de ressources pour le calcul des tables de routage, il impacte le routage proprement dit et par conséquent le trafic réseau.

Lorsqu'un routeur BGP reçoit des mises à jour en provenance de plusieurs systèmes autonomes décrivant différents chemins vers une même destination, il choisit le meilleur itinéraire pour l'atteindre et le propage vers ses voisins.

La décision de routage est fondée sur plusieurs attributs, notamment les suivants :

- **AS-path.** Liste les numéros de tous les AS qu'une mise à jour doit traverser pour atteindre une destination.
- **Origin.** Donne des informations sur l'origine de la route. Ces informations peuvent être IGP (la route annoncée provient du même système autonome que l'annonceur), EGP (la route est apprise et ne provient pas du même système autonome) ou Incomplète (la route est apprise d'une autre manière).
- **Next hop.** Contient l'adresse IP du routeur vers lequel l'information doit être émise pour atteindre le réseau.
- **Weight.** Utilisé dans le processus de sélection de chemin lorsqu'il existe plus d'une route vers une même destination. On définit alors un poids. L'attribut de poids est local au routeur et n'est pas propagé dans les mises à jour de routage.
- **Local preference.** Rôle similaire à l'attribut de poids, si ce n'est que l'attribut de préférence locale fait partie des informations de mise à jour de routage.
- **Multi-exit discriminator.** Indication aux routeurs voisins externes concernant le chemin à privilégier vers un AS lorsque celui-ci possède plusieurs points d'entrée (*via* les différents routeurs externes de l'autre AS).
- **Community.** Utilisé pour grouper des destinations auxquelles des décisions de routage peuvent être appliquées.

Les routes apprises par les sessions eBGP d'un système autonome doivent être propagées au sein du système autonome par le biais de sessions iBGP. Il s'agit de maintenir une vue cohérente de l'ensemble des routes externes au système autonome pour l'ensemble des routeurs.

La spécification initiale de BGP suppose qu'un graphe complet (modèle « complet ») de sessions iBGP soit configuré au sein du système autonome pour distribuer les routes interdomaines.

Par conséquent, il doit y avoir :

$\frac{n \times (n - 1)}{2}$ sessions iBGP au sein d'un système autonome si n est le nombre de routeurs.

La raison à cela est que les sessions iBGP ne redistribuent pas les routes apprises en iBGP afin d'éviter les phénomènes de bouclage. Par exemple, pour un réseau contenant 100 routeurs, il serait nécessaire de configurer de l'ordre de 5 000 sessions iBGP au total dans les configurations des routeurs.

Le modèle réflecteur de routes a été proposé pour réduire le nombre de configurations des sessions iBGP.

Sachant que le sous-graphe associé aux réflecteurs de routes doit être complet, il doit y avoir :

$\frac{n \times (n - 1)}{2}$ sessions iBGP entre les réflecteurs de routes.

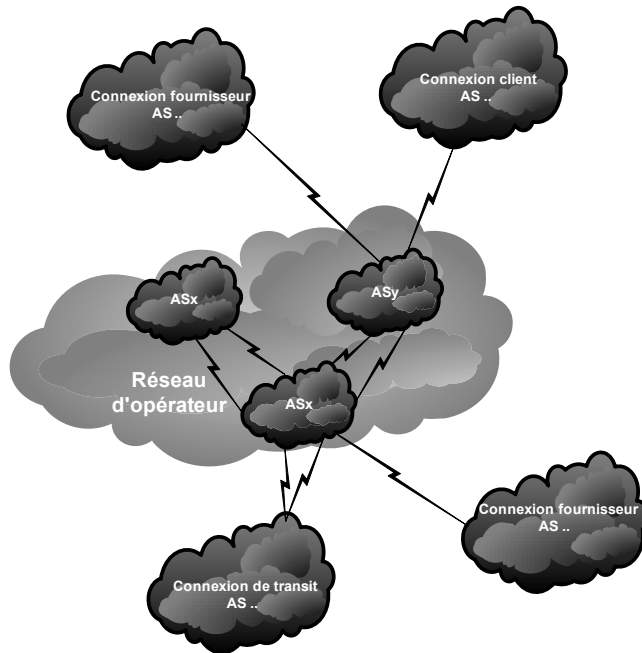
Cependant, le nombre de réflecteurs de routes nécessaires est, par architecture, très inférieur comparé au nombre de routeurs dans le système autonome.

Les niveaux d'architecture illustrés aux figures 11.4 à 11.7 peuvent être définis à l'aide de BGP :

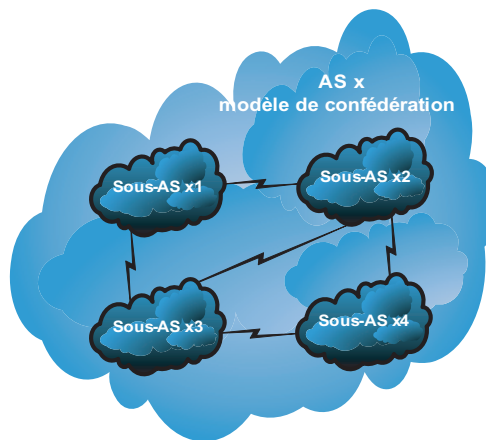
- Le premier niveau d'architecture (*voir figure 11.4*) définit les différentes interconnexions entre systèmes autonomes. Le découpage en différents AS dépend de nombreux paramètres, tels que le nombre d'équipements réseau, la localisation géographique, etc. Pour des raisons de disponibilité et de résilience des connexions entre AS, le nombre de sessions d'interconnexion entre les systèmes autonomes de l'opérateur doit être supérieur au minimum à deux sessions. La figure illustre une topologie type des interconnexions d'un réseau d'un opérateur de télécommunications, avec ses partenaires et ses clients au niveau des AS.
- Le second niveau d'interconnexion (*voir figure 11.5*) concerne le découpage d'un système autonome en sous-systèmes autonomes, ou SubAS. Ce type d'architecture est un modèle de type confédération, qui ne semble plus être utilisé du fait de la difficulté de maintenir la cohérence des configurations. La figure illustre une topologie type des interconnexions des sous-systèmes autonomes au sein d'un système autonome.
- Le dernier niveau d'architecture (*voir figures 11.6 et 11.7*) se situe soit au niveau d'un système autonome, soit dans un sous-système autonome. Il est composé de plusieurs

Figure 11.4

Topologie des interconnexions des systèmes autonomes d'un réseau BGP

**Figure 11.5**

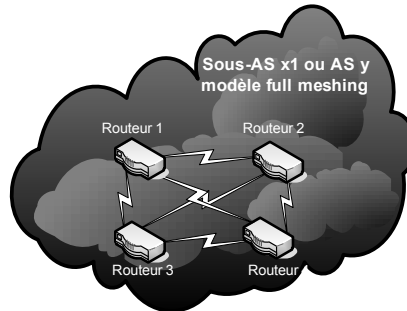
Topologie des interconnexions des sous-systèmes autonomes d'un réseau BGP



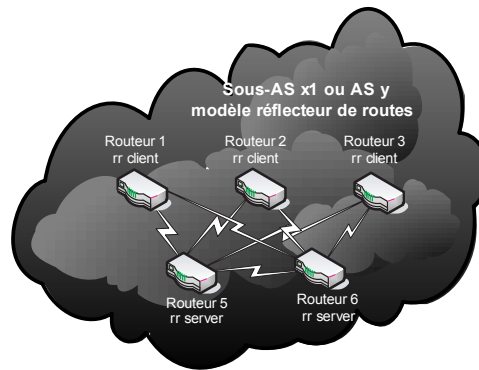
modèles possibles. Dans le modèle full-meshing, tous les équipements ont une session BGP les uns avec les autres. Dans le modèle de type réflecteur de routes, tous les équipements ont une session BGP avec des équipements dédiés au routage, appelés réflecteurs de routes. Dans les deux cas, les graphes doivent être connexes (il existe un chemin entre tous les nœuds du graphe). Une combinaison des deux topologies est possible. Un réflecteur de routes est un routeur BGP qui peut redistribuer sur des sessions iBGP les routes qu'il a apprises d'autres sessions iBGP.

Figure 11.6

Topologie des interconnexions des équipements réseau d'un réseau BGP (full-meshing)

**Figure 11.7**

Topologie des interconnexions des équipements réseau d'un réseau BGP (réflecteur de routes)



Un réflecteur de routes est un routeur BGP qui peut redistribuer sur des sessions iBGP les routes qu'il a apprises d'autres sessions iBGP. Un réflecteur de routes a des voisins clients et des voisins non clients (les voisins non clients sont considérés ici comme des réflecteurs de routes). Un réflecteur de routes reçoit des routes de tous ses voisins iBGP et utilise son processus de décision BGP afin de déterminer les meilleures routes pour joindre chaque destination. Si la meilleure route a été reçue sur une session iBGP avec un voisin client, le réflecteur de route annonce cette route à tous ses voisins iBGP. Si la route a été reçue d'un voisin non client, la route n'est annoncée qu'aux voisins clients.

Mécanismes de sécurité du routage externe

Les échanges de routes avec des AS externes constituent un point névralgique de la sécurité d'un réseau.

Nous détaillons dans cette section les techniques permettant de renforcer la sécurité de tels échanges.

Contrôle par secrets partagés

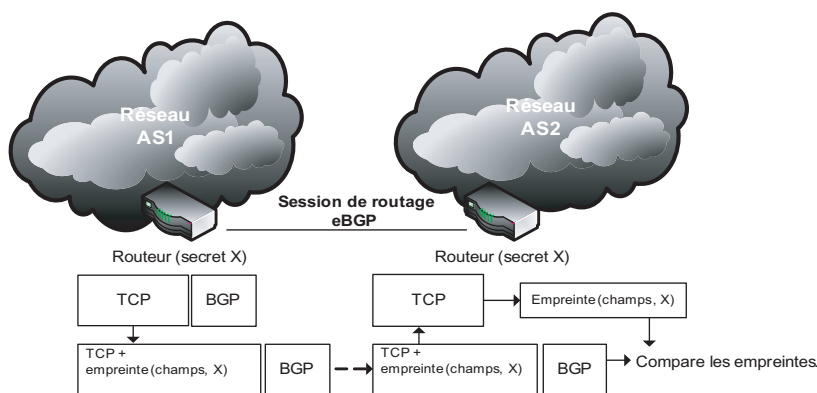
Le contrôle d'une session de routage BGP entre deux routeurs peut être réalisé par l'option d'empreinte MD5 véhiculée dans les paquets TCP. Il s'agit de vérifier en point-à-point les annonces de routes échangées entre deux routeurs à l'aide d'un secret partagé,

ou clé secrète. Ce contrôle doit être mis en œuvre en priorité sur les sessions eBGP, qui présentent le plus de risques pour un AS.

Sachant que les deux routeurs possèdent un secret partagé, une empreinte fondée sur une fonction de hachage (MD5, SHA-x, etc.) est générée pour contrôler les échanges de routes. Quand un routeur émet un paquet IP contenant des données BGP, une empreinte est calculée et insérée dans le paquet TCP, puis vérifiée par l'autre routeur BGP, comme l'illustre la figure 11.8.

Figure 11.8

Contrôle du routage par les secrets partagés



Cette empreinte est calculée à partir de la clé secrète et de champs constants qui n'ont pas été modifiés par le processus d'acheminement du paquet, notamment les suivants :

- adresse IP source ;
- adresse IP destination ;
- en-tête TCP sans les options avec un checksum à 0 ;
- données du segment TCP ;
- secret partagé ou clé secrète (distribué par un canal sécurisé).

L'empreinte est insérée dans le champ Options du paquet TCP, permettant de mettre en œuvre un mécanisme de contrôle d'une session de routage BGP. En revanche, elle ne permet pas d'authentifier le chemin pris par une route ni l'origine de la route.

Les différents secrets partagés permettent aussi de créer des groupes distincts ou périmètres de sécurité entre les sessions iBGP et les diverses sessions eBGP.

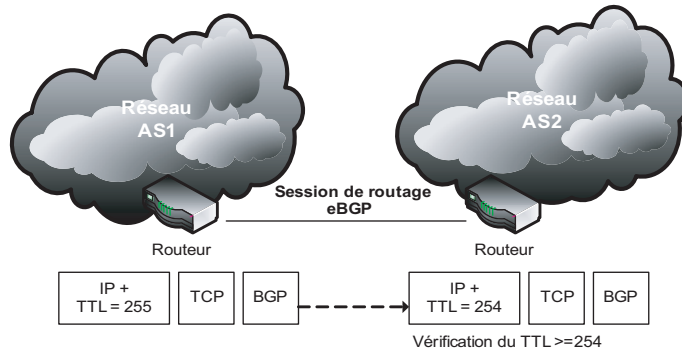
Contrôle par les TTL

Une autre méthode pour contrôler une session de routage BGP consiste à mettre en place un contrôle du TTL (Time To Live) contenu dans les paquets IP échangés par la session de routage BGP. Ce contrôle doit être mis en œuvre en priorité sur les sessions eBGP, qui comportent le plus de risques pour un AS.

Partant du principe que les sessions de routage BGP entre deux routeurs sont généralement directes, les paquets IP contenant des informations de routage BGP émis par un routeur doivent arriver à l'autre routeur avec un $TTL = TTL - 1$ (voir figure 11.9).

Figure 11.9

Contrôle du routage par les TTL



Comme une annonce de routes entre deux routeurs correspond chaque fois à un nouveau paquet IP, le TTL du paquet IP émis est par défaut égal à 255. Si l'autre routeur reçoit des annonces de routes ayant un TTL qui n'est pas égal à 254, il peut en conclure que ce n'est pas le routeur avec lequel il a une session de routage qui a émis cette annonce.

Ce contrôle permet de mettre en œuvre un mécanisme de contrôle d'une session de routage BGP. En revanche, il ne permet pas plus que le précédent d'authentifier le chemin annoncé par une route ni l'origine de la route.

Contrôle des annonces de routes eBGP

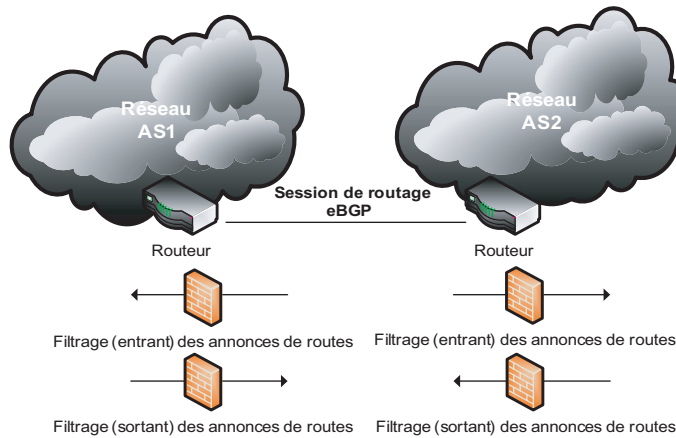
Les annonces de routes peuvent être soumises à une réelle politique de routage définie par l'administrateur d'un système autonome (opérateur de télécommunications). Cette politique peut à la fois s'appliquer aux annonces de routes émises vers un système autonome (routes transmises à l'intérieur d'un AS) et aux annonces de routes qu'émet le système autonome (routes émises à l'extérieur d'un AS), comme l'illustre la figure 11.10.

Cette politique de routage définit des règles de contrôle ou de filtrage fondées notamment sur les éléments suivants :

- Listes de filtrage associées aux valeurs des systèmes autonomes. Par exemple, telle route ne peut être annoncée que par la liste des systèmes autonomes suivants.
- Listes de filtrage associées aux préfixes annoncés ou émis. Par exemple, certains préfixes ne doivent pas être annoncés (RFC 1918).
- Contrôles de l'instabilité des routes. Par exemple, si un préfixe fait l'objet de mises à jour incessantes, il peut être mis en quarantaine afin de protéger le processus de routage BGP.

Figure 11.10

Contrôle du routage par les annonces de routes



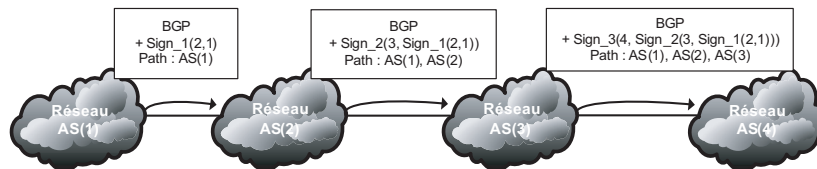
Contrôle de l'authentification des routes

Bien qu'il existe un certain nombre d'éléments de configuration permettant de renforcer la sécurité des sessions BGP, deux problèmes fondamentaux subsistent. Le premier consiste à authentifier l'origine d'une route et le second à authentifier le chemin pris par une route. Quelques initiatives ont vu le jour pour répondre à ces problématiques.

La première initiative, sBGP (secure-BGP), consiste à déployer un système à clé publique dans lequel chaque système autonome possède un certificat électronique. Les sessions de routage BGP s'établissent *via* le protocole IPsec. Lors de l'annonce d'une route, chaque système autonome vérifie le chemin émis et signe à son tour avec sa clé privée le chemin s'il doit l'annoncer à un autre système autonome, comme l'illustre la figure 11.11 (les signatures s'empilent comme les couches d'un oignon).

Figure 11.11

Contrôle du routage par l'authentification sBGP

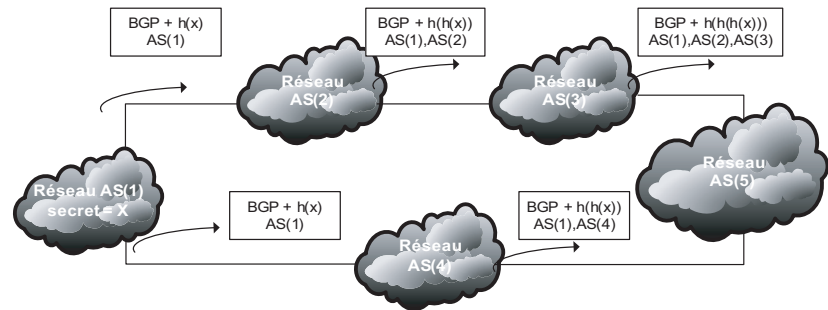


La deuxième initiative exploite le fait que le déploiement d'un système à clé publique ajouté aux impacts cryptographiques sur les processeurs des routeurs limitent une mise en œuvre rapide d'un tel système. « Listen and Whisper » propose notamment une méthode de contrôle des annonces de routes limitant au minimum les impacts sur le temps processeur des routeurs. L'idée consiste à fournir un mécanisme permettant de vérifier la consistance des annonces de routes. Par exemple, à la figure 11.12, l'AS(e) reçoit deux annonces de routes par deux chemins différents.

Lors de l'initialisation de l'annonce d'une route, l'AS(a) génère un secret X et utilise une fonction de hachage pour ajouter une empreinte à ses annonces de routes. Chaque AS

Figure 11.12

Contrôle du routage par
l'authentification Whisper



traversé génère une nouvelle empreinte fondée sur l'empreinte précédente. Si l'AS(e) reçoit deux annonces de routes r et s , de longueurs respectives k et l (représentant le nombre d'AS traversés, $k > l$) et d'empreintes y_r et Y_s , il peut vérifier la consistance de la route en réalisant le calcul $h^{k-l}(y_s) = y_r$.

Si cette solution n'impacte que faiblement les temps processeur des routeurs, elle ne permet pas d'authentifier de manière sûre l'origine d'une route.

La troisième initiative, SoBGP (Secure origin BGP), émanant de Cisco, veut répondre aux mêmes besoins de sécurité que la solution sBGP, mais avec une approche différente, qui nécessite de déployer une nouvelle couche de serveurs pour contrôler les certificats et les chemins associés aux routes.

Enfin, l'initiative IRV (Interdomain Routing Validation) consiste à ne pas modifier le protocole BGP et à proposer une architecture de serveurs spécifique permettant de valider les informations de routage interdomaine hors bande.

En dépit de toutes ces initiatives, aucune de ces solutions n'est actuellement mise en œuvre.

Règles de sécurité pour l'architecture de routage BGP

Les règles de sécurité à considérer pour l'architecture de routage BGP sont les suivantes :

- L'architecture de routage est clairement documentée et justifiée. Cela recouvre le découpage en différents systèmes et sous-systèmes autonomes.
- La topologie de routage est décrite dans les documents de l'ingénierie.
- Les échanges de tables de routage lors d'une session BGP sont authentifiés.

La commande suivante :

```
neighbor {ip-address | peer-group-name} password string
```

permet de définir un mot de passe pour une session BGP.

- Les échanges de tables de routage lors d'une session BGP sont authentifiés.

La commande suivante :

```
neighbor ip-address ttl-security hops hop-count
```

permet de définir un contrôle BGP fondé sur le TTL.

- Un filtrage sur les numéros de systèmes autonomes est défini et mis en place pour les interconnexions avec les autres opérateurs réseau ou fournisseurs de services réseau. Ce filtrage couvre les échanges de routes du réseau vers l'extérieur (filtre out) et de l'extérieur vers le réseau (filtre in).

La commande suivante :

```
neighbor {ip-address | peer-group-name} filter-list access-list-number
{in | out}
```

permet de définir un filtre sur les systèmes autonomes pour une session BGP.

- Un nombre maximal de préfixes est défini afin de s'assurer que le nombre d'entrées dans les tables de routage reste sous contrôle.

La commande suivante :

```
neighbor {ip-address | peer-group-name} maximum-prefix maximum
[threshold] [warning-only]
```

limite le nombre de préfixes IP annoncés.

- Un filtrage sur les classes d'adresses IP non autorisées est défini et mis en place, notamment sur les classes IANA réservées (filtrages out et in).

La commande suivante :

```
neighbor {ip-address | peer-group-name} prefix-list prefix-listname {in
| out}
```

permet de définir un filtre sur les adresses IP pour une session BGP.

- Un filtrage fondé sur l'attribut de communauté du protocole BGP (filtrages out et in) est défini et mis en place.

La commande suivante :

```
neighbor {ip-address | peer-group-name} route-map route-map-name {in |
out}
```

permet d'appliquer une politique de route-map sur une session BGP.

La commande suivante :

```
route-map map-tag [[permit | deny] | [sequence-number]]
```

permet de définir une politique de route-map.

La commande suivante :

```
set community {community-number [additive]} | none
```

permet d'ajouter des attributs BGP.

La commande suivante :

```
set comm-list community-list-number | community-list-name delete
```

permet d'ajouter ou d'enlever des attributs BGP.

La commande suivante :

```
match as-path path-list-number
```

permet de faire correspondre des systèmes autonomes.

La commande suivante :

```
match community standard-list-number|expanded-list-number|community-
list-name [exact-match]
```

permet de faire correspondre une liste d'attributs BGP.

- Un filtrage sur l'instabilité des mises à jour de routes (*dampening*) est défini et mis en place.

La commande suivante :

```
bgp dampening [half-life reuse suppress max-suppress-time] [route-map
map]
```

permet de définir une politique prenant en compte les instabilités des mises à jour de route.

Toutes ces mesures de sécurité protègent le réseau d'éventuelles attaques de routage, sans toutefois apporter de sécurité totale du routage réseau. Elles ne permettent pas d'authentifier le chemin pris par une route ni l'origine de la route. Il est de surcroît possible de détourner du trafic par le routage à des fins de vol d'information.

La principale faiblesse des protocoles de routage actuels vient du fait qu'ils n'intègrent aucune brique de sécurité. Sachant que toute attaque ou perturbation du routage peut impacter directement la disponibilité du réseau et de ses services, il est primordial de considérer les protocoles de routage comme des éléments-clés de la sécurité d'un réseau.

Les protocoles de routage multicast

Dans les réseaux IP, les paquets sont généralement acheminés d'une seule source vers un seul récepteur, de proche en proche, par des routeurs (mode unicast IP). Cependant, pour les applications telles que la diffusion de contenus audio/vidéo nécessitant que les paquets IP soient délivrés à de multiples destinations, l'émission d'une copie de chaque paquet IP à chaque destinataire atteint ses limites lorsque le nombre de récepteurs est important (la bande passante réseau nécessaire augmente, la même donnée étant transportée de multiples fois sur les mêmes liens).

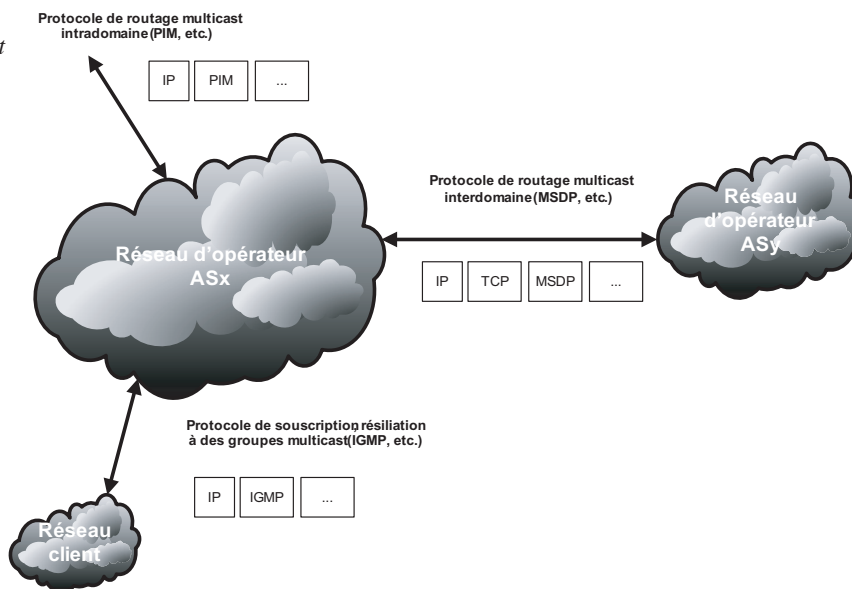
Le multicast répond à cette problématique en fournissant une méthode efficace pour le transport des communications multipoint-à-multipoint. Ce mode de diffusion permet à une source d'émettre une seule copie de son trafic à destination de plusieurs récepteurs. C'est alors au réseau de répliquer de façon optimale le trafic au plus près des récepteurs en créant des arbres de distribution (arbre spécifique, arbre partagé).

Le multicast IP est de plus en plus couramment déployé, à la fois dans Internet et dans les réseaux privés, pour fournir des services de diffusion de contenu multimédia nécessitant de diffuser des données de façon simultanée à un ensemble d'abonnés. Il permet d'économiser de précieuses ressources de bande passante et de capacités réseau et allège la charge des applications de diffusion, qui n'ont plus à émettre autant de copies du programme à diffuser qu'elles ont de destinataires.

Des protocoles très différents doivent être activés au niveau du réseau pour mettre en œuvre un service de diffusion multicast, notamment les suivants (*voir figure 11.13*) :

- protocoles d'accès multicast tels que IGMP (Internet Group Management Protocol), MLD (Multicast Listener Discovery), Proxy IGMP/MLD, snooping IGMP/MLD, GMRP (Generic Attribute Registration Protocol-Multicast Registration Protocol), etc. ;
- protocoles de routage multicast intradomaines tels que PIM-SM (Protocol Independent Multicast Sparse Mode), PIM-DM (Protocol Independent Multicast Dense Mode), DVMRP (Distance Vector Multicast Routing Protocol), MOSPF (Multicast Open Shortest Path Forwarding), etc. ;
- protocoles de routage interdomaines tels que MSDP (Multicast Source Discovery Protocol), BGMP (Border Gateway Multicast Protocol), PIM-SSM (Protocol Independent Multicast-Source-Specific Multicast), etc.

Figure 11.13

Les protocoles multicast

Afin de supporter des communications de groupe, les trois mécanismes distincts suivants doivent être définis et mis en œuvre au niveau de la couche réseau :

- **Adressage.** Il doit y avoir une adresse multicast IP permettant de communiquer avec un groupe de récepteurs plutôt qu'avec un seul récepteur. Cette adresse doit permettre d'identifier un ensemble de destinataires faisant partie d'un groupe spécifique. Un mécanisme doit permettre d'associer cette adresse multicast IP à l'adresse multicast de la couche liaison de données, quand elle existe. Le listing suivant détaille l'attribution d'adresses multicast :

```

...
224.0.1.6  NSS, Name Service Server.
224.0.1.7  AUDIOWEBS - Audio News Multicast.
224.0.1.8  SUN NIS+ Information Service.
224.0.1.9  MTP, Multicast Transport Protocol.
224.0.1.10 IETF-1-LOW-AUDIO. 224.0.1.11 IETF-1-AUDIO.
224.0.1.12 IETF-1-VIDEO.
224.0.1.13 IETF-2-LOW-AUDIO.
224.0.1.14 IETF-2-AUDIO.
224.0.1.15 IETF-2-VIDEO.
224.0.1.16 MUSIC-SERVICE.
224.0.1.17 SEANET-TELEMETRY.
224.0.1.18 SEANET-IMAGE.
224.0.1.19 MLOADD.
224.0.1.20 Any private experiment.
224.0.1.21 DVMRP on MOSPF.
224.0.1.22 SVRLOC.
224.0.1.23 XINGTV.

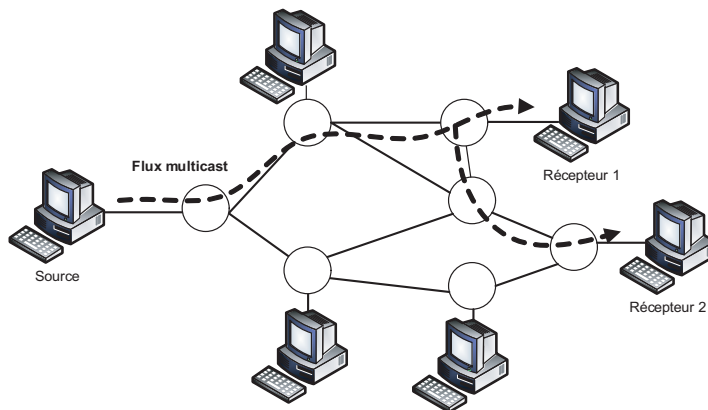
```

224.0.1.24 microsoft-ds.
224.0.1.25 nbc-pro.
224.0.1.26 nbc-pfn.
224.0.1.27 lmsc-calren-1.
...

- **Enregistrement dynamique.** Il doit exister un mécanisme permettant à des terminaux de joindre ou de quitter une communication de groupe. Faute de cela, le réseau ne peut savoir quels sous-réseaux ont besoin de recevoir le trafic d'un groupe en particulier.
- **Routing multicast.** Le réseau doit être capable de calculer et de construire des arbres de distribution multicast des destinataires permettant à des sources d'envoyer des paquets vers tous les récepteurs. Ces arbres de distribution permettent d'assurer que :
 - Le trafic multicast atteint tous les destinataires qui ont joint le groupe multicast.
 - Le trafic multicast n'est pas transmis vers des réseaux dans lesquels il n'y a pas de récepteur (sauf s'il s'agit d'un réseau de transit permettant d'atteindre d'autres destinataires).
 - Une seule copie d'un même paquet est transmise sur un lien réseau donné, même s'il y a de multiples destinataires connectés à ce lien.
 - Une source de trafic n'a pas à répliquer de paquets et envoie une seule copie de chaque paquet de données à destination de multiples récepteurs.

La figure 11.14 illustre la topologie d'un seul groupe multicast et le processus de diffu-

Figure 11.14
*Exemple de diffusion
multicast*



sion du trafic sur un arbre de distribution multicast. La source de trafic envoie une seule copie de ses données multicast, lesquelles sont répliquées par les routeurs multicast de manière à atteindre tous les terminaux membres du groupe (récepteurs 1 et 2) ayant souscrit à ce flux multicast.

Les protocoles de routage multicast sont classifiés en deux grandes familles (dense et éparse) en fonction de la distribution attendue des membres du groupe multicast sur le réseau :

- Dans les topologies réseau de type dense, il est supposé que la répartition des récepteurs d'un groupe donné est dense et homogène sur l'ensemble d'un domaine réseau.
- Dans les topologies réseau de type épars, il est supposé que la répartition des récepteurs d'un groupe donné n'est pas homogène et est donc fortement dispersée sur l'ensemble d'un domaine réseau. Afin de préserver les ressources du réseau, il est important de pouvoir restreindre le trafic multicast et de l'empêcher d'être diffusé vers des régions du réseau où il n'y a pas de membre.

Bien que les protocoles en mode épars soient plus complexes à administrer opérationnellement, ils ont prouvé leur efficacité pour des réseaux de grande taille. PIM-SM (Sparse Mode), le protocole de routage multicast le plus largement répandu aujourd'hui, est un protocole de routage multicast intradomaine en mode épars. Il utilise des arbres partagés ou RPT (Rendez-vous Point Tree) par groupe, qui ont pour racine un routeur particulier, appelé Rendez-vous Point (RP). Le rôle du routeur RP est de servir de point de rendez-vous aux sources et aux récepteurs d'une diffusion multicast donnée. Les sources émettent leurs flux multicast vers le routeur RP, qui les retransmet vers les récepteurs de ce flux.

L'architecture de PIM-SM définit un autre routeur particulier, appelé DR (Designated Router). Il s'agit d'un seul routeur élu et connecté à un sous-réseau LAN, dont le rôle est de détecter les sources de trafic multicast et d'initier périodiquement les procédures d'enregistrement auprès du routeur RP. Le DR permet aussi de maintenir à jour la table des groupes actifs, de déclencher le cas échéant les opérations d'ajout ou de suppression de branches de l'arbre RPT et de transmettre le trafic multicast sur le LAN à destination de ses récepteurs locaux.

D'une manière générale, les attaques possibles sur un service de diffusion multicast peuvent être classifiées de la manière générique suivante :

- **Selon leur type.** On distingue les attaques par déni de service, qui visent à engorger les capacités des réseaux ou à saturer les ressources des routeurs, et les attaques mettant à profit les vulnérabilités des protocoles par falsification de messages de signalisation multicast.
- **Selon leur cible.** On peut distinguer les attaques par déni de service dirigées contre le plan de transfert des routeurs et celles dirigées contre le plan de contrôle des routeurs.
- **Selon leur origine.** Ces attaques peuvent provenir soit du cœur de réseau, soit de l'accès. Le cœur du réseau contient l'ensemble des routeurs multicast qui exécutent les protocoles de routage multicast. Le réseau d'accès est constitué des équipements qui exécutent les protocoles d'abonnement/désabonnement aux flux de diffusion multicast.

Voici deux types d'attaques possibles :

- Une source malveillante pourrait attaquer un arbre de distribution multicast existant en injectant un trafic parasite (des paquets quelconques dont l'adresse de destination est

l'adresse multicast du groupe visé) sur le groupe de diffusion et violerait ainsi l'intégrité du flux de diffusion légitime en ajoutant ce trafic parasite à la communication de groupe existante. Ce trafic parasite, qui est reçu par tous les récepteurs du groupe, vise à perturber la diffusion existante du flux légitime.

- Un assaillant pourrait souscrire à des milliers d'adresses de groupes et à des milliers d'adresses sources. L'envoi de requêtes IGMP/MLD par l'assaillant déclencherait de nombreux événements dans le protocole de routage multicast associé. L'énorme quantité d'entrées dans la table de routage multicast peut pénaliser les flux légaux. Cette attaque consomme aussi des ressources mémoire dans les équipements réseau gérant le trafic multicast afin de maintenir les états multicast créés et traiter les messages de routage multicast. Elle est particulièrement dangereuse pour les équipements réseau situés aux racines (ou proches des racines) des arbres de distribution multicast puisque ce sont ceux qui ont à maintenir le plus d'états multicast.

L'envoi par l'assaillant de requêtes IGMP/MLD déclenche des événements dans le protocole de routage multicast PIM déployé par l'opérateur. Le DR envoie des messages PIM Join afin de créer/prolonger les arbres de distribution multicast jusqu'au terminal assaillant. Cela crée une énorme quantité d'entrées dans la table de routage multicast TIB des routeurs, dont les limites peuvent être vite atteintes, empêchant les routeurs de fonctionner normalement et pénalisant tous les autres flux légaux.

Sachant qu'une entrée multicast $(*,G)$ dans un routeur nécessite environ 300 octets, auxquels il faut ajouter environ 150 octets par interface de sortie OIF et 20 octets supplémentaires par timer, une entrée multicast (S,G) dans un routeur nécessite environ 250 octets, auxquels il faut ajouter environ 150 octets par interface de sortie OIF et 20 octets supplémentaires par timer.

Si un attaquant provoque l'émission de 100 messages PIM Join(S,G) différents par seconde vers la même source S pendant 260 secondes (avant qu'une entrée multicast ait pu expirer), le nombre d'entrées multicast créées sur l'ensemble des routeurs multicast compris entre le site de l'attaquant et le routeur connectant la source S est égal à $100 \times 260 = 26\,000$ entrées, soit un espace mémoire nécessaire de :

$$26\,000 \times (250 + 150 + 20) = 11 \text{ Mo.}$$

Si dix attaquants provoquent l'émission de 100 messages PIM Join($*,G$) différents par seconde vers différentes sources pendant 260 secondes (avant qu'une entrée multicast ait pu expirer), le nombre d'entrées multicast créées sur le RP est égal à $10 \times 100 \times 260 = 260\,000$ entrées, soit un espace mémoire nécessaire de :

$$260\,000 \times (300 + 150 + 20) = 122 \text{ Mo.}$$

Ces exemples mettent en évidence que si des mécanismes spécifiques ne sont pas mis en place dans le réseau multicast pour empêcher ces attaques, ou à tout le moins en limiter les effets, elles peuvent très facilement impacter les ressources mémoire et processeur des routeurs multicast.

Les contre-mesures possibles pour limiter de tels impacts sont de nature diverse, comme l'illustrent les règles de sécurité suivantes.

De nombreux travaux sont en cours afin de renforcer la sécurité des flux multicast en tenant compte de la nature distribuée et dynamique des membres de groupes multicast.

Aujourd'hui, bien qu'il existe un ensemble de protocoles de routage multicast matures et disponibles, très peu d'opérateurs de télécommunications ont déployé une telle fonctionnalité, alors même que la constante augmentation du nombre d'opérateurs supportant le multicast témoigne d'une demande grandissante de la part des clients pour des services de diffusion multicast.

La supervision réseau SNMP

SNMP (Simple Network Management Protocol) est un protocole de gestion de réseau qui permet de contrôler un réseau à distance en interrogeant ses équipements, en modifiant leur configuration et en observant différentes informations liées à l'émission de données. Il peut en outre être utilisé pour gérer à distance logiciels et bases de données.

SNMP est devenu un standard TCP/IP, et son utilisation est universelle. Au même titre que HTTP, FTP ou SSH, il utilise une syntaxe abrégée ASN.1 (Abstract Syntax Notation One) par l'entremise d'une MIB (Management Information Base) pour définir les informations de management. Ces informations sont répertoriées en nombres entiers représentant des noms selon une architecture hiérarchisée respectant la syntaxe ASN.1. SNMP est bâti sur une architecture client-serveur.

La figure 11.15 illustre le principe de fonctionnement du protocole SNMP, qui repose sur la couche réseau UDP afin d'offrir ses services de supervision.

Le protocole SNMP fonctionne sur le principe des requêtes-réponses. Des alertes asynchrones peuvent être générées par des agents SNMP lorsqu'ils veulent avertir les systèmes d'administration du réseau d'un problème.

Il existe quatre sortes de requêtes et deux sortes de réponses.

Les requêtes sont les suivantes :

- GetRequest : pour obtenir une variable ;
- GetNextRequest : pour obtenir la variable suivante ;
- GetBulk : pour rechercher un ensemble de variables regroupées ;
- SetRequest : pour modifier la valeur d'une variable.

Les réponses sont les suivantes :

- GetResponse : pour permettre à l'agent de retourner la réponse au NMS (Network Management System) ;
- NoSuchObject : pour informer le NMS de l'indisponibilité de la variable.

Il existe trois versions du protocole SNMP : SNMP v1, SNMP v2 et SNMP v3. La version 2 est beaucoup plus complexe que la 1 et contient, entre autres, un niveau hiérar-

chique d'administration, avec un administrateur central. La version 3 comprend des modules de sécurité spécifiques.

Les risques viennent surtout des faiblesses du protocole SNMP lui-même, qui n'a pas été conçu pour être sécurisé dans ses versions 1 et 2.

Ses principales faiblesses sont les suivantes :

- Les transactions ne sont pas chiffrées.
- L'authentification est réalisée par un mot de passe appelé « communauté », qui est transmis en clair dans les transactions.
- Les deux modes de droits d'accès globaux sont en lecture seule et lecture-écriture pour une communauté SNMP donnée.
- Il est possible de limiter la vue sur une MIB pour une communauté SNMP donnée.
- SNMP est fondé sur le protocole UDP et permet facilement d'usurper des adresses IP.

SNMP v3 n'est pas une mise à jour des versions 1 ou 2 mais doit être employé en conjonction avec elles pour leur offrir une couche ou des briques de sécurité.

La figure 11.16 détaille les deux sous-modules de sécurité offerts par cette version v3.

Les mécanismes de sécurité offerts par SNMP v3 sont les suivants :

- Authentification par utilisateur et chiffrement grâce au modèle USM (User-based Security Model). Ces services de sécurité reposent sur des clés privées, qui ne sont pas accessibles par les requêtes SNMP. Ils s'appuient sur les fonctions de hachage HMAC/MD5-96 ou HMAC/SHA-96, qui prennent en compte les clés privées. Pour le chiffrement, c'est l'algorithme DES qui est utilisé, mais d'autres algorithmes cryptographiques sont à prévoir avec l'évolution du protocole.
- Contrôle d'accès sur plusieurs niveaux, grâce notamment au VACM (View-based Access Control Model), qui limite l'accès à un domaine d'une MIB tout en spécifiant les droits d'accès en lecture et écriture.

Pour une utilisation à des fins internes de gestion du réseau, il semble plus avantageux d'établir des tunnels IPsec d'administration avec les équipements, qui protègent tous les protocoles des couches supérieures, tels que SNMP, plutôt que d'utiliser le protocole SNMP v3.

Si SNMP doit être ouvert à l'extérieur de l'entreprise pour des raisons connues et acceptées, l'évolution vers une version v3 permet de maîtriser les droits d'accès plus sérieusement que les versions 1 et 2.

Dans tous les cas, des droits d'écriture ne doivent jamais être donnés, même à titre temporaire, en dehors de l'entité en charge de la gestion du réseau.

Règles de sécurité pour l'architecture de routage multicast relatif à l'accès

Les règles de sécurité à considérer pour l'architecture de routage multicast relatif à l'accès sont les suivantes :

- Filtrer et autoriser de manière statique les seules sources connues et légitimes en configurant des access-lists appliquées aux adresses des groupes et/ou des sources multicast des paquets multicast :

```
/* Filtrer et autoriser les sources */
ip access-list extended authorized_Sources&Groups
    permit ip @ipS @netS @ipG @netG
    deny ip any any
```

```
interface x
    ip igmp access-group authorized_Sources&Groups
```

- Filtrer et autoriser de manière statique les seuls récepteurs connus et légitimes en configurant des access-lists appliquées aux adresses IP source des messages des protocoles IGMP ou MLD :

```
/* Filtrer et autoriser les seuls récepteurs */
access-list authorised_group permit @ip1 @net1
access-list authorised_group permit @ip2 @net2
access-list authorised_group deny any
```

```
interface x
    ip igmp access-group authorised_group
```

- Configurer et appliquer des limitations aux protocoles de découverte et de gestion de groupes multicast (IGMP, MLD) afin de limiter le nombre d'états multicast au niveau global ou sur une interface donnée :

```
/* Niveau global pour igmp ou mld */
ip igmp limit number1
ip mld state-limit number2
```

```
/* Niveau interface pour igmp ou mld */
interface x
    ip igmp limit number3
```

```
interface y
    ip mld limit number4
```

- Configurer et appliquer des limitations en débit aux sources de trafic afin de limiter la quantité de trafic multicast acceptée par seconde, en entrée ou en sortie, sur une interface donnée d'un routeur :

```
/* Limitation en débit */
interface x
ip multicast rate-limit {in/out} group-list authorised_group source-list
authorised_source packetrate
```

```
/* Contrôle des groupes */
access-list authorised_group permit @ipG @netG
/* Contrôle des sources */
access-list authorised_source permit @ipS @netS
```

- Sécuriser l'échange des messages de gestion et de découverte des groupes multicast à l'aide de protocoles tels que IPsec.

Règles de sécurité pour l'architecture de routage multicast relatif au routage intradomaine

Les règles de sécurité à considérer pour l'architecture de routage multicast relatif au routage intradomaine sont les suivantes :

- Contrôle par configuration statique des adresses IP des voisins PIM :

```
access-list access-list-name permit @ip
access-list access-list-name deny any
```

```
interface x
  ip pim neighbor-filter access-list-name
```

- Filtrage statique au niveau d'un DR, avec contrôle par access-lists des adresses des groupes et/ou des sources multicast des paquets multicast :

```
ip access-list extended access-list-name
  permit ip @ipS @netS @ipG @netG
  deny ip any any
```

```
interface x
  ip access-group access-list-name in
```

- Filtrage des sources autorisées afin de restreindre au niveau du RP l'espace d'adresses source duquel on accepte des messages PIM Register :

```
access-list access-list-name permit @ip
access-list access-list-name deny any
```

```
ip pim accept-register {list access-list-name | route-map map-name}
```

- Filtrage en entrée sur une interface de tous les paquets PIM fondés sur le champ protocole :

```
ip access-list extended filtrage-PIM
  deny 103 any any
  permit ip any any
```

```
interface x
  ip access-group filtrage-PIM in
```

- Filtrage des sources et groupes à usage interne au domaine multicast par définition de « frontières multicast » :

```
access-list access-list-name permit @ip
access-list access-list-name deny any
```

```
interface x
  ip multicast boundary access-list-name
```

- Contrôle au niveau d'une interface d'un routeur multicast du débit maximal auquel une source peut émettre du trafic sur un groupe :

```
ip multicast rate-limit {in | out} group-list liste-groupe source-list
liste-source rate
```

```
access-list liste-groupe permit @ip
access-list liste-groupe deny @ip
```

```
access-list liste-source permit @ip
access-list liste-source deny any
```

- Limitation du nombre de messages PIM Register par entrée (S,G) encapsulés par seconde par un routeur DR :

```
ip pim register-rate-limit register-rate
```
- Configuration du nombre maximal d'états multicast (*,G) et (S,G) qui peuvent être créés dans la table de routage multicast d'un routeur :

```
ip multicast route-limit routes-number
```

Figure 11.15
Représentation en couches
du protocole SNMP

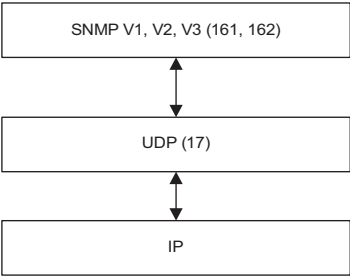
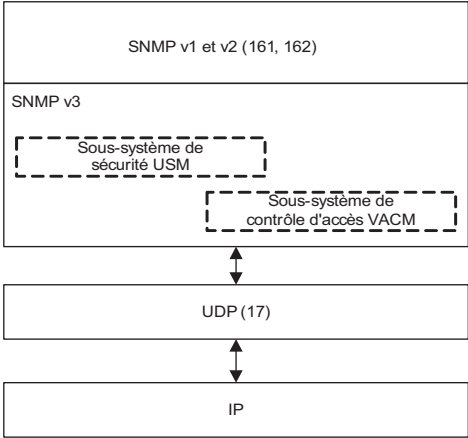


Figure 11.16
Représentation des modules
du protocole SNMP



Mise à l'heure des équipements réseau NTP

La mise à jour sur une même base de temps des horloges des équipements réseau est primordiale pour la corrélation et la correction des problèmes réseau et pour les investigations de sécurité.

Il ne faut pas confondre l'heure des équipements avec les horloges utilisées afin de synchroniser les flux de données entre les équipements afin d'éviter le fameux effet de gigue.

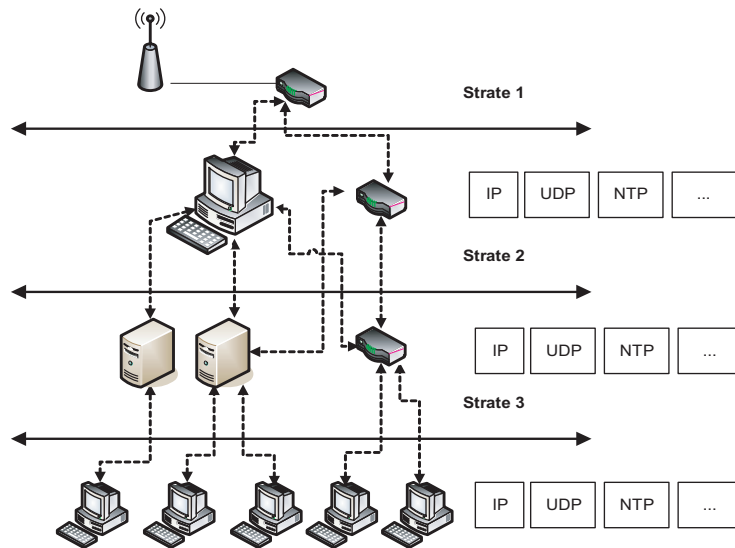
Le protocole NTP (Network Time Protocol) permet de synchroniser des systèmes entre eux, en dépit du fait que le protocole IP utilisé comme vecteur de transport fonctionne en mode non connecté, et donc sans mise à jour en temps réel des horloges.

NTP est construit sur une hiérarchie de couches, ou strates, qui agissent comme autant de vecteurs pour la synchronisation des horloges. Il est possible de définir jusqu'à 15 niveaux de couches, le 16^e niveau correspondant à une horloge non synchronisée.

Les systèmes associés à un niveau de strate 1 se synchronisent généralement sur des récepteurs radio branchés ou sur toute source sûre donnant des mesures du temps. Les systèmes associés à un niveau de strate 2 se synchronisent sur les systèmes de strate 1. Il en va de même pour les autres couches ou strates, comme illustré à la figure 11.17.

Figure 11.17

Hiérarchies des strates NTP



Une architecture NTP est bâtie à la fois sur une source d'horloge sûre et sur des niveaux de strates limités :

- Pour la source d'horloge, on s'appuie sur des antennes GPS (Global Positioning System) et non sur les sources NTP que peut offrir Internet.
- Pour les niveaux de strate, seule une étude permet de définir et de limiter au minimum le nombre de niveaux.

Dans la plupart des implémentations, il est possible de contrôler les échanges de données NTP entre une source et un serveur à l'aide d'un mot de passe partagé. En revanche, ce contrôle ne permet pas d'authentifier au sens strict du terme les échanges de données ni la confidentialité.

La résolution de noms DNS

Chaque interface d'équipement connectée à un réseau TCP/IP est identifiée au moyen d'une adresse IP unique dans le réseau. Le service DNS (Domain Name Service) permet d'associer un nom à tout équipement ayant une adresse IP.

Bien que l'utilisation des noms à la place des adresses IP ne soit pas requise par le protocole IP, l'usage des noms s'est peu à peu imposé dans le réseau Internet.

Pour que le système fonctionne, il faut qu'existe, soit au niveau de l'équipement, soit ailleurs dans le réseau, une correspondance entre nom et adresse. Pour la correspondance locale sur l'équipement, on renseigne un fichier, nommé *hosts* sur les systèmes Unix. Pour la correspondance distante, le DNS met en œuvre une base de données hiérarchique, qui peut être répartie sur plusieurs serveurs, avec répartition de charge et distribution des mises à jour de la base de données.

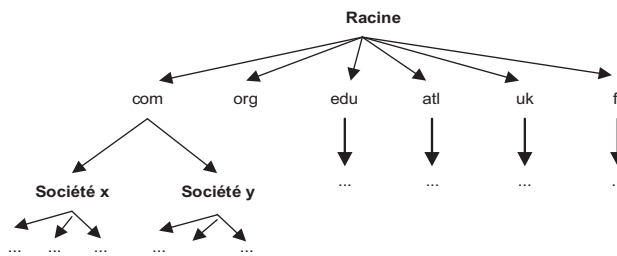
Un serveur de noms est responsable de la mise à jour des correspondances nom/adresse IP des systèmes de son domaine. Il est appelé *Authoritative Server*, ou serveur d'autorité pour le domaine. Un serveur peut déléguer l'autorité d'un ou de plusieurs sous-domaines à d'autres serveurs. Ces derniers deviennent serveurs d'autorité pour ces sous-domaines. Aucun serveur ne possède d'informations complètes sur les domaines et sous-domaines du réseau, y compris les serveurs racines. Les serveurs pointent en fait sur les serveurs de noms qui détiennent ces informations.

Un serveur de noms gère la base de données de son domaine et la liste des serveurs de noms situés jusqu'à deux niveaux de domaines en dessous de lui. Si un serveur se trouve dans l'impossibilité de répondre à une résolution nom/adresse IP, il retransmet la requête au serveur de noms ayant cette information.

La figure 11.18 illustre la hiérarchie des domaines de noms avec le domaine racine et les sous-domaines qui lui sont rattachés.

Figure 11.18

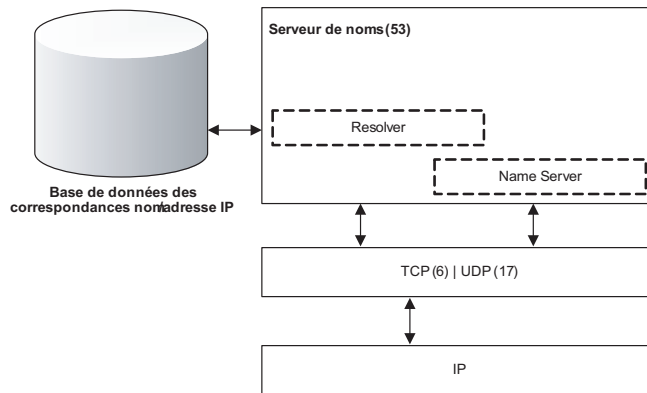
Hiérarchie des domaines de noms



Le protocole DNS est décrit par diverses RFC de l'IETF et fait l'objet de constantes améliorations. Une future version sécurisée est notamment annoncée. Il utilise la pile protocolaire IP/TCP ou IP/UDP. Son implémentation sur un serveur est composée d'un module *Resolver*, contenant des bibliothèques de routines permettant de poser des questions aux serveurs de noms, et d'un module *Name Server*, qui exécute le processus répondant aux questions de correspondance nom/adresse IP.

La figure 11.19 illustre les modules d'un serveur de noms.

Figure 11.19
Modules d'un serveur de noms



Il existe trois types de serveurs de noms :

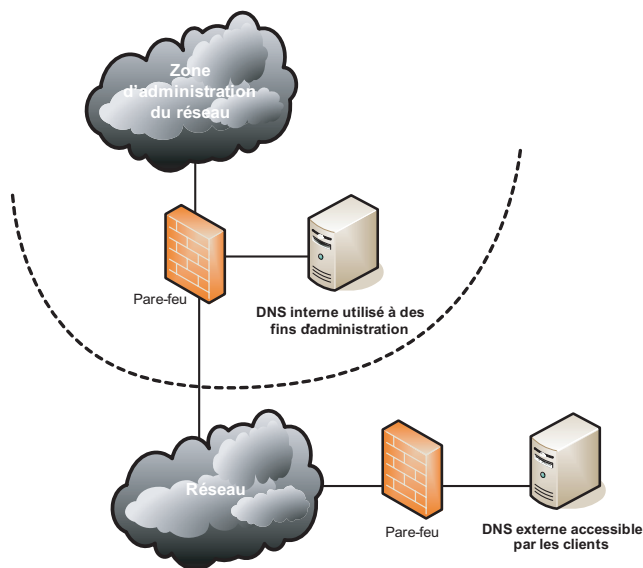
- **Serveur primaire.** C'est le serveur d'autorité sur le domaine. Il tient à jour un fichier, appelé fichier de zone, qui établit les correspondances entre les noms et les adresses IP des hosts de sa zone.
- **Serveur secondaire.** Reçoit régulièrement, par transfert de zone, les fichiers de la base de données DNS concernant la zone qu'il sert. Il est capable de répondre aux requêtes de noms IP (partage de charge) et de secourir le serveur primaire en cas de panne.
- **Serveur cache.** Pour répondre en temps réel aux demandes de résolution de noms de domaines, le serveur cache ne constitue sa base d'information qu'à partir des réponses des serveurs de noms. Il inscrit les correspondances nom/adresse IP dans un cache, avec une durée de validité (TTL) limitée et n'a aucune autorité sur le domaine. Il n'est pas responsable de la mise à jour des informations contenues dans son cache. En conséquence, la mise à jour d'une zone sur un serveur primaire peut ne pas se refléter immédiatement sur un serveur cache, car celui-ci attendra la fin de la durée de validité (TTL) pour aller interroger à nouveau le serveur primaire et s'apercevoir que l'information a été modifiée.

Sachant qu'une attaque sur un serveur DNS peut impacter immédiatement le trafic du réseau et de ses services, la sécurisation d'un tel service est cruciale. Une différence doit être faite entre les serveurs DNS à vocation de gestion interne et externe du réseau. De plus, des serveurs de noms différents doivent être déployés. Pour chaque système, le système d'exploitation doit être sécurisé au maximum. Chaque serveur DNS doit être déployé derrière un pare-feu à filtrage dynamique et sur une section de LAN entièrement réservée à cet effet, comme illustré à la figure 11.20.

Sachant que le protocole DNS est un pilier pour le réseau, des évolutions de sécurité ont été proposées afin de définir le protocole DNSsec (extensions de sécurité au protocole DNS). Les services rendus par DNSsec permettent de garantir la sécurité des données et

Figure 11.20

Séparation des serveurs de noms à un usage d'administration réseau



des transactions de données et d'offrir une architecture de distribution de clés reposant sur des algorithmes cryptographiques.

La sécurité offerte côté serveur offre les fonctionnalités suivantes :

- Chaque zone génère un ensemble de paires de clés privées/publiques.
- Les parties privées des clés signent les informations (RRsets) faisant partie intégrante de la zone.
- Les signatures sont stockées dans le fichier de zone avec les données qu'elles authentifient.
- Les parties publiques des clés sont publiées dans le fichier de zone et peuvent faire l'objet de requêtes DNS classiques.

Côté client, la connaissance de la clé publique d'une zone permet de vérifier les signatures et de contrôler ainsi l'authenticité et l'intégrité des informations contenues dans la zone.

Cela nécessite cependant la connaissance des clés de toutes les zones avec lesquelles le resolver est susceptible de communiquer.

En résumé

La gestion d'un réseau est constituée de nombreux domaines, tous aussi importants pour assurer de bout en bout la sécurité du réseau et de ses services.

Nous avons détaillé dans cette partie un ensemble de techniques permettant de renforcer la sécurité des accès, des authentifications, etc. Ces techniques peuvent être mises en

œuvre pour satisfaire les exigences dictées par la politique de sécurité. Rappelons une fois encore que c'est la politique de sécurité et ses objectifs qui doivent être à l'amont des techniques, et non le contraire.

La partie suivante aborde les contrôles de sécurité externe et interne qui permettent d'établir des tableaux de bord de la sécurité réseau. Cette étape vise avant tout à vérifier l'application de la politique de sécurité.