



Laboratoire
Supinfo des Technologies
Cisco

CCNP 1 - Essentiel

Building Scalable Cisco Internetworks

Auteurs : GUILLEMOT Erwan & ROBIN Eric

Relecture : ROBIN Eric

Version 1.1 – 11 Août 2004



SUPINFO - Ecole Supérieure d'Informatique de Paris
23. rue de Château Landon 75010 Paris
Site Web : <http://www.supinfo.com>

Laboratoire SUPINFO des Technologies Cisco

Site Web : www.labo-cisco.com – E-mail : labo-cisco@supinfo.com

Ce document est la propriété de SUPINFO et est soumis aux règles de droits d'auteurs

Table des matières

1. Gestion de la croissance d'un réseau modulaire.....	4
1.1. Problèmes créés par la congestion réseau.....	4
1.2. Symptômes de la congestion réseau.....	4
1.2.1. Trafic excessif.....	5
1.2.2. Paquets perdus.....	5
1.2.3. Retransmission des paquets.....	5
1.2.4. Tables de routage incomplètes.....	5
1.2.5. Liste des serveurs incomplète.....	5
1.2.6. Des erreurs dans l'application du protocole Spanning-Tree.....	6
1.3. Création d'un réseau respectant les critères requis.....	6
1.3.1. Modèles de conception de réseaux.....	6
1.3.2. Modèle à 3 couches.....	7
1.4. Réduction du trafic réseau.....	8
1.4.1. Listes de contrôle d'accès.....	8
1.4.2. Interface nulle.....	9
1.5. Priorités entre les trafics.....	10
1.6. Optimisation CPU et méthodes supplémentaires de contrôle de trafic.....	10
1.6.1. Fast, Autonomous et Silicon Switching.....	10
1.6.2. Emplacement client/serveur.....	11
1.6.3. ip helper address.....	11
1.6.4. Le tunneling sur IP.....	12
2. Adressage IP.....	13
2.1. Bases de l'adressage IP.....	13
2.2. Prefix Routing / CIDR.....	13
2.2.1. Problèmes d'adressage sur le réseau mondial.....	13
2.2.2. Calcul du masque de sur-réseau pour le CIDR.....	13
2.2.3. Diminution des tables de routage des routeurs de l'Internet.....	14
2.3. VLSM (Variable-Length Subnet Mask).....	15
2.3.1. Introduction.....	15
2.3.2. Rappels formels sur le subnetting.....	15
2.3.3. Concevoir un plan d'adressage selon la technique VLSM.....	15
2.3.4. Considérations sur les RFC 950 et 1878.....	16
2.3.5. Allocation des adresses VLSM.....	17
2.4. Agrégat de routes.....	18
2.5. Configuration de l'agrégat de routes.....	19
2.5.1. Agrégat automatique.....	19
2.5.2. Agrégat manuel.....	19
2.5.3. Sous-réseaux discontinus.....	20

3. Routage IP	21
3.1. Protocoles routés et protocoles de routage.....	21
3.2. Table de routage.....	22
3.3. Fonctions de commutation et de routage	23
3.4. Protocoles de routage à vecteur de distance et à état des liens	24
3.4.1. Vecteur de distance	25
3.4.2. Etat des liens.....	25
3.5. Système autonome - Protocoles de routage intérieurs et extérieurs.....	26
3.6. Redistribution de routes	26
3.7. Distribute Lists.....	29
3.8. Route Maps	30
4. Protocole EIGRP	32
4.1. Caractéristiques.....	32
4.2. Termes et définitions	33
4.3. Métriques	34
4.4. Protocole Hello	35
4.4.1. Neighbor Table.....	36
4.4.2. Topology Table	36
4.5. DUAL	37
4.5.1. Choix d'un successeur.....	38
4.5.2. Ajout d'un réseau dans la Topology Table	38
4.5.3. Suppression d'une route ou routeur de la Topology Table	39
4.5.4. Recherche d'un chemin alternatif vers un réseau distant	40
4.6. Fonctionnement avec IPX.....	41
4.7. Commandes	41
4.8. Configuration.....	43
4.8.1. Configuration pour IP.....	43
4.8.2. Configuration pour IPX.....	43
4.8.3. Préoccupation concernant la bande passante et configuration sur un réseau NBMA	44

1. Gestion de la croissance d'un réseau modulaire

Lors de la création d'un réseau d'entreprise, il est indispensable de prendre en compte différents paramètres. En effet, la phase la plus importante est le regroupement des besoins des utilisateurs. De plus, il sera indispensable de mesurer la structure existante, si elle existe, et les flux de données transitant sur le réseau afin d'optimiser les réponses apportées aux besoins du réseau.

Voici les éléments auxquels doit répondre un réseau :

- Fiabilité
- Réactivité
- Efficacité
- Adaptabilité
- Accessibilité
- Evolutivité

1.1. Problèmes créés par la congestion réseau

L'utilisation du réseau dans les entreprises a augmenté de façon exponentielle. Il est important de permettre à son réseau de croître en même temps que les besoins des utilisateurs.

Les applications deviennent de plus en plus complexes, de plus en plus gourmande en bande passante et fonctionnent de plus en plus en mode client/serveur.

De plus, une mauvaise implémentation d'un réseau peut aboutir à un réseau non évolutif et donc incapable de répondre aux besoins de l'organisation. Les différents facteurs représentent les principales causes de congestion réseau.

Voyons les différents symptômes de la congestion.

1.2. Symptômes de la congestion réseau

Voici les principales conséquences d'une congestion réseau :

- Un trafic excessif
- Des paquets perdus
- Une retransmission des paquets
- Des tables de routage incomplètes
- Des listes de serveurs incomplètes
- Des erreurs dans l'application du protocole Spanning-Tree

1.2.1. Trafic excessif

Le trafic excessif se caractérise par une incapacité du média à fournir les moyens fonctionnels nécessaires au flux de données. A ce moment, se pose le problème intrinsèque de la technologie Ethernet, architecture de loin la plus utilisée. En effet, l'accès au média est non déterministe, et chaque interface réseau doit attendre le "silence" avant d'émettre. Un trafic excessif peut donc empêcher les utilisateurs d'utiliser le réseau de façon fiable, et leurs trames de données feront l'objet de collisions fréquentes.

Une collision se traduit par l'émission d'un signal de bourrage qui a pour effet de stopper toute communication au sein du domaine de collision, jusqu'à la fin du signal de bourrage et de la durée minimale entre chaque trame (9,6ms pour de l'Ethernet).

Au bout de 16 tentatives, si l'interface réseau n'a toujours pas réussi à émettre la trame, elle envoie un rapport d'erreur au processus générateur de la trame (typiquement l'application). Ainsi, l'application renvoie la trame, et ceci aboutit à une congestion plus importante (notion de cercle vicieux).

1.2.2. Paquets perdus

Un des effets de la congestion est que tous les paquets ne peuvent transiter via le réseau. Les files d'attente et les mémoires tampon dans les équipements saturent et doivent ainsi "jeter" les paquets, ce qui a pour effet de créer un timeout sur les dispositifs finaux (stations de travail par exemple).

1.2.3. Retransmission des paquets

Lorsque les paquets sont perdus, les couches responsables de l'intégrité du message (couches 4 et 7) retransmettent les paquets perdus. Ceci cause une augmentation dangereuse de la congestion réseau.

1.2.4. Tables de routage incomplètes

Ces problèmes de trafic excessif, paquets jetés, retransmission de paquets affectent aussi les transmissions des tables de routage qui utilisent le même média que les paquets utilisateur.

Les tables de routage peuvent donc ne pas arriver à destination, ou arriver incomplètes ce qui implique une incohérence dans la vue de la topologie par les routeurs et à l'augmentation du temps de convergence.

1.2.5. Liste des serveurs incomplète

De la même manière que pour les mises à jour de routage, les paquets de mises à jour des serveurs peuvent être perdus. Il en résulte une inaccessibilité des serveurs pour les utilisateurs.

1.2.6. Des erreurs dans l'application du protocole Spanning-Tree

Le protocole Spanning-Tree permet au réseau commuté (couche 2) de maintenir une topologie avec redondance de liens sans effet de boucle de commutation. Afin d'organiser l'arborescence Spanning-Tree, les commutateurs et/ou les ponts échangent des messages appelés BPDU (Bridge Protocol Data Unit). Ainsi, un lien redondant est bloqué.

Si les messages BPDU n'arrivent plus à destination, le lien redondant bloqué risque de redevenir actif et donc utilisable à la fin de la période maximale (Max Age timer). Ceci a pour conséquence de causer des boucles de commutation et des tempêtes de broadcast.

1.3. Création d'un réseau respectant les critères requis

1.3.1. Modèles de conception de réseaux

Bien qu'il soit nécessaire d'appréhender la congestion sur un réseau, et donc de mettre en place des solutions pour la réduire, il est encore plus important de construire un réseau évolutif, permettant de croître selon les besoins des utilisateurs.

Il existe deux structures de modèle de réseaux :

- Hiérarchique :
 - Réseau divisé en couches
 - Fonction(s) précise(s) associée(s) à chaque couche
- Maillée :
 - Topologie linéaire
 - Tous les dispositifs ont les mêmes fonctions

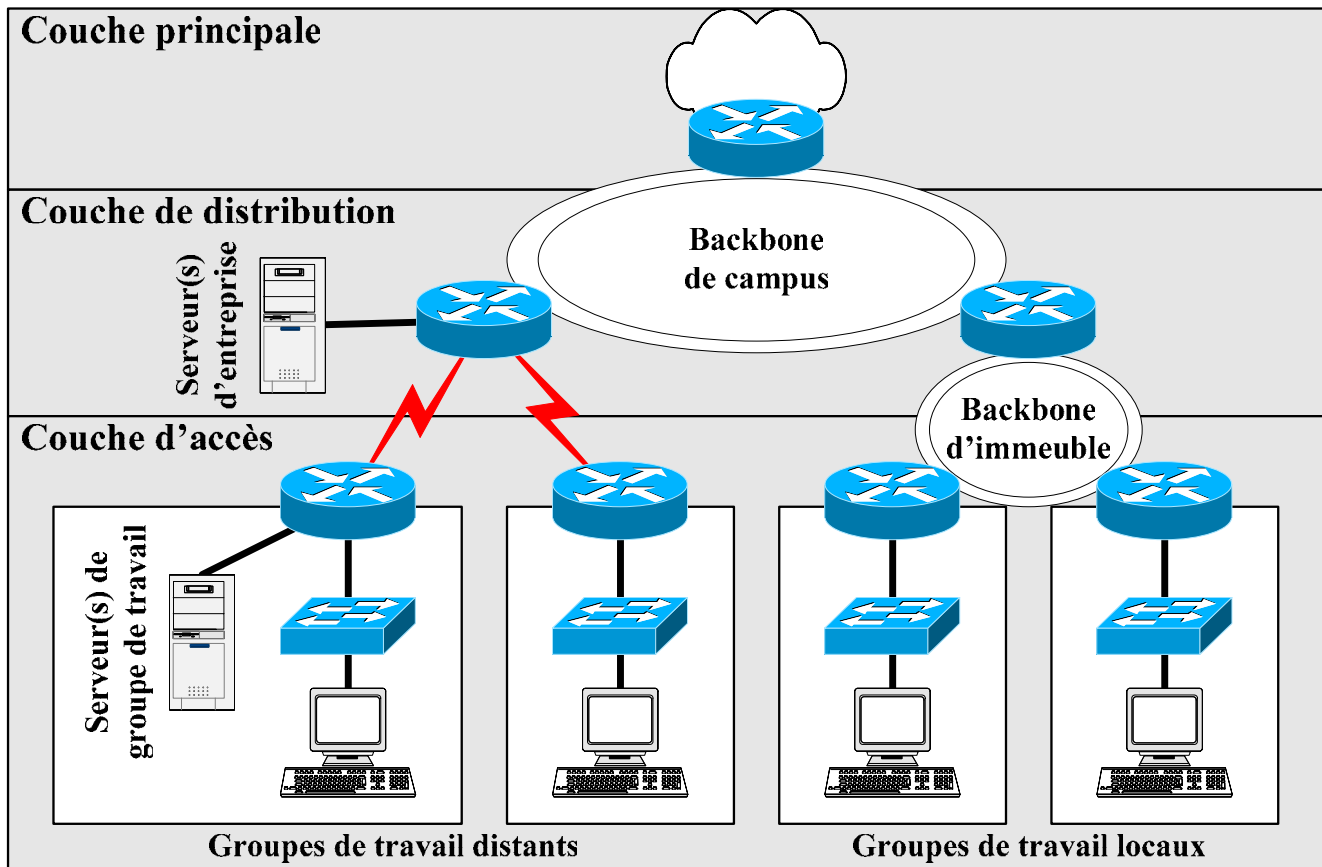
L'intérêt d'utiliser un modèle de réseau hiérarchique lors de la conception est de :

- Faciliter les modifications et la compréhension du réseau (réseau modulaire)
- Limiter les coûts et la complexité des mises à niveau du réseau (appliquées à un sous-ensemble uniquement)
- Limiter les coûts de construction et d'élaboration du réseau
- Faciliter l'identification des points de défaillance

L'utilisation d'un modèle hiérarchique procure des avantages tels que :

- Evolutivité
- Facilité de mise en œuvre
- Facilité de dépannage
- Prévisibilité
- Prise en charge de protocoles
- Facilité de gestion

1.3.2. Modèle à 3 couches



Les couches de ce modèle sont :

- Couche principale (centrale) : Assure l'optimisation du transport entre les sites
- Couche de distribution : Assure une connectivité fondée sur les politiques
- Couche d'accès : Permet aux utilisateurs et aux groupes de travail d'accéder au réseau

La couche principale :

- Assure la communication (la plus rapide possible) entre les sites éloignés
- Comporte habituellement des liaisons point-à-point
- Aucun hôte présent, que des unités de communication
- Services présents (Frame Relay, T1/E1, SMDS) loués auprès d'un fournisseur de services
- Ne s'occupe pas du filtrage ou de la sécurité
- Exigence de chemins redondants pour la continuité de service en cas de panne
- Fonctionnalités des protocoles de routage très importantes (partage de charge, convergence rapide)
- Utilisation efficace de la bande passante

La couche distribution :

- Fournit des services à plusieurs LAN au sein d'un WAN (backbone de campus)
- C'est l'emplacement du backbone du WAN (de type Fast Ethernet)
- Sert à interconnecter des immeubles
- Emplacement des serveurs d'entreprise (DNS, messagerie centralisée)
- A pour rôle de définir les frontières (sous la forme de politiques)
- Prend en charge le filtrage (ACL) et le routage inter VLAN

La couche d'accès :

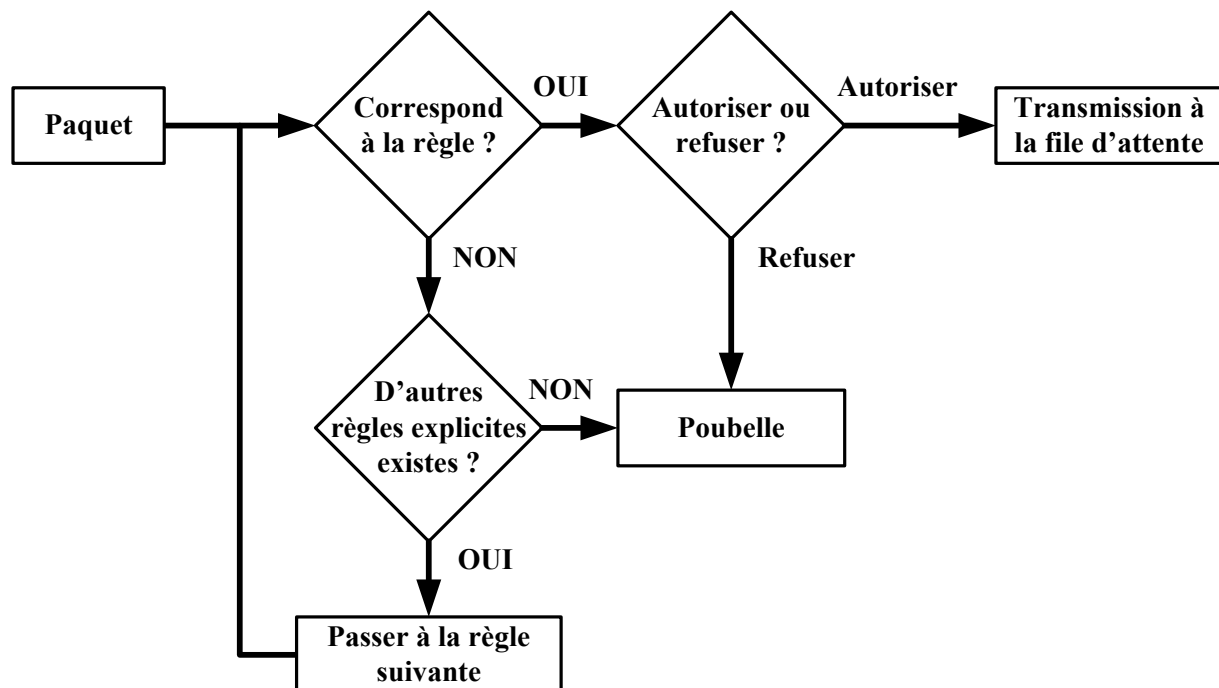
- Partie LAN du réseau
- Emplacement des hôtes (utilisateurs)
- Emplacement des serveurs de groupe de travail (stockage des fichiers, impression)
- Possibilité d'utiliser des ACLs afin de déterminer les besoins précis d'un groupe d'utilisateurs
- Partage et/ou commutation de la bande passante, microsegmentation et VLAN
- Regroupement des utilisateurs selon leur fonction, leurs besoins
- Isolation du trafic de broadcast destiné à un groupe de travail ou à un LAN

1.4. Réduction du trafic réseau

1.4.1. Listes de contrôle d'accès

Les listes de contrôle d'accès (ACLs) sont des instructions qui expriment une liste de règles, imposées par l'opérateur, donnant un contrôle supplémentaire sur les paquets reçus et transmis par le routeur. Les ACLs sont capables d'autoriser ou d'interdire des paquets, que ce soit en entrée ou en sortie vers une destination.

Elles opèrent selon un ordre séquentiel et logique, en évaluant les paquets à partir du début de la liste d'instructions. Si le paquet répond au critère de la première instruction, il ignore le reste des règles et il est autorisé ou refusé. Il ne peut y avoir qu'une seule ACL par protocole, par interface et par sens.



Une ACL est identifiable par son numéro ou son nom, attribué suivant le protocole et le type :

- ACL standard
- ACL étendue
- ACL nommée

Une ACL standard permet d'autoriser ou d'interdire des adresses spécifiques ou bien un ensemble d'adresses ou de protocoles.

Une ACL étendue permet de faire un filtrage plus précis qu'une ACL standard, elle permet également d'effectuer un filtrage en fonction du protocole.

Depuis la version 11.2 d'IOS, il est possible d'utiliser les ACLs nommées. Les ACLs nommées permettent l'identification par des chaînes alphanumériques plutôt que par la représentation numérique actuelle.

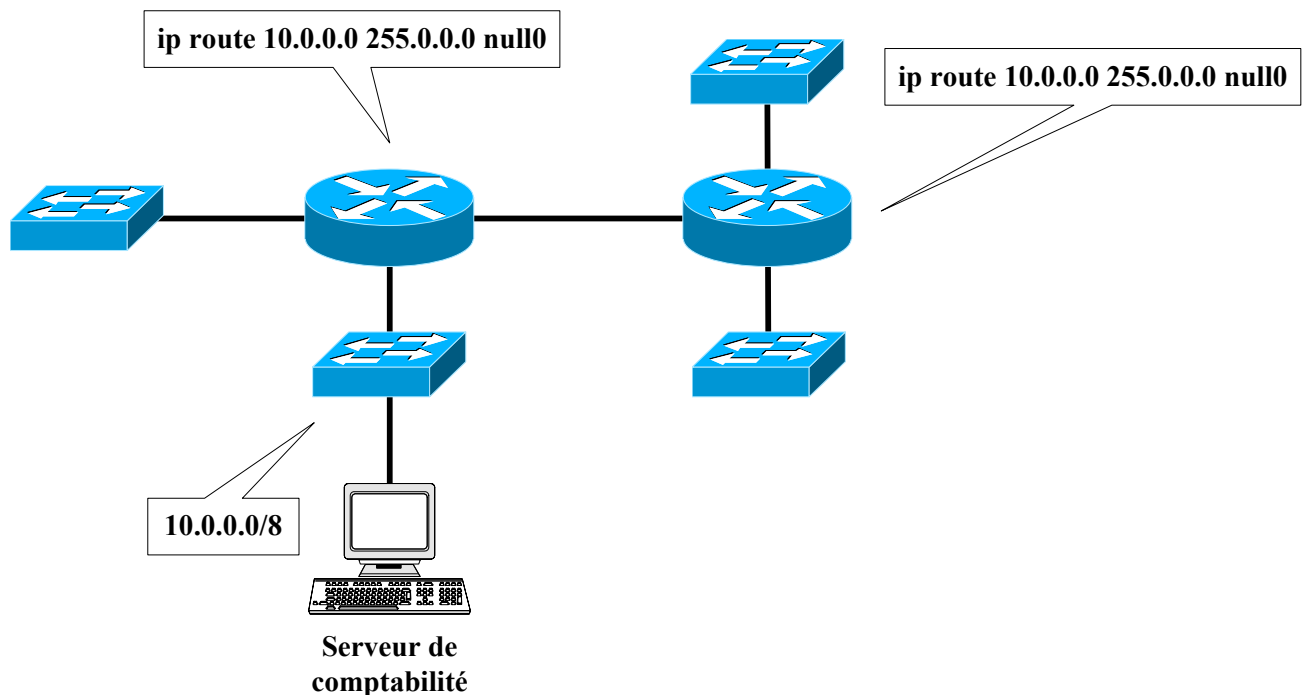
Vous pouvez utiliser les ACLs nommées dans les situations suivantes :

- Identifier intuitivement les listes de contrôle d'accès à l'aide d'un code alphanumérique
- Configurer plus de 99 ACLs standard et plus de 100 ACLs étendues dans un routeur pour un protocole donné

1.4.2. Interface nulle

L'interface nulle est une interface virtuelle qui n'a aucune existence physique. Tous les paquets qui sont envoyés sur cette interface disparaissent. Il s'agit d'une bonne alternative aux ACLs coûteuses en ressources processeur.

On ne les trouve qu'en tant qu'interface locale de sortie pour une route statique. Prenons pour exemple la commande **ip route 10.0.0.0 255.0.0.0 null0**. Dans cet exemple, tous les paquets qui arrivent au routeur ayant pour destination le réseau 10.0.0.0 seront envoyés vers l'interface nulle et donc ne pourront accéder à ce réseau.



Au vu de cette topologie, personne ne pourra accéder au réseau où se situe le serveur de la comptabilité, excepté les utilisateurs appartenant au réseau 10.0.0.0.

1.5. Priorités entre les trafics

Il existe plusieurs types de priorités. Elles sont implémentées au niveau d'une interface et sont appliquées à la file d'attente de cette interface.

Par défaut, la technique de Weighted Fair Queuing est appliquée sur les interfaces ayant une bande passante inférieure ou égale à celle d'une T1. La méthode FIFO est appliquée par défaut pour les interfaces possédant une bande passante supérieure.

Le processus de file d'attente analyse les caractéristiques du trafic sur le lien, selon la taille et le type de paquet transmis pour distinguer le trafic interactif du trafic des fichiers utilisateur.

Voici les 4 méthodes de priorité configurées manuellement à l'aide des listes de contrôle d'accès :

- **Priority Queuing** : Cette méthode scinde l'interface du trafic sortant en 4 canaux virtuels. Une importance (ou priorité) est affectée à chacune de ces interfaces virtuelles. Selon le type de trafic qui doit sortir de l'interface, il est dirigé vers tel ou tel canal virtuel, qui possédera sa propre priorité.
- **Custom Queuing** : L'interface est divisée en plusieurs sous-files d'attente. Chaque file d'attente a un seuil définissant le nombre d'octets qui peuvent être envoyés avant que la file d'attente suivante commence à être activée. Il est ainsi possible de déterminer le pourcentage de bande passante affecté à chaque protocole.
- **Class-Based Weighted Fair Queuing (CBWFQ)** : Cette méthode étend les possibilités de la méthode Weighted Fair Queuing en fournissant un support pour les classes de trafic utilisateur. Il faut définir des classes de trafic basées sur différents critères (protocoles, ACLs, interfaces de trafic entrant, etc.). Les paquets qui correspondent à ces critères constituent le trafic pour cette classe. Une file d'attente est réservée pour chaque classe, et le trafic appartenant à cette classe est dirigé vers cette file d'attente.
- **Low-Latency Queuing (LLQ)** : Cette méthode apporte une file d'attente prioritaire au CBWFQ, dans le sens où l'on définit le trafic le plus important (delay sensitive), telle que la voix ou les données. Ainsi, le trafic qui est considéré comme le plus important sera transmis avant que toutes les autres files d'attente commencent à se vider.

1.6. Optimisation CPU et méthodes supplémentaires de contrôle de trafic

Dans les premières versions de Cisco IOS (antérieures à 10.3) des ACLs complexes, empêchaient la mise en cache des tables de routage. Désormais, la mise en cache des tables de routage est possible puisque des méthodes, telles que Fast, Autonomous et Silicon Switching ont été mises en place et améliorées au fur et à mesure des versions de Cisco IOS.

1.6.1. Fast, Autonomous et Silicon Switching

Ces méthodes ont été créées afin de permettre au routeur d'envoyer les données le plus rapidement possible. Ces techniques mettent en cache les décisions de routage que le routeur a pris ; ceci implique qu'un paquet provenant d'une même source vers une même destination ne passera pas par le processus de routage, et pourra être transmis directement vers l'interface de sortie.

1.6.2. Emplacement client/serveur

L'emplacement des clients par rapport aux serveurs est un facteur extrêmement critique dans le design d'un réseau.

En effet, la donne actuelle est la centralisation des serveurs dans une ferme afin d'avoir une administration centralisée. Cependant, un problème est aussitôt posé dans le sens où les routeurs sont des pares feu de broadcast et qu'ils empêchent la connexion des clients aux serveurs si cette connexion doit se faire par trames de broadcast.

La solution apportée est la technique des "**ip helper address**".

1.6.3. ip helper address

L'**ip helper address** remplace l'adresse de broadcast reçue par une interface par une adresse précisée lors de la configuration. Cela fonctionne donc sur le principe de la translation d'adresses.

Un **ip helper address** est configuré sur une interface pour le trafic entrant. L'adresse de destination peut être un serveur en particulier ou un sous-réseau.

On peut mettre en place plusieurs **ip helper address**, ce qui représente une bonne méthode lors de la mise en place de serveurs secondaires (de secours).

L'**ip helper address** transmet par défaut les broadcasts pour les ports UDP suivants :

- Time (37)
- TACACS (49)
- DNS (53)
- BOOTP Server (67)
- BOOTP Client (68)
- TFTP (69)
- NetBIOS Name Server (137)
- NetBIOS Datagram Service (138)

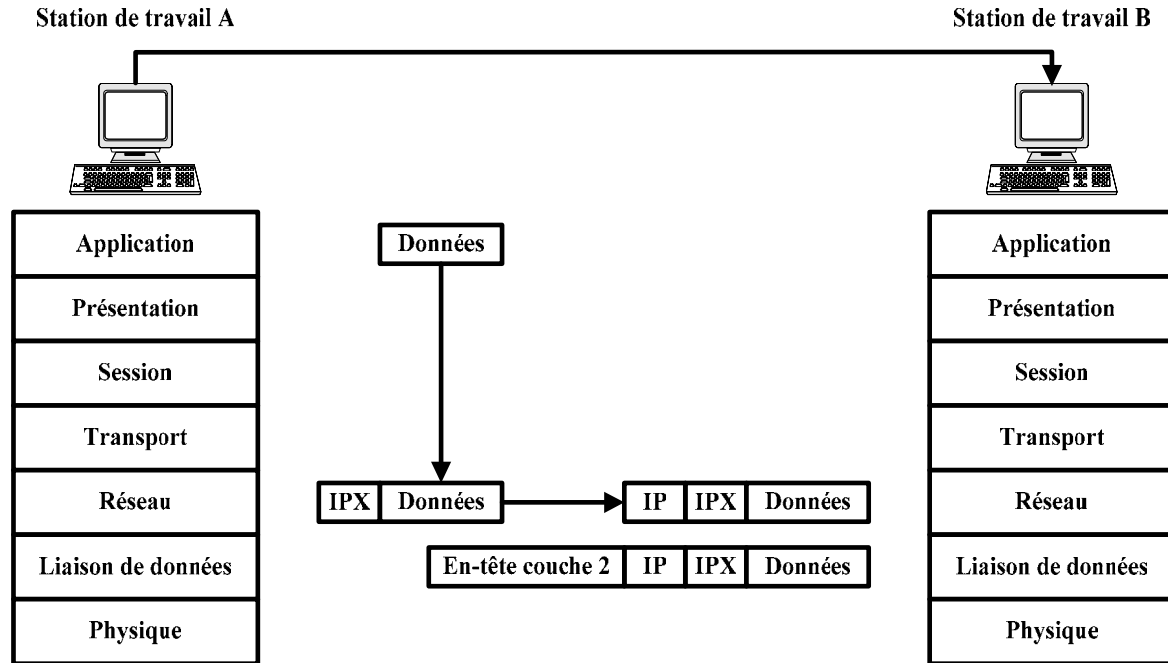
En addition des **ip helper address**, on peut mettre en place des instructions dites "**ip forward-protocol**" qui permettent d'identifier un type de port en particulier. Il est ainsi possible d'ajouter des types de trafic (numéros de port) qui pourront être ajoutés aux **ip helper address** ou être supprimés.

Commandes utilisées :

- **ip helper address {adresse}**
 - Mode de configuration d'interface
 - Le paramètre **adresse** peut être une adresse d'hôte ou une adresse de broadcast.
- **ip forward-protocol {udp [port] | nd | sdns}**
 - Mode de configuration d'interface
 - Le paramètre **udp [port]** permet de spécifier le trafic, via son numéro de port UDP, à ajouter.

1.6.4. Le tunneling sur IP

Faire du tunneling consiste à encapsuler un protocole d'une couche spécifique du modèle OSI dans un autre protocole de la même couche ou d'une couche supérieure.



Dans cet exemple, les données de la couche application passent à travers le modèle OSI, jusqu'à la couche 3 où le paquet est encapsulé selon le protocole choisi (IPX dans ce cas). Ensuite ce paquet est encapsulé dans un paquet IP, passe à la couche 2 pour recevoir l'entête de trame.

L'utilisation d'un tunnel n'est pas justifiée pour optimiser le réseau ou en faciliter la compréhension. L'administrateur de la couche principale n'a plus à se soucier des variations des différents protocoles d'extrémités.

De plus, le trafic peut utiliser les avantages du protocole IP et de ses algorithmes de routage. Les deux paires distantes communiquent à travers un simple lien point à point, bien qu'ils soient séparés par plusieurs routeurs.

Cependant, quelques considérations doivent être prises en compte lors de la création d'un tunnel IP :

- Le délai et la latence créés par le tunnel peuvent causer des timeouts.
- Il faudra autant de tunnels que de liens, puisque le tunnel est vu comme un lien point à point.
- Le tunnel peut causer des incohérences dans les tables de routage. Il sera en effet vu comme un lien préféré (1 ou 2 sauts maximum), alors qu'il peut représenter dans la réalité un lien passant par 10 routeurs.

Pour configurer un tunnel :

- Rentrer en mode de configuration de l'interface du tunnel :
 - **interface tunnel {numéro de l'interface}**
- Spécifier l'interface source et destination du tunnel :
 - **tunnel source {numéro de l'interface | adresse IP}**
 - **tunnel destination {nom d'hôte | adresse IP}**

Il faut faire ces manipulations sur les deux routeurs d'extrémités, et l'adresse IP source et l'adresse IP destination du premier routeur doivent correspondre à l'adresse IP destination et à l'adresse IP source du second routeur respectivement.

2. Adressage IP

2.1. Bases de l'adressage IP

Une adresse IP est une adresse généralement exprimée en notation décimale pointée de 32 bits. Le modèle TCP/IP reste l'ensemble de protocoles de communication de l'Internet dans le sens où il s'agit du modèle le plus évolutif, le plus fiable et le plus simple à implémenter. En effet, les adresses IP ne possèdent pas de partie réseau fixe, ce qui permet d'organiser des regroupements d'adresses, essentiels au fonctionnement du réseau mondial.

2.2. Prefix Routing / CIDR

Prefix routing, mieux connu sous le nom de CIDR (Classless InterDomain Routing), est possible grâce aux nouveaux protocoles de routage qui incluent le masque de sous réseau dans les mises à jour de routage.

Tous les protocoles de routage IP sont Classless exceptés RIPv1 et IGRP.

2.2.1. Problèmes d'adressage sur le réseau mondial

Pour les petites entreprises (50 hôtes) qui souhaitent être connectées à Internet, une adresse de classe C était indispensable. Une adresse de classe C permet à 253 hôtes d'être connectés. Pour les entreprises de moyenne taille (Plus de 255 hôtes et moins de 65000) il était nécessaire d'obtenir une adresse de classe B.

On se rend donc bien compte qu'il y a alors un gâchis dans l'attribution des classes d'adresses. Le CIDR apporte une réponse à ce problème.

Le principe du CIDR est de regrouper des classes contiguës d'adresses IP afin de fournir au client la quantité d'adresses IP la plus précise possible par rapport à ses besoins, et de diminuer ainsi la taille des tables de routage grâce à un masque de sur-réseau.

2.2.2. Calcul du masque de sur-réseau pour le CIDR

On commence par définir le nombre d'utilisateurs sur le réseau.

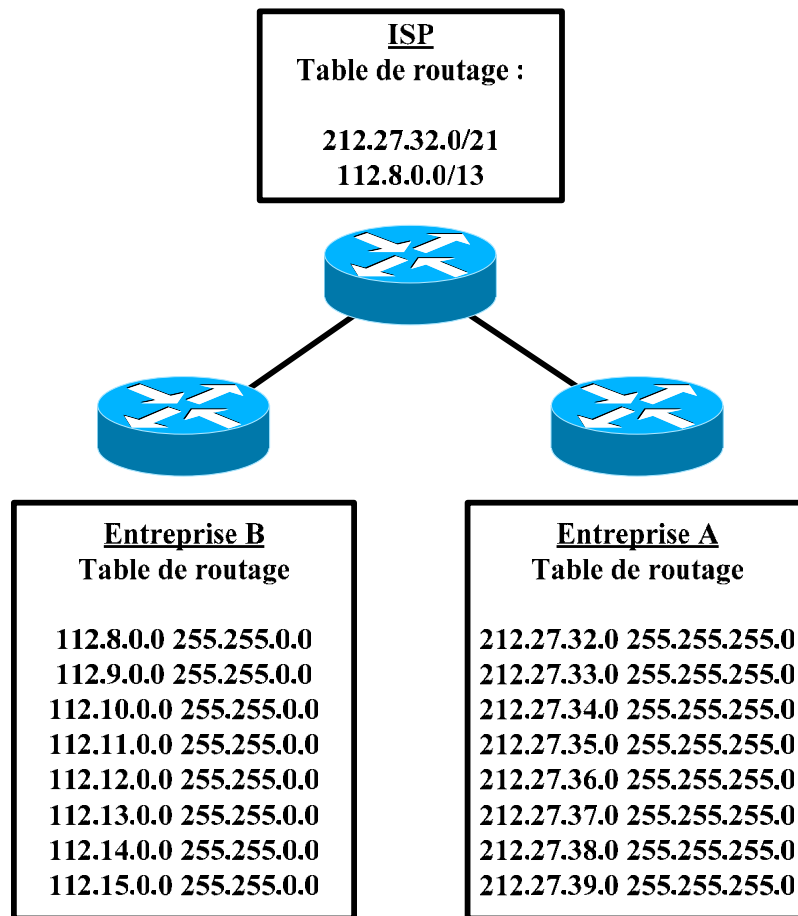
On calcule le nombre de bits nécessaires pour coder le nombre d'hôtes.

On emprunte le nombre nécessaire de bits à la partie hôte.

On met ces bits à 0 et tous les bits précédents à 1.

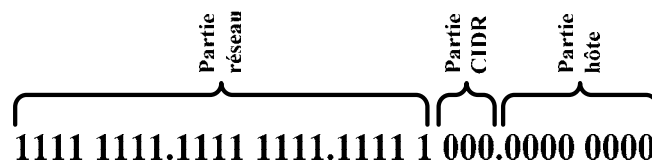
On convertit en décimal.

2.2.3. Diminution des tables de routage des routeurs de l'Internet



Pour trouver les blocs contigus d'adresses IP, il suffit de compter le nombre de bits empruntés à la partie réseau de l'adresse IP. Soit x ce nombre, on aura des blocs contigus de 2^x adresses.

Dans le schéma ci-dessus, on retrouve 8 blocs d'adresses contigus dans chaque entreprise. On a dû emprunter 3 bits ($2^3 = 8$). Pour l'entreprise A, les adresses IP sont des classes C ; on garde donc les 8 derniers bits à 0 pour la partie cliente, et l'on a emprunté 3 bits à la partie réseau. Le masque de sur-réseau est donc le suivant :



On se rend donc compte qu'au niveau de la table de routage de l'ISP, le nombre d'entrées n'est pas de 16 mais de 2. Il y a donc un intérêt très important dans la mise en place du CIDR :

- Réduction du trafic utilisé dans les échanges de tables de routage
- Diminution de l'utilisation du CPU des routeurs
- Meilleure flexibilité

2.3. VLSM (Variable-Length Subnet Mask)

2.3.1. Introduction

VLSM est utilisé au niveau d'une organisation et représente une extension au CIDR, qui lui est conçu pour le réseau mondial. Il apporte une flexibilité indispensable aux besoins des entreprises dont les segmentations d'utilisateurs ne sont pas homogènes. Cette technique permet d'utiliser au mieux les adresses IP disponibles et d'assurer un design hiérarchique fiable en permettant l'agrégat de routes et une meilleure documentation. Ceci se caractérise par une évolution du masque de sous-réseau selon l'endroit où l'on se trouve dans l'organisation.

Les protocoles supportant le VLSM sont les suivants :

- RIPv2
- OSPF
- BGP
- IS-IS
- EIGRP

2.3.2. Rappels formels sur le subnetting

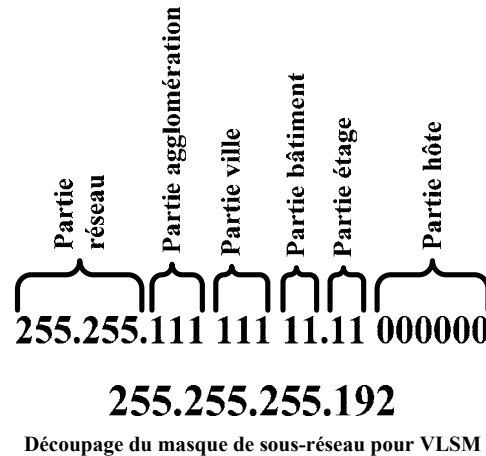
Les points importants concernant le subnetting sont les suivants :

- Il faut toujours prendre en compte la règle suivante : on ne peut pas utiliser les portions de sous-réseaux formées que de 0 ou de 1 binaires. Une fois cette règle appliquée, il n'est plus nécessaire de l'appliquer aux sous-réseaux enfants. Cette règle est appelée la règle du 2^n-2 .
- Un sous-réseau peut contenir des hôtes ou être redéfini en plusieurs sous-réseaux.
- Le protocole de routage utilisé doit inclure le masque de sous-réseaux dans les mises à jour de routage.
- Les bits d'ordre le plus haut des sous-réseaux devant être regroupés en agrégat doivent être identiques.
- Les décisions de routage sont appliquées sur l'intégralité du sous-réseau.
- Le routeur, pour prendre ses décisions de routage, va du sous-réseau le plus spécifique au sous-réseau le plus général.

2.3.3. Concevoir un plan d'adressage selon la technique VLSM

Pour concevoir un plan d'adressage hiérarchique, il faut procéder comme suit :

- Recenser le nombre total d'utilisateurs sur le réseau (prévoir une marge pour favoriser l'évolutivité du réseau).
- Choisir la classe d'adresse la plus adaptée à ce nombre.
- Partir du plus haut de l'organisation (couche principale) et descendre au plus près des utilisateurs (couche accès).
- Décompter les entités au niveau de chaque couche. Par exemple, on retrouve 3 grandes agglomérations, 7 villes, au moins 2 bâtiments dans chaque ville, au moins 2 étages par bâtiments et au moins 30 utilisateurs par étage.
- Pour chacune de ces entités, réserver le nombre de bits nécessaire en prévoyant l'évolutivité du réseau (par exemple, il pourrait y avoir deux nouvelles agglomérations).
- Calculer le masque de sous-réseau à chaque niveau de l'organisation.



2.3.4. Considérations sur les RFC 950 et 1878

Ces deux RFC (Request For Comment), nommées "Internet Standard Subnetting Procedure" et "Variable-Length Subnet Table For IPv4", établissent qu'on ne doit pas retrouver tous les bits à 0 ou à 1 dans les portions d'adresses suivantes :

- La portion Internet (Partie Classful de l'adresse)
- La portion sous-réseau
- La partie hôte

La règle pour calculer le nombre de sous-réseaux ou le nombre d'hôtes disponibles est $2^n - 2$ (n étant le nombre de bits) selon la RFC 950.

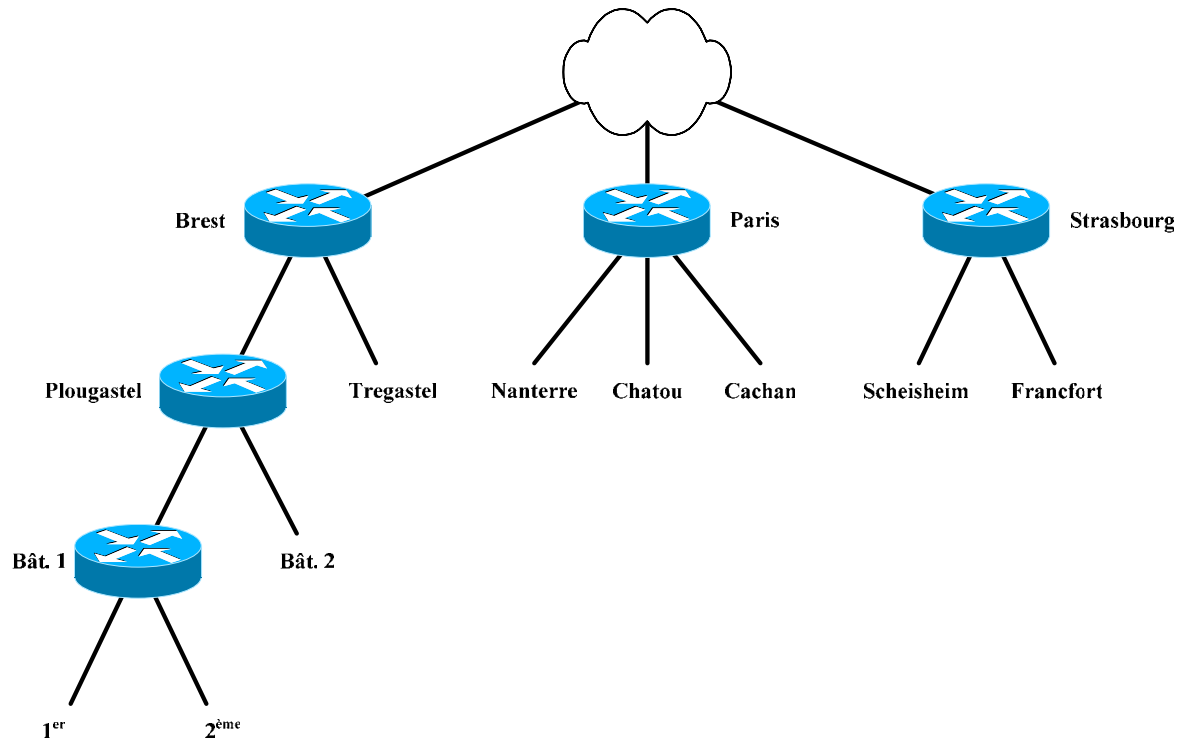
Cette règle est désormais légèrement modifiée au regard de la RFC 1878. Le nombre de sous-réseaux est calculé par la formule 2^n , où n est le nombre de bits étendus (partie au-delà de la partie Classful) du masque de sous-réseau VLSM. Il est alors possible d'utiliser les adresses ne comportant que des 0 pour le sous-réseau. La formule devient alors $2^n - 1$.

La commande pour permettre l'utilisation des sous-réseaux ne contenant que des 0 (**ip subnet-zero**) est active par défaut à partir de la version 12.0 de Cisco IOS. Cependant, il faudra prendre en compte le fait que certains systèmes tels que Sun Solaris 4.x n'intègrent pas cette fonctionnalité. La règle $2^n - 2$ reste vraie pour la portion Internet et pour la portion hôte.

Pour le VLSM la règle $2^n - 2$ doit être appliquée une seule fois sur la partie sous-réseau, c'est-à-dire sur un des découpages (peu importante lequel). Dans l'exemple précédent, la règle a été appliquée au niveau de la partie agglomération.

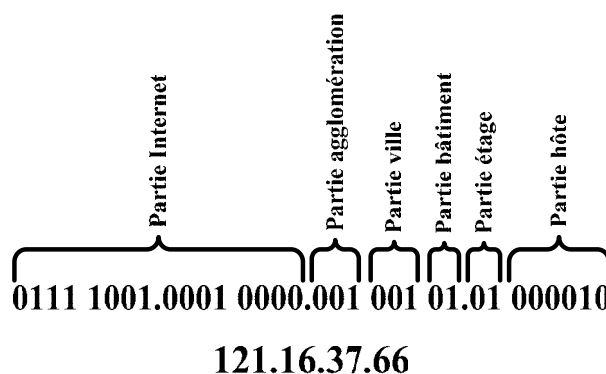
2.3.5. Allocation des adresses VLSM

Une fois le design du réseau établi et les règles VLSM prises en compte, il est très simple de déterminer les plages d'adresses IP pour chaque réseau segmenté.



Prenons l'exemple ci-dessus et voyons à quoi pourrait ressembler l'adresse d'un utilisateur au 1^{er} étage du bâtiment 1 à Plougastel dépendant de l'agglomération de Brest. La partie Classful de l'adresse comprend les deux premiers octets (définis par le masque de sous-réseau). Il faudra donc prendre une adresse de classe B. Choisissons 121.16.0.0 :

- 3 bits sont disponibles pour la partie agglomération (c'est dans cette partie que l'on a choisi d'appliquer la règle 2^n-2). Prenons 001 pour l'agglomération de Brest.
- 3 bits sont disponibles pour la partie Ville. Prenons 001 pour Plougastel.
- 2 bits sont disponibles pour la partie Bâtiment. Prenons 01 pour le bâtiment 1.
- 2 bits sont disponibles pour la partie Etage. Prenons 01 pour le premier étage.



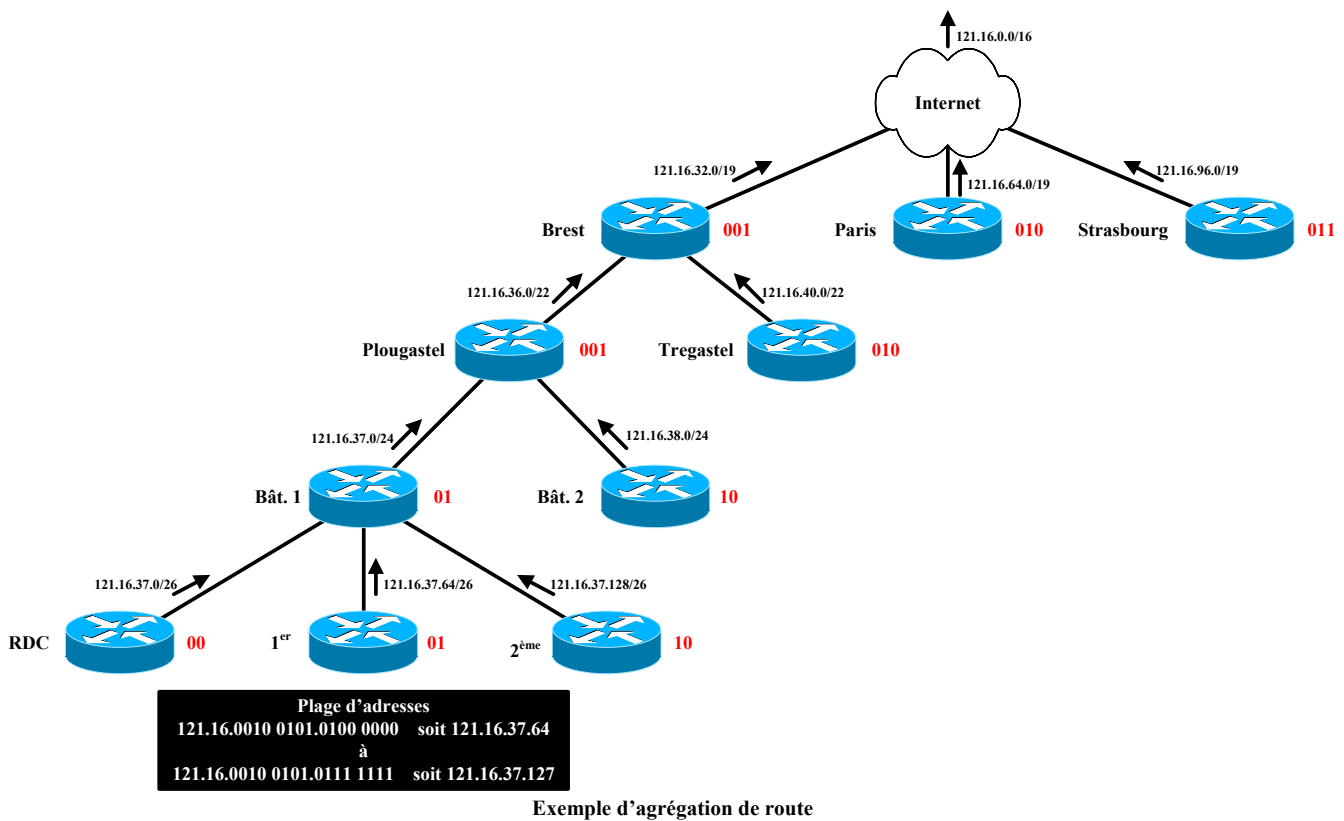
Exemple d'adresse - Masque : 255.255.255.192

2.4. Agrégat de routes

Lorsque la conception réseau est fondée sur un modèle hiérarchique, il est aisé de comprendre les avantages, en terme de réduction du trafic et de la taille des tables de routage.

L'agrégat de route permet de regrouper une multitude de réseaux en une seule adresse réseau. VLSM et CIDR utilisent les mêmes principes, à la différence que VLSM est une extension du CIDR au niveau d'une organisation. Plus on se trouve haut dans la hiérarchie du réseau, plus les tables de routage sont générales. Les masques de sous-réseaux sont de plus en plus généralistes du fait qu'ils ont agrégé plusieurs sous-réseaux en un seul.

Ces sous-réseaux agrégés sont souvent appelés sur-réseaux ou routes agrégées.



Avantages de l'agrégat de routes :

- **Réduction de la taille des tables de routage** : Les tables de routage sont plus petites, et demandent ainsi moins de bande passante et moins de temps processeur.
- **Simplification du calcul des algorithmes de routage** : Le calcul des chemins réseau est plus simple, donc plus rapide du fait des tables de routage contenant moins d'entrées.
- **Les changements topologiques du réseau sont cachés** : Les sous-réseaux les plus proches des utilisateurs sont cachés derrière leur masque de sous-réseau. Les routeurs les plus proches de la couche principale ne sont donc pas avertis des modifications au niveau des tables de routage des routeurs inférieurs puisque ces routes sont agrégées. L'avantage est donc qu'il n'y a pas de surplus de demandes en bande passante ; mais ceci implique un inconvénient puisqu'un réseau peut être injoignable mais considéré inversement par les routeurs supérieurs. Ceci se traduit donc généralement par un trafic transitant sur le backbone du réseau alors que la destination est injoignable.

2.5. Configuration de l'agrégat de routes

Les différents protocoles de routage gèrent l'agrégat de route de façons différentes. Nous verrons les différentes étapes de configuration au niveau de l'étude de chaque protocole de routage concerné.

2.5.1. Agrégat automatique

Les anciens protocoles de routage, tels que RIPv1 ou IGRP, agrègent automatiquement les adresses aux frontières Classful. Ceci est intrinsèque du fait que ces protocoles n'envoient pas le masque de sous-réseau dans les mises à jour de routage. Quand une mise à jour de routage arrive sur une interface du routeur, ce dernier regarde s'il a une interface appartenant à la même partie réseau. Si tel est le cas, le routeur applique à cette mise à jour le masque de sous-réseau configuré au niveau de cette interface. Dans le cas contraire, il applique le masque de sous-réseau par défaut, à savoir celui aux frontières Classful.

L'agrégation automatique est activée par défaut pour tous les protocoles de routage, excepté OSPF. On ne peut désactiver cette agrégation automatique que sur les protocoles Classless.

Ceci s'applique grâce à la commande suivante (mode configuration du protocole de routage) : **no auto-summary**

2.5.2. Agrégat manuel

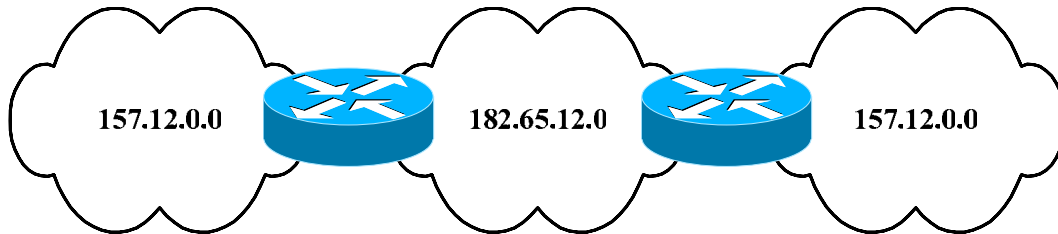
Les protocoles de routage Classless envoient le masque de sous-réseau dans leur mise à jour de routage. Ceci permet donc l'utilisation de VLSM et de la mise en place de l'agrégation de routes. Lorsqu'une mise à jour de routage arrive sur une interface du routeur, ce dernier assigne le masque au sous-réseau particulier. Lorsque le routeur cherche une entrée dans la table de routage, il se base sur l'entrée la plus proche du sous-réseau cherché (masque de sous-réseau le plus long vers le sous-réseau particulier).

Ceci implique :

- Un design hiérarchique évolutif
- L'agrégation de route
- La possibilité d'avoir des sous-réseaux discontinus

2.5.3. Sous-réseaux discontigus

Un réseau discontigu est un réseau dans lequel on retrouve des sous-réseaux contigus séparés par un réseau dont la partie Classful n'appartient pas à ces réseaux contigus. Ceci peut arriver dans le cadre d'une conception volontaire ou encore lors de la rupture de liens dans une topologie. La même adresse Internet apparaît plusieurs fois dans le réseau mais à différents endroits.



Exemple d'un réseau discontigu

Si le réseau n'utilise pas de protocole de routage Classless, le masque de sous-réseau par défaut est employé et les entrées de tables de routage ont des chemins multiples vers une même destination (partie Classful). Ceci aura pour effet de mettre en place dans la plupart des cas un partage de charge incohérent (si les routes ont un coût identique), et donc des connexions intermittentes (flapping).

Pour éviter ceci, on utilisera un protocole de routage Classless.

Si on retrouve un réseau discontigu dans l'organisation, il est primordial de désactiver l'agrégat de route (ou de ne pas le configurer), puisque cette méthode regroupe les sous-réseaux entre eux. On peut donc imaginer qu'un routeur dans le nuage 182.65.12.0 reçoive une mise à jour de routage ne comprenant qu'une seule entrée pour le réseau agrégé 157.12.0.0. Ceci est particulièrement vrai lorsque l'on utilise le protocole EIGRP qui agrège automatiquement les sous-réseaux. Lorsque l'on utilise ce protocole, l'agrégat se configure au niveau de l'interface, et on peut ainsi préciser quelles interfaces appliquent ou pas l'agrégat de sous-réseaux.

3. Routage IP

3.1. Protocoles routés et protocoles de routage

Avant de commencer à rentrer dans le vif du sujet, à savoir étudier les buts et caractéristiques générales des protocoles de routage, nous allons voir les définitions suivantes :

- **Protocole routé** : Il définit le format des paquets et fournit les informations d'adressage.
 - **Protocole routable** : L'administrateur gère l'adressage (Distinction des parties hôte et réseau de la plage d'adresse).
 - **Protocole non routable** : Ce type de protocole ne fournit pas les mécanismes nécessaires à la distinction des réseaux.
- **Protocole de routage** : Il utilise les informations fournies par un protocole routé contenues dans les paquets afin de prendre des décisions concernant la redirection de ces paquets.

La liste des protocoles routés suivante présente les protocoles les plus connus :

Nom du protocole routé	Protocole routable ?
IP	Oui
IPX	Oui
Appletalk	Oui
CLNP	Oui
NetBEUI	Non

Les buts principaux d'un protocole de routage concernent :

- La complétion d'une table de routage en fonction de critères
- Le maintien à jour de la pertinence de cette table de routage
- La communication avec les dispositifs de routage voisins pour l'échange d'informations

Il existe deux grands types de protocoles de routage pour IP :

- **Classful** (RIPv1 et IGRP) :
 - Aucune information concernant le masque de sous-réseau n'est envoyée dans les mises à jour destinées aux routeurs voisins.
 - L'agrégation de routes s'effectue automatiquement et obligatoirement aux frontières des classes d'adresses.
 - Pour les réseaux directement connectés, le masque de sous-réseau appliqué est celui configuré sur l'interface locale.
 - Le masque de sous-réseau doit être identique pour toutes les interfaces d'un même routeur appartenant à une même classe d'adresses.
 - VLSM n'est pas supporté.
- **Classless** (RIPv2, EIGRP, OSPF, Integrated IS-IS et BGP) :
 - Les masques de sous-réseau sont inclus dans les mises à jour de routage envoyées aux voisins.
 - VLSM est supporté.
 - L'agrégation de routes peut être automatiquement faite aux frontières des classes d'adresses ou configurées administrativement (CIDR).

3.2. Table de routage

La table de routage contient les informations nécessaires pour chaque réseau de destination connu par le routeur. Les décisions de routage seront toujours prises par rapport au contenu de cette table de routage, les protocoles de routage n'étant présents à ce niveau que pour la remplir. Elle contient les champs suivants :

- **Destination :**
 - Il existe par défaut une entrée spécifique par destination.
 - Ceci peut aller jusqu'à six, de manière à permettre la répartition de charge sur plusieurs liens (généralement en utilisant la technique du round-robin). Ces entrées doivent obligatoirement avoir un prochain saut différent.
 - Il ne peut exister qu'une seule entrée dans la table de routage pour une destination donnée passant par un même prochain saut.
- **Interface de sortie :**
 - C'est l'interface locale du routeur vers laquelle le paquet sera commuté.
- **Métrique :**
 - Il s'agit d'une valeur numérique, utilisée par les protocoles de routage, qui permet la sélection du meilleur chemin et qui est basée sur des critères propres à chaque protocole de routage :

Protocole de routage	Métrique
RIP	Nombre de sauts
IGRP & EIGRP	Bande passante, délai, charge, fiabilité & MTU
OSPF	Coût
IS-IS	Coût

- La métrique d'une route pour un réseau directement connecté est égale à 0.
- Plus la métrique est petite, meilleure est la route.
- **Distance administrative :**
 - Cette valeur numérique permet d'indiquer un ordre de préférence entre les différents protocoles lorsque plusieurs d'entre eux concourent pour une même entrée dans la table de routage. En effet, il est presque impossible de comparer objectivement les informations fournies par différents protocoles de routage en utilisant leurs métriques calculées avec des critères différents.
 - Plus la distance administrative est petite, plus le protocole est considéré comme prioritaire.
 - Les différentes valeurs à connaître sont :

Protocole	Distance administrative
Directement connecté	0
Statique	1
EIGRP summary route	5
External BGP	20
EIGRP	90
IGRP	100
OSPF	110
IS-IS	115
RIP	120
EGP	140
External EIGRP	170
Internal BGP	200
Réseau inconnu	255

- **Prochain saut :**
 - Cela correspond à l'adresse du prochain routeur sur le chemin pour atteindre le réseau de destination.

- **Moyen d'apprentissage :**

- Ce champ explicite la méthode d'apprentissage pour chaque entrée dans la table de routage, en nous précisant le protocole de routage qui nous a informé de cette entrée :

Code	Protocole
C	Directement connecté
S	Statique
I	IGRP
R	RIP
B	BGP
D	EIGRP
D EX	External EIGRP
O	OSPF
O IA	OSPF inter-area
O N1	OSPF NSSA external type 1
O N2	OSPF NSSA external type 2
O E1	OSPF external type 1
O E2	OSPF external type 2
i	IS-IS
i L1	IS-IS level-1
i L2	IS-IS level-2
*	Candidat par défaut

Il existe une seule table de routage par protocole routé et par routeur, quel que soit le nombre de protocoles de routage configurés sur ce même routeur. Ceci induit qu'il n'y a pas de limite sur le nombre de protocoles de routage pour un même protocole routé, sachant que les protocoles de routage concourront pour le meilleur chemin.

Les commandes suivantes permettent de visualiser et gérer la table de routage pour le protocole IP :

- **show ip route [réseau] [masque]** : Afficher la table de routage IP.
- **clear ip route {* | {réseau [masque]}}** : Supprimer une ou plusieurs entrées de la table de routage IP.
- **ip classless** : Active (actif par défaut) la prise en charge des informations ne respectant pas le découpage d'adresses en classes. Ceci permet d'activer le support des masques de sous-réseau et d'une route par défaut.

3.3. Fonctions de commutation et de routage

Il faut bien faire la distinction entre les fonctions de commutation et de routage d'un routeur. Il s'agit en réalité de deux fonctions complètement différentes et complémentaires :

- **Routage** : Prise de décision (traitement logiciel)
- **Commuation** : Application de la décision de routage (traitement matériel)

La fonction de routage a pour but d'apprendre la topologie logique du réseau, représentée dans la table de routage, et de prendre des décisions basées sur cette vue de la topologie. Ceci permet de déterminer l'interface de sortie pour les paquets entrants. Les décisions portent sur les sujets suivants :

- Le protocole routé doit être configuré sur le routeur.
- La table de routage doit posséder une entrée pour le réseau de destination (entrée spécifique ou réseau/route par défaut).
- Le réseau de destination doit être accessible.
- Le meilleur chemin doit être choisi pour atteindre le réseau de destination.
- Les chemins redondants doivent être utilisés si configurés et existants.
- L'interface de sortie doit être définie afin de placer le paquet dans la file d'attente de cette interface.

La fonction de commutation a pour intérêt principal le déplacement des données au travers du routeur. Ce traitement se fait matériellement et s'opère une fois que la décision de routage est prise. La grande caractéristique de cette fonction de commutation est la rapidité d'exécution. Les opérations suivantes sont effectuées :

- Vérification de la validité des trames
- Vérification des critères de taille des trames
- Vérification du CRC des trames
- Désencapsulation des trames entrantes et recherche de l'adresse de destination dans la mémoire tampon
- Création de l'en-tête et en-queue de trame pour les paquets sortants et transfert des trames résultantes vers la file d'attente de l'interface de sortie

La relation entre la fonction de routage et celle de commutation peut être améliorée en instruisant une mémoire tampon des décisions de routage (Route Cache). Ceci permet d'économiser considérablement les ressources processeur du routeur tout en accélérant le transfert des paquets au travers de ce dernier. Les entrées de cette mémoire tampon incluent les informations suivantes :

- Un préfixe IP
- L'interface de sortie
- L'en-tête de trame à utiliser pour les paquets à transférer

Ceci est faisable en utilisant une des techniques suivantes :

- Fast Switching
- Autonomous Switching
- Silicon Switching
- CEF (Cisco Express Forwarding)

CEF va un peu plus loin dans ce principe. En effet, chaque interface possède sa propre instance du CEF et sa propre mémoire tampon des décisions de routage, qui est appelée Forwarding Information Base (FIB).

Lorsqu'une de ces techniques est mise en œuvre, il devient alors impossible d'effectuer une répartition de charge sur plusieurs chemins pour une même destination. En effet, des en-têtes différents devraient être utilisés pour des paquets similaires, ce qui est contradictoire avec le principe de Route Cache.

3.4. Protocoles de routage à vecteur de distance et à état des liens

Les protocoles de routage peuvent être distingués par le type d'algorithme de routage qu'ils utilisent. Il en existe deux types génériques :

- Vecteur de distance
- Etat des liens

Le tableau suivant énumère les différents protocoles de routage et le type d'algorithme associé :

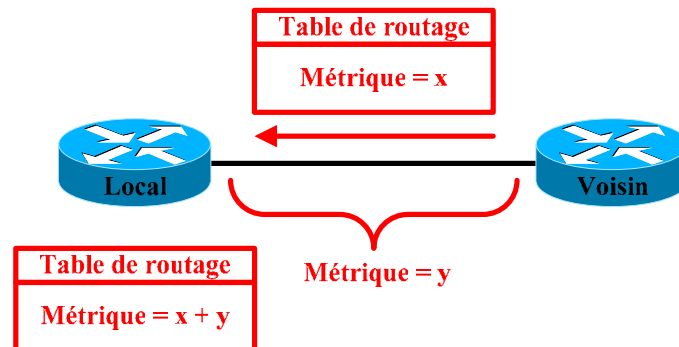
Protocole de routage	Algorithme
RIP	Vecteur de distance
IGRP	Vecteur de distance
EIGRP	Vecteur de distance évolué (Hybride)
OSPF	Etat des liens
IS-IS	Etat des liens

Ces algorithmes diffèrent principalement par leur vision de la topologie du réseau ainsi que par leur processus de transmission des informations dans les mises à jour de routage.

3.4.1. Vecteur de distance

Cet algorithme possède une vision de la topologie du réseau qui est basée sur celle de ses voisins. En effet, les mises à jour de routage envoyées par les protocoles de routage à vecteur de distance contiennent directement la table de routage du routeur émetteur. Le récepteur n'a donc pour seul travail que de récupérer ces informations, de garder les entrées pertinentes et de modifier les métriques.

La métrique locale pour une entrée dans la table de routage a pour valeur le résultat de l'addition entre la métrique incluse dans la mise à jour reçue par un voisin et de la valeur de la métrique vers ce voisin.



De plus, les mises à jour possèdent des caractéristiques précises :

- Elles sont envoyées périodiquement
- Elles contiennent directement toutes les entrées de la table de routage de l'émetteur
- Elles sont émises en broadcast (sauf exceptions telles que RIPv2 et EIGRP)

La sélection du meilleur chemin, qui sera inclus dans la table de routage, se fait en utilisant l'algorithme de Bellman Ford. Ce dernier se base sur les métriques de chaque chemin. Une exception existe pour les protocoles de routage à vecteur de distance propriétaires, tels que IGRP et EIGRP de Cisco.

3.4.2. Etat des liens

Cet algorithme exploite le principe du plus court chemin d'abord (Shortest Path First). Ce principe est basé sur l'utilisation :

- D'une table de données topologiques
- De l'algorithme de Dijkstra
- D'un arbre du plus court chemin d'abord (SPF Tree)

Les mises à jour de routage des protocoles à état des liens possèdent de grandes différences comparées à celles des protocoles à vecteur de distance :

- Elles sont uniquement envoyées lors de modifications topologiques (Triggered Updates).
- Elles contiennent des informations topologiques (Link State Advertisements).
- Elles sont incrémentielles.
- Elles sont émises en multicast sur des adresses spécifiques.

La propagation d'informations topologiques permet de ne pas baser ses décisions de routage sur une vision du réseau donnée par les voisins ainsi que d'être plus efficace au niveau de la pertinence de l'information. En effet, l'état d'un seul lien peut affecter plusieurs routes. Les ressources utilisées sont alors plus orientées processeur plutôt que bande passante sur le réseau.

Les protocoles de routage à état des liens développent des relations de voisinage avec les routeurs adjacents. Ces relations sont maintenues en permanence via l'émission réception de messages. L'intérêt principal est de connaître l'existence d'un voisin avec qui converser ainsi que son état et, par conséquent, l'état des routes passant par lui.

3.5. Système autonome - Protocoles de routage intérieurs et extérieurs

Un système autonome est, par définition, l'ensemble des dispositifs interconnectés régis par la même administration. Cela permet de délimiter la responsabilité du routage à un ensemble défini.

Ces systèmes autonomes sont identifiés par un numéro de système autonome qui est chiffré sur 16 bits. Ce numéro est unique dans le monde et permet d'identifier une organisation aux yeux du reste du monde informatique. Il est attribué par le Network Information Center (NIC). Pour les protocoles de routage imposant l'indication d'un tel numéro et se trouvant dans un réseau privé, ce numéro de système autonome peut être choisi arbitrairement dans la plage de valeurs allant de 64512 et 65535.

Cette notion de système autonome crée donc une nouvelle distinction entre les protocoles de routage :

- **Protocoles de routage intérieurs (IGP)** : Protocoles ayant pour mission principale le routage à l'intérieur d'un système autonome
- **Protocoles de routage extérieurs (EGP)** : Protocoles permettant le routage entre les systèmes autonomes

Les protocoles de routage intérieurs voient un système autonome comme un seul et unique protocole de routage. De ce point de vue, si plusieurs protocoles de routage existent dans un même système autonome, chaque protocole considérera le protocole adjacent comme externe.

Les protocoles de routage sont donc classifiés ainsi :

Classification	Protocoles
IGP	RIP, IGRP, EIGRP, OSPF et IS-IS
EGP	EGP et BGP

Typiquement, la convergence d'un réseau est restreinte au système autonome. Le temps de convergence dépend donc du protocole utilisé dans le système autonome.

3.6. Redistribution de routes

Ce principe de système autonome introduit la nécessité de pouvoir communiquer entre ces systèmes autonomes ou, plus simplement, de pouvoir échanger les informations de routage entre les IGP. Pour cela, il existe la méthode classique d'utilisation d'un protocole de routage extérieur, tel que BGP.

Il existe une autre méthode qui est plus réservée à la communication des protocoles de routage intérieurs dans un même système autonome, à savoir la redistribution de routes entre ces derniers. Le principe est extrêmement simple et consiste à injecter les routes provenant d'un protocole en tant que routes externes dans un autre protocole.

La redistribution entre certains protocoles s'effectue automatiquement entre :

- EIGRP et RTMP (Appletalk)
- EIGRP et IPX-RIP (IPX)
- IGRP et EIGRP (IP), s'ils ont le même numéro de système autonome

La redistribution de routes entre protocoles peut créer des inconvénients majeurs en cas de mauvaise configuration et/ou d'implémentation :

- **Une décision de routage mauvaise ou moins efficace** de par la disparité entre les métriques de chaque protocole redistribué. Le choix d'une route moins efficace est désigné par le choix du chemin sous-optimal (Suboptimal Path).
- **L'apparition d'une boucle de routage**, lorsque la redistribution bidirectionnelle est mise en place et que les routes originellement redistribuées dans un autre protocole de routage reviennent par ce dernier.
- **Un temps de convergence accru** à cause de la disparité des technologies mises en œuvre.

Il est néanmoins facile de remédier à ces inconvénients en configurant plus précisément la redistribution entre chaque protocole. Ceci inclut :

- Le changement de la métrique
- Le changement de la distance administrative
- L'utilisation de routes par défaut
- L'utilisation d'interfaces passives avec des routes statiques
- L'utilisation de Distribute Lists

Un autre moyen est de contrôler les mises à jour de routage. Il existe plusieurs méthodes de contrôle :

- Interfaces passives
- Routes par défaut
- Routes statiques
- Interfaces nulles
- Distribute Lists
- Route Maps

Une interface passive est une interface qui ne participe pas au processus de routage pour un protocole particulier. Cela a deux effets différents, en fonction du protocole de routage mis en place :

- **Protocole de routage à vecteur de distance** : Les mises à jour de routage ne sont pas envoyées mais elles sont écoutées.
- **Protocole de routage à état des liens** : L'établissement d'une relation de voisinage n'est pas possible sur le lien concerné. Par conséquent, aucune mise à jour de routage n'est envoyée ni reçue sur ce lien.

Une interface nulle est une interface abstraite, et qui ne mène qu'au néant. Elle n'existe qu'en tant qu'interface locale de sortie dans une route statique et elle peut avoir plusieurs utilités :

- **Alléger le traitement processeur** en remplaçant des ACLs standards par des routes statiques utilisant une interface nulle.
- **Introduire des routes** dans un autre protocole de routage, via redistribution. Ceci permet entre autres de redistribuer des routes d'un protocole de routage supportant et utilisant VLSM vers un autre protocole ne le supportant pas en agrégeant ces routes.

Les Distribute Lists et les Route Maps seront discutées plus loin dans ce chapitre.

Un point important à prendre en compte lors d'une redistribution de routes d'un protocole vers un autre est la métrique donnée aux routes injectées. La valeur peut être indiquée directement avec la commande de configuration de la redistribution ou à l'aide d'une commande séparée.

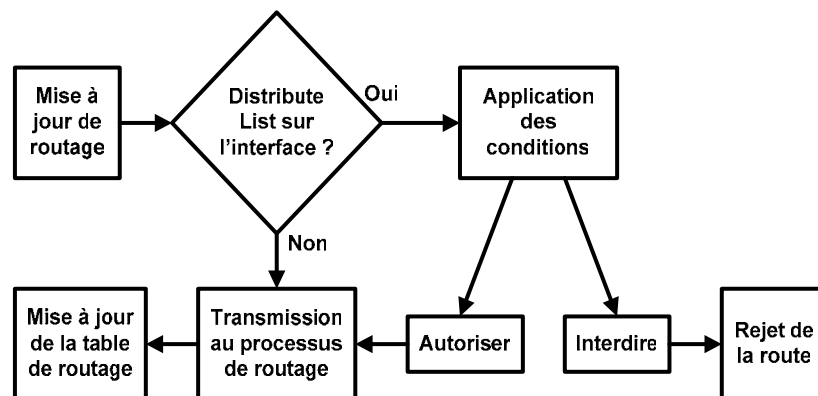
Les commandes à utiliser pour chaque action expliquée plus haut sont les suivantes :

- **redistribute {protocole} [processus] {level-1 | level-1-2 | level-2} [metric {valeur}] [metric-type {1 | 2}] [match {internal | external {1 | 2}}] [tag {valeur}] [route-map [map-tag]] [weight {valeur}]**
 - Mode de configuration du protocole de routage
 - **protocole** peut prendre les valeurs **connected**, **bgp**, **eigrp**, **egp**, **igrp**, **isis**, **iso-igrp**, **mobile**, **ospf**, **static** et **rip**. Cela correspond au protocole dont les routes seront redistribuées dans le protocole de routage courant.
 - **processus** correspond au numéro de système autonome pour les protocoles IGRP, EIGRP, EGP et BGP. Pour OSPF, il correspond à l'identifiant de processus (Process ID). Ce paramètre n'existe pas pour RIP.
 - Les options **level-1**, **level-1-2** et **level-2** n'existent que pour IS-IS et permettent de définir le type de routes qui seront redistribuées.
 - La valeur indiquée après le mot clé **metric** permet de spécifier la métrique des routes qui seront redistribuées. Pour IGRP et EIGRP, il ne faut pas indiquer une valeur mais cinq valeurs (Bande passante, délai, fiabilité, charge et MTU).
 - **metric-type** est une option disponible uniquement pour OSPF et permet de spécifier le type de route externe qui sera injectée par redistribution. La valeur par défaut est 2.
 - L'option **match** est disponible lorsque l'on redistribue le protocole OSPF dans le protocole de routage courant. Elle permet de spécifier le type de routes OSPF qui seront redistribuées.
 - Le paramètre **tag** est une valeur décimale sur 32 bits incluse dans chaque route externe. Par défaut, BGP et EGP utilisent la valeur numérique du système autonome distant, alors que zéro est utilisé comme valeur par défaut pour les autres protocoles.
 - Une Route Map peut être utilisée pour filtrer les routes qui seront redistribuées. Il faut pour cela utiliser le mot-clé **route-map** suivi du nom de la Route Map désirée. Si aucun nom n'est spécifié, alors aucune route ne pourra être redistribuée.
 - Le paramètre **weight** permet de spécifier l'attribut correspondant lors d'une redistribution dans BGP. On peut spécifier une valeur allant de 0 à 65535.
- **default-metric {BP} {délai} {fiabilité} {charge} {MTU}**
 - Mode de configuration des protocoles de routage IGRP et EIGRP
 - Permet de configurer la métrique (appelée Seed Metric) que prendront toutes les routes injectées dans IGRP ou EIGRP par redistribution.
- **default-metric {valeur}**
 - Mode de configuration du protocole de routage
 - Idem à la commande précédente, sachant qu'elle est utilisable pour les autres protocoles de routage IP.
- **distance eigrp {distance interne} {distance externe}**
 - Mode de configuration du protocole de routage EIGRP
 - Permet de modifier la valeur de la distance administrative pour les routes internes et externes.
- **distance {distance} [{adresse} {masque générique}] [n° ACL | nom] [ip]**
 - Mode de configuration du protocole de routage
 - Cette commande est pour tous les autres protocoles de routage.
 - La valeur pour la distance administrative peut aller de 10 à 255, sachant que 255 signifie que la destination est injoignable. Les valeurs allant de 0 à 9 sont réservées.
 - Il est possible d'indiquer l'adresse réseau pour laquelle on souhaite modifier la valeur de la distance administrative. Cela est possible grâce aux paramètres **adresse** et **masque générique**.
 - On peut aussi spécifier le numéro ou le nom d'une ACL standard. Ceci permet d'appliquer la distance administrative voulue uniquement aux routes correspondant aux autorisations de l'ACL.
 - L'option **ip** indique que seules les routes IP d'IS-IS seront modifiées.
- **passive-interface {type} {numéro}**
 - Mode de configuration du protocole de routage
 - Indique l'interface à rendre passive pour ce protocole de routage.

- **ip route {préfix} {masque} {adresse | interface} [distance] [tag {tag}] [permanent]**
 - Mode de configuration globale
 - Le troisième paramètre peut être soit l'adresse IP du prochain saut, soit l'interface locale de sortie.
 - Le quatrième paramètre permet d'explicitement la distance administrative pour cette route (valeur par défaut = 1).
 - Le paramètre **tag** définit cette route statique comme une valeur match pour des Route Maps.
 - Le dernier paramètre, **permanent**, force la route à ne pas être supprimée de la table de routage même si l'interface locale associée devenait non fonctionnelle.
- **ip route 0.0.0.0 0.0.0.0 {adresse | interface} [distance] [permanent]**
 - Mode de configuration globale
 - Définit une route statique par défaut sur le routeur local
- **ip default-network {préfix}**
 - Mode de configuration globale
 - Cette commande génère une route par défaut pour être envoyée dans les mises à jour de routage. Elle ne crée pas de route par défaut dans la table de routage du routeur local sauf si le réseau candidat par défaut indiqué en paramètre est un réseau directement connecté.
 - Plusieurs instances de cette commande peuvent coexister, sachant que seule la meilleure, en termes de métrique et distance administrative, sera considérée pour être la route par défaut dans la table de routage.

3.7. Distribute Lists

L'amélioration des performances dans la communication réseau porte sur plusieurs points principaux. L'un d'eux concerne la réduction des mises à jour de routage. L'un des moyens pour y parvenir est l'utilisation d'ACLs appliquées aux mises à jour de routage. Ceci est appelé des Distribute Lists.



Algorithme de traitement des Distribute Lists

Le principe de fonctionnement est très simple. Lorsqu'un paquet de mise à jour de routage est reçu ou envoyé, le routeur regarde si une Distribute List est configurée. Si elle est configurée, les routes contenues dans cette mise à jour seront filtrées par l'ACL correspondante. Finalement, celles étant autorisées par ce filtrage seront envoyées comme normalement au processus de routage afin de mettre à jour par la suite la table de routage.

Les commandes existantes pour l'implémentation de Distribute Lists sont les suivantes :

- **distribute-list {n° ACL | nom} in [{type} {numéro}]**
 - Mode de configuration du protocole de routage
 - Applique une Distribute List pour les mises à jour de routage entrantes
 - Le premier paramètre peut être soit le numéro soit le nom de l'ACL standard qui sera utilisée pour effectuer le filtrage.
 - Le paramètre optionnel permet de définir une interface spécifique pour cette Distribute List.
- **distribute-list {n° ACL | nom} out [interface | processus | n° AS]**
 - Mode de configuration du protocole de routage
 - Le premier paramètre peut être soit le numéro soit le nom de l'ACL standard qui sera utilisée pour effectuer le filtrage.
 - Le paramètre optionnel explicite l'interface, le processus de routage ou le numéro de système autonome sur lequel la Distribute List sera appliquée.
 - Pour OSPF, il n'est pas possible d'appliquer la Distribute List en fonction d'une interface.

Ces commandes ne permettent que d'appliquer la Distribute List. Il faut donc préalablement créer l'ACL standard qui sera utilisée par cette Distribute List.

3.8. Route Maps

Les Route Maps sont des suites de critères avec des actions à accomplir. Les critères sont définis par des instructions **match** alors que les actions à accomplir sont sous la forme d'instructions **set**.

Elles peuvent être utilisées afin de :

- Contrôler la redistribution
- Contrôler et modifier les informations de routage
- Définir des politiques dans la politique de routage. Dans ce cas, la Route Map sera configurée sur l'interface entrante.

Les caractéristiques des Route Maps sont les suivantes :

- Elles possèdent une liste de critères, représentés par des instructions **match**.
- Les paquets ou les routes qui correspondent aux critères se voient appliquer un traitement particulier en utilisant des instructions **set**.
- Plusieurs Route Maps peuvent exister avec le même nom. Ceci s'appelle une collection de Route Maps. Cette collection est considérée comme une unique Route Map.
- Chaque partie dans une Route Map est identifiée avec un numéro de séquence.
- A l'instar d'une ACL, la première correspondance dans les critères est utilisée. Les suivantes ne seront pas analysées.
- Une instruction **match** peut contenir plusieurs conditions. Au moins une de ces conditions doit être vérifiée afin que l'instruction **match** soit vérifiée. Il s'agit d'un OU logique.
- Une Route Map peut contenir plusieurs instructions **match**. Dans ce cas, toutes les instructions **match** devront être vérifiées. Il s'agit d'un ET logique.
- Le numéro de séquence permet de définir l'ordre dans lequel seront examinées les différentes conditions.
- Il existe une interdiction implicite à la fin d'une Route Map, comme pour les ACLs.

La commande permettant de créer une Route Map est :

- **route-map {nom} [permit | deny] [seq-n°]**
 - Mode de configuration globale
 - Permet de passer dans le mode de configuration de la Route Map
 - Le paramètre **permit** ou **deny** indique l'action principale à effectuer en cas de correspondance des critères **match** de cette Route Map.
 - **seq-n°** est le numéro de séquence pour cette Route Map dans une collection (défaut = 10).

Il faut maintenant voir les différents critères et actions (tous depuis le mode de configuration de la Route Map) :

- **match ip address [n° ACL | nom] [...]**
 - Utilisation d'une ou plusieurs ACLs standard ou étendues pour l'examen des paquets entrants
- **match length {min} {max}**
 - Définition d'un critère de longueur du paquet. Les valeurs **min** et **max** sont inclusives.
- **set default interface {type} {numéro} [...]**
 - Cette instruction fournit une liste d'interfaces de sortie par défaut si aucune entrée explicite n'existe pour le réseau de destination dans la table de routage.
 - Si plusieurs interfaces sont indiquées, alors la première interface fonctionnelle sera utilisée.
- **set interface {type} {numéro} [...]**
 - Cette instruction précise l'interface de sortie pour les paquets correspondants si une entrée existe pour le réseau de destination dans la table de routage.
 - Si plusieurs interfaces sont indiquées, alors la première interface fonctionnelle sera utilisée.
- **set ip default next-hop {adresse} [...]**
 - Cette instruction fournit une liste de voisins par défaut si aucune entrée explicite n'existe pour le réseau de destination dans la table de routage.
 - Si plusieurs voisins sont indiqués, alors le premier voisin disponible sera utilisé.
- **set ip next-hop {adresse} [...]**
 - Cette instruction fournit une liste de voisins à utiliser comme prochain saut si une entrée existe pour le réseau de destination dans la table de routage
 - Si plusieurs voisins sont indiqués, alors le premier voisin disponible sera utilisé.
- **set ip precedence {priorité}**
 - Ceci permet de modifier la valeur des bits de priorité dans le champ "Type de Service" de l'en-tête IP des paquets correspondant.
- **set ip tos {TOS}**
 - Cette instruction indique la valeur du type de service IP dans le champ "Type de Service" de l'en-tête IP.

Il ne reste plus qu'à appliquer la Route Map fraîchement créée :

- **ip policy route-map {nom}**
 - Mode de configuration d'interface
 - Applique une Route Map sur cette interface
- **ip route-cache policy**
 - Mode de configuration d'interface
 - Active le Fast Switching pour le routage basé sur les politiques sur cette interface
 - Désactivé par défaut
 - Nécessite qu'une Route Map soit appliquée sur cette interface pour fonctionner
 - Les actions **set ip default** et **set interface** ne pas supportées.

La commande **show route-map [nom]** permet de visualiser les critères spécifiques aux Route Maps.

4. Protocole EIGRP

4.1. Caractéristiques

EIGRP (Enhanced IGRP), protocole propriétaire Cisco, est une version améliorée d'IGRP qui utilise la même technologie à vecteur de distance. Les améliorations portent principalement sur :

- Les propriétés de convergence
- L'efficacité des opérations du protocole

Les changements apportés correspondent à beaucoup des caractéristiques des protocoles de routage à état des liens, et ont pour buts de faciliter l'évolutivité et d'accélérer le temps de convergence des réseaux. De ce fait, il est référencé dans la catégorie des protocoles de routage hybride, ou, d'après Cisco, à vecteur de distance évolué.

Les caractéristiques principales d'EIGRP sont :

- Protocole de routage Classless, avec support du VLSM
- Algorithme DUAL
- Mises à jour incrémentales, avec adressage multicast, et de façon fiable (via RTP)
- Utilisation de la bande passante réduite par rapport à IGRP
- Utilisation d'une métrique composite
- Découverte de voisins
- Principe de successeur, avec de multiples FS
- Agrégation de routes manuelle
- Etat des routes (Active et Passive)
- Partage de charge entre chemins n'ayant pas les mêmes métriques
- Compatibilité avec IGRP

Cisco identifie quatre composants principaux pour EIGRP :

- Modules dépendants du protocole routé
- RTP
- Découverte et restauration de voisins
- DUAL

Une des grandes spécificités d'EIGRP par rapport aux autres protocoles de routage est donc, d'après son premier composant principal, son support de plusieurs protocoles routés :

- IP
- IPX
- AppleTalk

Pour chaque protocole routé utilisé, EIGRP maintient 3 tables distinctes :

- Table de voisinage (Neighbor Table)
- Table de topologie (Topology Table)
- Table de routage (Route Table)

4.2. Termes et définitions

EIGRP utilise beaucoup de termes génériques et spécifiques que nous détaillons et définissons ci-dessous :

- **Neighbor (voisin)** : Routeur voisin directement connecté qui utilise aussi EIGRP.
- **Neighbor Table (table de voisinage)** : Table contenant une liste de tous les voisins. Cette table est élaborée en fonction des informations contenues dans les Hello reçus par les voisins.
- **Route Table (table de routage)** : Table de routage pour un protocole routé précis.
- **Topology Table (table de topologie)** : Table contenant tous les réseaux appris par les voisins. Cette table sert à remplir la table de routage en fonction de certains critères.
- **Hello** : Message utilisé pour découvrir les voisins et les maintenir dans la table de voisinage.
- **Update** : Paquet du protocole Hello contenant les informations sur les changements du réseau.
- **Query** : Paquet du protocole Hello demandant aux voisins l'existence d'un FS.
- **Reply** : Paquet du protocole Hello répondant à un paquet Query.
- **ACK (accusé de réception)** : Paquet du protocole Hello accusant réception des autres messages du protocole Hello. Le fenêtrage de RTP est fixé à 1. Ceci signifie que chaque paquet Update doit être suivi d'un ACK.
- **Holdtime** : Valeur incluse dans les paquets Hello indiquant le temps qu'un routeur attend un signe d'un voisin avant de le considérer comme indisponible. Sa valeur est 3 fois celle de l'intervalle de transmission des messages Hello. Passé ce délai, le voisin sera considéré comme mort.
- **Smooth Round Trip Time (SRTT)** : Temps en millisecondes (ms) nécessaire à un paquet d'être envoyé à un voisin puis à une réponse d'être reçue. Il sert à calculer la valeur du RTO.
- **Retransmission Timeout (RTO)** : Temps d'attente en millisecondes (ms) pour un paquet ACK avant retransmission du paquet d'origine. Sa valeur est calculée en fonction de la valeur du SRTT.
- **Reliable Transport Protocol (RTP)** : Condition de délivrance d'un paquet par séquence avec garantie.
- **Diffusing Update ALgorithm (DUAL)** : Algorithme appliqué sur la table de topologie pour converger le réseau.
- **Advertised Distance (AD)** : Métrique diffusée par un voisin dans sa mise à jour de routage. Elle correspond à la métrique depuis ce voisin, connu localement comme le prochain saut.
- **Reported Distance (RD)** : Autre nom pour l'AD.
- **Feasible Distance (FD)** : Plus petite métrique pour une destination donnée. C'est la métrique pour la route actuellement dans la table de routage.
- **Feasible Condition (FC)** : Condition vérifiée quand un voisin informe une AD plus petite que la FD du routeur local pour une même destination.
- **Feasible Successor (FS)** : Voisin vérifiant la FC. Il est potentiellement éligible en tant que successeur.
- **Successor** : Voisin utilisé comme prochain saut pour une destination donnée. C'est le FS ayant la plus petite métrique.
- **Active** : Etat d'une route lorsqu'une modification de réseau apparaît, et qu'il n'y a pas de FS dans la table de topologie. Le routeur interroge alors ses voisins pour connaître de plausibles routes alternatives.
- **Passive** : Etat normal pour une route opérationnelle dans la table de routage.
- **Stuck In Active (SIA) (aussi appelé Query Scoping)** : Etat d'un routeur lorsqu'une route reste active après dépassement d'un certain temps.

4.3. Métriques

Les métriques sont très similaires à celles employées par IGRP. La grande différence est que la valeur métrique est maintenant un nombre sur 32 bits. Les décisions prises peuvent donc être plus fines ou détaillées.

Il peut y avoir jusqu'à 6 routes pour une même destination dans la table de routage, et que ces routes peuvent être de 3 types :

- **Internal** : Route interne à l'AS
- **Summary** : Routes internes mises sous la forme d'un unique agrégat de routes
- **External** : Route externe à l'AS qui a été redistribuée dans l'AS EIGRP (inclus aussi les routes statiques redistribuées)

La formule pour le calcul d'une métrique EIGRP est la suivante :

$$\text{Métrique} = (K1 \times \text{Bandwidth} + K2 \times \text{Bandwidth} \div (256 - \text{Load}) + K3 \times \text{Delay}) + K5 \div (\text{Reliability} + K4)$$

Les différents paramètres de cette formule sont les suivants :

- **K1** : Coefficient rattaché à la bande passante (valeur par défaut = 1)
- **K2** : Coefficient rattaché à la charge (valeur par défaut = 0)
- **K3** : Coefficient rattaché au délai (valeur par défaut = 1)
- **K4** : Coefficient rattaché à la fiabilité (valeur par défaut = 0)
- **K5** : Coefficient rattaché au MTU (valeur par défaut = 0)
- **Bandwidth** : Valeur correspondant à la plus petite bande passante de liaison entre les hôtes source et destination. Cette valeur est calculée avec la formule $10^7 \div \text{BP} \times 256$, avec BP la bande passante exprimée en Kbps.
- **Load** : Charge sur la liaison. C'est un pourcentage binaire dont la valeur peut aller de 0 à 255.
- **Delay** : Délai de transmission sur le chemin exprimé en microsecondes (μs). C'est la somme des délais de toutes les liaisons entre les hôtes source et destination. Cette valeur est calculée via la formule $\Sigma_{\text{délais}} \times 256$.
- **Reliability** : Fiabilité de la liaison. C'est aussi un pourcentage binaire dont la valeur peut aller de 0 à 255 et qui est déterminée par le ratio entre le nombre de paquets corrects et le nombre de paquets transmis sur le média.

Ainsi, avec les valeurs par défaut, on arrive à la formule simplifiée suivante :

$$\text{Métrique} = \text{Bandwidth} + \text{Delay}$$

$$\text{Métrique} = (10^7 \div \text{BP} + \Sigma_{\text{délais}}) \times 256$$

On peut donc remarquer que, avec les paramètres par défaut, une métrique d'EIGRP est 256 fois plus grande qu'une métrique d'IGRP pour une même destination.

4.4. Protocole Hello

Le protocole Hello permet l'échange des informations de routage entre les routeurs utilisant le protocole EIGRP ainsi que la découverte dynamique des voisins. Certains messages utilisent RTP afin d'assurer la bonne réception des informations.

Les paquets du protocole Hello utilisant le multicast se servent de l'adresse 224.0.0.10 pour leur transmission.

Plusieurs types de messages, ou plus précisément paquets, existent et se distinguent de part leur utilité :

- **Hello**
 - Emis périodiquement
 - Non orienté connexion
 - Toutes les 5 secondes sur les liaisons LAN
 - Toutes les 60 secondes sur les liaisons WAN
- **Update**
 - Contient les informations des différents réseaux connus par un routeur EIGRP. Ces informations sont à destination de ces voisins, afin qu'ils puissent compléter leur table de topologie.
 - Orienté connexion avec RTP
 - S'il s'agit d'un nouveau voisin, alors le ou les paquets Update envoyés vers ce voisin sont en unicast. Dans les autres cas, le paquet Update est envoyé en multicast.
- **Query**
 - Requête vers un voisin en vue d'obtenir des informations sur les différents réseaux connus par ce dernier. Celui-ci répondra, via un ou plusieurs paquets Reply.
 - Envoyé lorsqu'une ou plusieurs destinations passent à l'état Active
 - Orienté connexion avec RTP
 - Ce type de paquet est toujours envoyé en multicast.
 - Ce type de paquet est généralement envoyé afin d'enquêter sur un réseau suspect (plus accessible, changement d'états et/ou de chemin, etc.).
- **Reply**
 - Identique à un paquet Update sauf que celui-ci est envoyé uniquement en réponse à un paquet Query.
 - Orienté connexion avec RTP
 - Ce paquet est un unicast vers le voisin ayant émis le paquet Query.
- **ACK**
 - Accusé de réception pour les paquets envoyés orientés connexion
 - Envoyé sous la forme d'unicast
 - C'est un paquet Hello sans données qui contient un numéro d'accusé de réception différent de 0.
 - Le fenêtrage a une valeur par défaut de 1. Ceci implique donc que chaque paquet Update, Query et Reply devront être suivi de ce paquet ACK de chaque voisin afin d'en assurer la remise à ces derniers. Le cas échéant, le paquet Update, Query ou Reply envoyé précédemment sera réémis en unicast.
 - Après 16 essais de retransmissions unicast, le routeur marquera le voisin incriminé comme mort.

La capacité à envoyer des retransmissions unicast diminue le temps qu'il faut pour construire les différentes tables, car tous les voisins n'ont pas à traiter et accuser réception de chaque retransmission.

4.4.1. Neighbor Table

Un routeur est considéré comme voisin si :

- Un paquet Hello ou ACK est reçu de ce voisin.
- Le numéro d'AS est identique pour les deux routeurs.
- Les paramètres de métrique sont les mêmes sur les deux routeurs.

La réception en continu des paquets Hello en provenance des voisins permet de maintenir à jour la table de voisinage, sachant que cette table contient les champs suivants :

- **Adresse** : Adresse de couche 3 du voisin
- **Interface** : Interface locale par laquelle le paquet Hello de ce voisin a été reçue
- **Holdtime** : Temps d'attente d'un signe de vie du voisin avant de le considérer comme mort
- **Uptime** : Temps écoulé depuis la découverte de ce voisin
- **Nombre de paquets en file d'attente (Q Count)** : Permet la visualisation d'une possible congestion vers ce voisin
- **Numéro de séquence** : Numéro de séquence pour les paquets (Utilisant RTP) entrants et sortants. EIGRP garde donc en mémoire deux numéros de séquence différents.
- **SRTT** : Temps nécessaire à un paquet d'être envoyé au voisin puis une réponse reçue de ce dernier
- **RTO** : Temps d'attente d'un ACK avant réémission

4.4.2. Topology Table

Cette table permet de garder en mémoire tous les réseaux accessibles par les différents voisins (y compris les dupliqués). Elle est complétée grâce aux paquets Update ou Reply (en réponse à un paquet Query) reçus des voisins et enregistre les paquets qui ont été envoyés par le routeur à ses voisins.

L'avantage de posséder la table de routage de tous les voisins dans cette table est la diminution de la surcharge réseau ainsi que des calculs. Ceci permet donc une convergence très rapide.

Cette table permet de gérer la sélection des routes à ajouter dans la table de routage parmi toutes celles disponibles en faisant appel à l'algorithme DUAL.

Elle contient les informations suivantes :

- Etat de la route (Active ou Passive)
- Qu'un paquet Update a été envoyé aux voisins
- Qu'un paquet Query a été envoyé aux voisins. Si ce champ est positif, alors au moins une route doit être marquée comme étant à l'état Active.
- Si un paquet Query a été envoyé, un autre champ indiquera si un paquet Reply a été reçu des voisins
- Qu'un paquet Reply a été envoyé en réponse à un paquet Query reçu d'un voisin
- Les réseaux distants
- Le masque (ou préfix) pour ces réseaux
- La métrique vers chaque réseau (FD)
- La métrique pour chaque réseau avertie par les voisins (AD)
- Le prochain saut pour chaque réseau
- L'interface locale par laquelle sortir pour atteindre ce prochain saut
- Les successeurs, à savoir le chemin jusqu'à la destination, exprimé en sauts

Les métriques incluses dans la table de topologie sont celles indiquées dans les paquets reçus par les voisins (AD). Cela signifie que c'est la table de routage qui calculera la métrique totale vers la destination.

Elle est mise à jour car le routeur obtient ou perd la connectivité directe avec un voisin ou car un changement topologique a été détecté grâce à la communication réseau d'EIGRP. Il existe trois raisons menant à la recalculation de cette table de topologie :

- **Un nouveau réseau est disponible :**
 - Un paquet Update avertit de l'existence d'un nouveau réseau.
 - Une interface locale devient fonctionnelle pour un protocole de couche 3 supporté par EIGRP, et ce dernier est configuré avec les commandes de réseaux appropriées.
- **Le routeur change le successeur** dans la table de topologie ainsi que dans la table de routage :
 - Un paquet Reply ou Query est reçu, modifiant ainsi une ou plusieurs entrées dans la table de topologie.
 - Il y a modification du coût pour une interface locale via configuration.
- **Un réseau devient inaccessible :**
 - Un paquet Update, Query ou Reply reçu informe la table de topologie qu'un réseau est inaccessible.
 - Aucun paquet Hello n'est reçu d'un voisin menant à ce réseau avant expiration du Holdtime.
 - Le réseau est directement connecté et l'interface du routeur perd le signal de porteuse.

4.5. DUAL

Cet algorithme a pour buts de maintenir la table de topologie à jour et de (re)créer la table de routage.

La mise à jour de la table de routage est effectuée différemment en fonction de l'état du ou des réseaux traités :

- **Passive** : Il y a une recherche dans la table de topologie d'une route acceptable pour remplacer l'ancienne présente dans la table de routage :
 - Toutes les entrées pour une même destination sont examinées afin de trouver tous les FS (ceux qui vérifient la FC, à savoir que leur AD doit être inférieure à la FD indiquée dans l'ancienne version de la table de routage).
 - Après examen, il existe au moins un FS.
 - Le FS proposant la plus petite AD sera alors choisi comme successeur à l'entrée non valide de l'ancienne table de routage.
- **Active** : Il n'y a pas de routes acceptables dans la table de topologie pour remplacer l'ancienne présente dans la table de routage. Le routeur interroge alors ses voisins via un paquet Query afin d'obtenir des informations sur des chemins possibles de remplacement :
 - Toutes les entrées pour une même destination sont examinées afin de trouver tous les FS (ceux qui vérifient la FC, à savoir que leur AD doit être inférieure à la FD indiquée dans l'ancienne version de la table de routage).
 - Après examen, il n'existe aucun FS. Le routeur passe en mode actif et envoie des paquets Query à ses voisins.
 - Si un ou plusieurs voisins répondent en indiquant une ou plusieurs nouvelles routes vérifiant la FC ($AD > FD$), alors les voisins menant à ces routes deviennent des FS.
 - Le FS proposant la plus petite AD sera alors choisi comme successeur à l'entrée non valide de l'ancienne table de routage.

DUAL offre plusieurs fonctionnalités :

- Choix d'un successeur
- Ajout d'un réseau dans la Topology Table
- Suppression d'une route ou routeur de la Topology Table
- Recherche d'un chemin alternatif vers un réseau distant

4.5.1. Choix d'un successeur

Tout commence par l'élection de FS, à savoir un ou plusieurs voisins pouvant être les prochains sauts pour une même destination. Ceci se fait en comparant leur AD pour cette destination et en appliquant la FC.

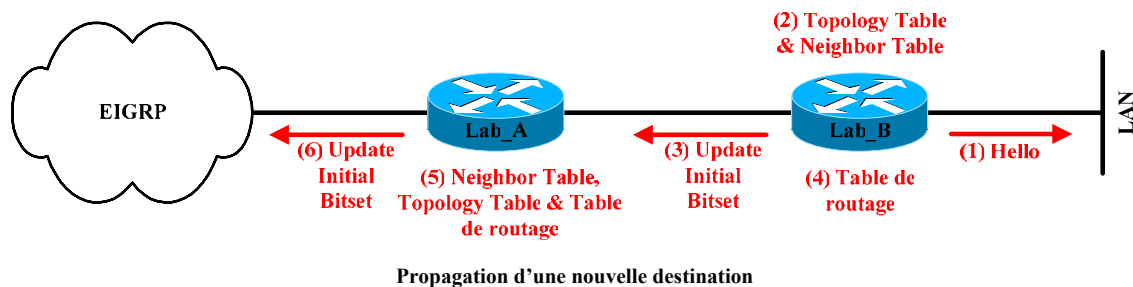
Cette FC identifie un FS si sa AD est inférieure à la FD. En d'autres termes, cela signifie qu'un voisin devient un successeur plausible si la métrique qu'il nous a communiquée (la métrique pour cette destination dans la table de routage de ce voisin) est strictement inférieure à la métrique pour cette destination de la table de routage du routeur local.

Ce principe de FC est la clé fondamentale d'EIGRP pour garder une table dénuée de boucle. En effet, si une route contient une boucle, alors l'AD sera plus grande que la FD.

Le successeur est le FS qui présente la plus petite AD.

A ce moment, DUAL n'a plus qu'à remplacer l'ancienne entrée dans la table de routage par celle correspondant au successeur pour cette destination.

4.5.2. Ajout d'un réseau dans la Topology Table

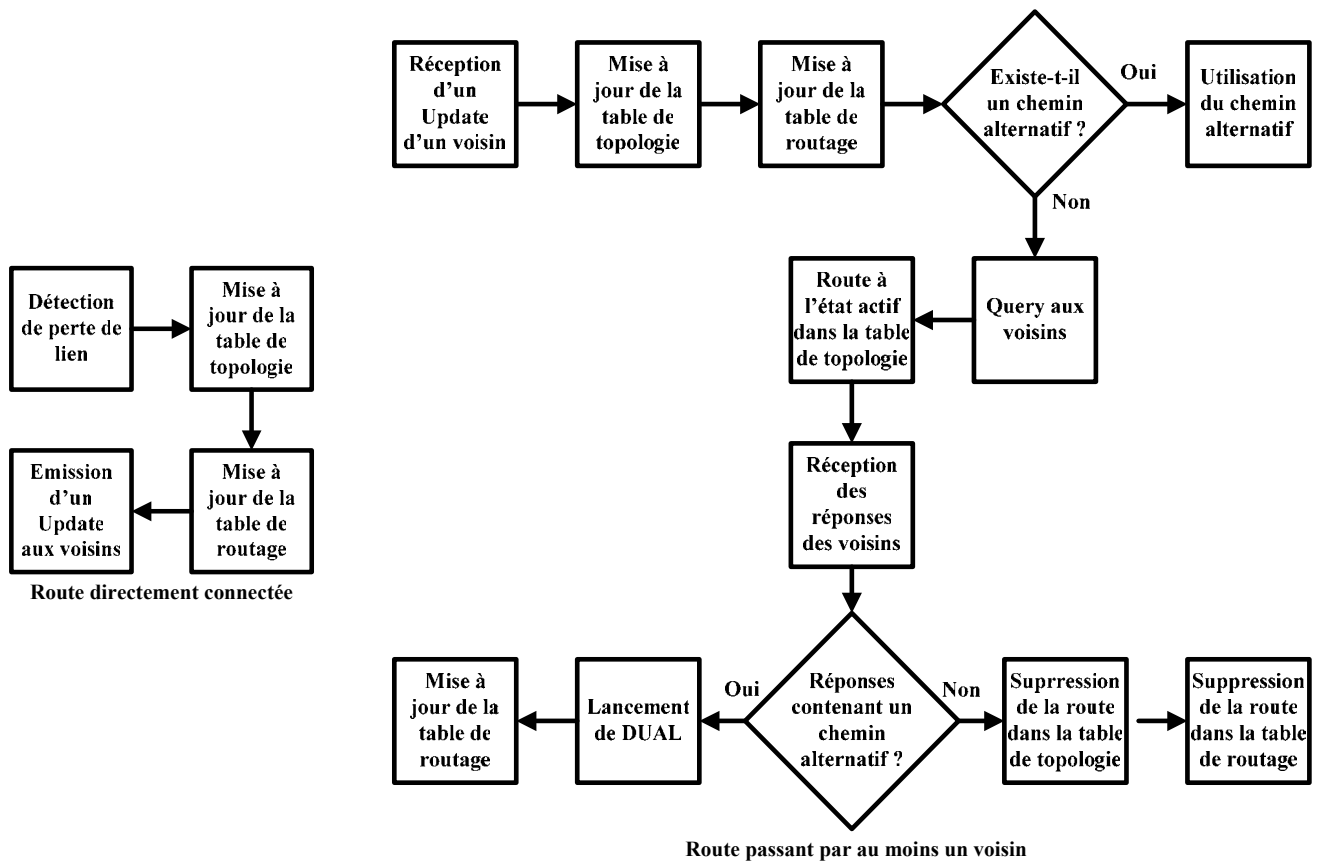


L'ajout d'une entrée dans la table de topologie puis, par succession, dans la table de voisinage et la table de routage se fait dans un ordre précis :

- **Etape (1) :** Dès qu'un routeur voit un nouveau réseau directement connecté de part la montée de son interface correspondante, il commence à envoyer des paquets Hello vers cette interface.
- **Etape (2) :** Si un routeur EIGRP se trouve sur ce nouveau réseau, alors le routeur local ajoutera ce nouveau voisin dans sa table de voisinage, sinon, aucune nouvelle entrée sera ajoutée. Par contre, quelles que soient les modifications faites dans la table de voisinage, une nouvelle entrée sera ajoutée dans la table de topologie, car un nouveau réseau vient d'apparaître.
- **Etape (3) :** Vu qu'il y a eu un changement topologique, le routeur local est obligé d'émettre un paquet Update vers tous ses voisins, les informant du nouveau réseau.
- **Etape (4) :** Le nouveau réseau est déjà ajouté à la table de topologie. Maintenant, il faut mettre à jour la table de routage. Le réseau ainsi ajouté sera marqué comme passif, car il est fonctionnel.
- **Etape (5) :** Les voisins du routeur local reçoivent le paquet Update. Ils mettent à jour le numéro de séquence dans leur table de voisinage et ajoutent le nouveau réseau dans leur table de topologie. Ils calculent la FD pour cette nouvelle destination et l'ajoutent à leur table de routage.
- **Etape (6) :** Ces voisins émettent à leur tour un paquet Update à tous leurs voisins respectifs, excepté celui par lequel est arrivé le changement topologique. Cela correspond à la mise en application du Split Horizon.

Initial Bitset : Ceci est présent dans l'en-tête EIGRP afin d'indiquer que les routes contenues dans le paquet Update sont de nouvelles routes, et non juste des modifications de routes existantes.

4.5.3. Suppression d'une route ou routeur de la Topology Table

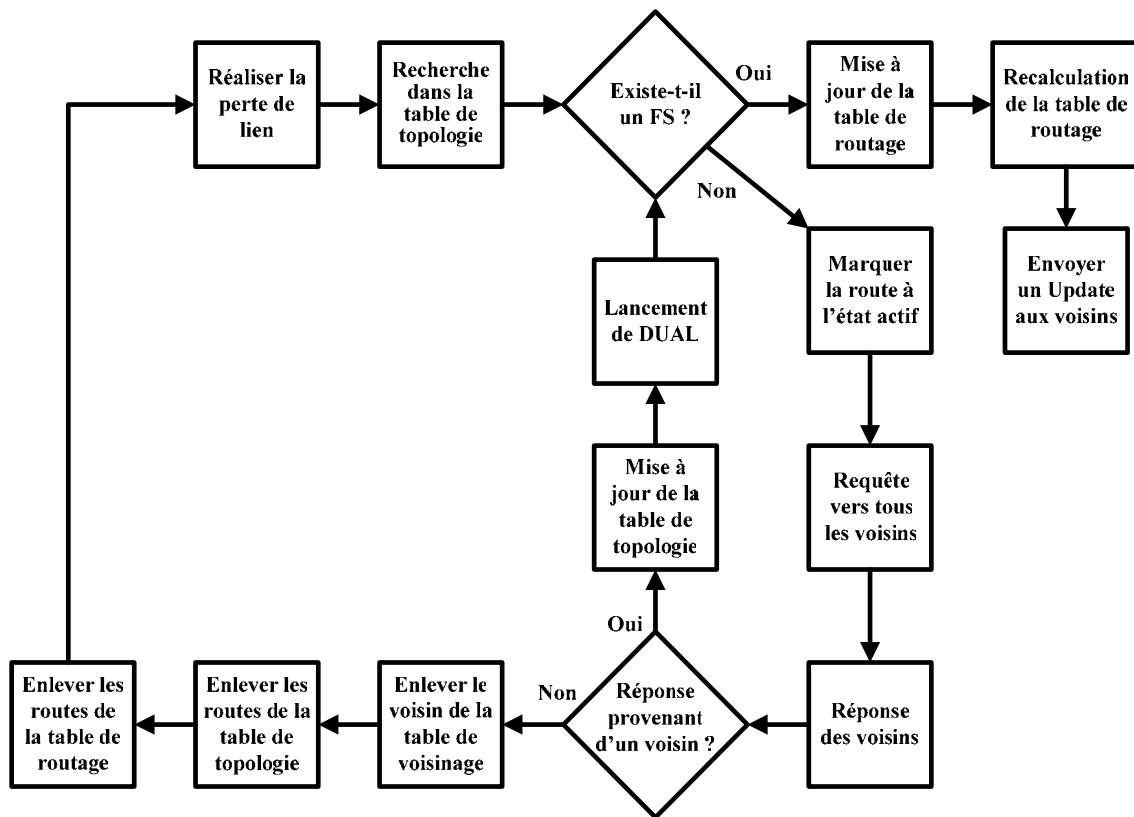


Le processus de suppression d'une route est différent en fonction de l'emplacement de cette route :

- **Route directement connectée au routeur :**
 - Détection de la perte de lien avec la route directement connectée via la non fonctionnalité de l'interface locale
 - Mise à jour des tables de topologie et de routage
 - Emission d'un paquet Update aux voisins
- **Route passant par au moins un voisin :**
 - Réception d'un paquet Update provenant d'un voisin informant la disparition d'une route
 - Mise à jour de la table de topologie ainsi que de la table de routage
 - Examen de la table de topologie pour trouver une hypothétique alternative :
 - S'il en existe une, alors elle est utilisée.
 - Sinon le routeur effectuera une requête à ses voisins via un paquet Query, et la route passera à l'état actif dans la table de topologie.
 - Lorsque toutes les réponses sont arrivées, les tables de voisinage et de topologie sont mises à jour :
 - Si un chemin alternatif est trouvé dans les réponses, alors DUAL est lancé afin de déterminer le meilleur chemin qui sera alors placé dans la table de routage.
 - Sinon les entrées pour cette destination inaccessible sont supprimées des tables de topologie et de routage.

Le processus de suppression d'un routeur est inclus dans l'algorithme de recherche d'un chemin alternatif vers un réseau distant.

4.5.4. Recherche d'un chemin alternatif vers un réseau distant



Algorithme de recherche d'un chemin alternatif vers un réseau distant

Cette recherche d'un chemin alternatif utilise l'algorithme ci-dessus :

- Réaliser la perte de lien pour une route et/ou un routeur via un paquet Update
- Chercher dans la table de topologie l'existence de FS :
 - Si un ou plusieurs FS existent, alors le meilleur est choisi comme successeur pour cette route. La table de routage est mise à jour avec les nouvelles informations et un paquet Update est envoyé aux voisins pour leur signaler les changements topologiques.
 - Sinon, la route passe à l'état actif et un paquet Query est envoyé aux voisins.
- Des paquets Reply sont reçus des voisins :
 - La table de topologie est mise à jour avec les nouvelles informations.
 - DUAL est lancé afin de trouver des FS.
 - L'algorithme repart à la deuxième étape, à savoir la recherche d'un FS dans la table de topologie.
- Aucune réponse n'est reçue d'un voisin :
 - Ce voisin est enlevé de la table de voisinage.
 - Les routes passant par ce voisin sont elles aussi enlevées successivement de la table de topologie puis de la table de routage.
 - Il y a donc perte de lien pour une ou plusieurs routes. L'algorithme reprend donc au début.

4.6. Fonctionnement avec IPX

Les avantages de l'utilisation d'IPX-EIGRP sur un réseau IPX sont les suivants :

- **Redistribution automatique** bidirectionnelle entre IPX-RIP et IPX-EIGRP
- **Taille du réseau augmentée** car IPX-EIGRP n'a pas de limitation liée au nombre de sauts maximum comme pour IPX-RIP (15 sauts)
- **Mises à jour SAP incrémentielles** lorsqu'un voisin existe sur une interface donnée. Ceci est automatiquement appliqué sur une interface série ou manuellement configuré sur une interface LAN. Si aucun voisin n'est détecté, les mises à jour SAP seront alors classiquement envoyées périodiquement.

La redistribution entre IPX-RIP et IPX-EIGRP se fait via les règles suivantes :

- Les routes IPX-EIGRP sont toujours préférées aux routes IPX-RIP, sauf quand le nombre de sauts inclut dans une route IPX-EIGRP externe est supérieur au nombre de sauts indiqué par IPX-RIP.
- Les routes IPX-RIP redistribuées sont des routes IPX-EIGRP externes. Le nombre de sauts de ces routes externes correspond au nombre de sauts de la route IPX-RIP lors de la redistribution.

4.7. Commandes

Les commandes de configuration d'EIGRP sont les suivantes :

- **router eigrp {n° AS}**
 - Mode de configuration globale
 - Active l'algorithme du protocole de routage pour IP.
 - Permet de passer en mode de configuration de ce protocole de routage.
- **network {réseau} [masque générique]**
 - Mode de configuration du protocole de routage
 - Spécifie la ou les interfaces interagissant avec ce protocole de routage. Une interface émettra et recevra donc des mises à jour de routage EIGRP si leur adresse IP fait partie du réseau indiqué en paramètre.
 - Inclut les informations concernant ces réseaux dans les mises à jour de routage transmises.
 - Le réseau indiqué en paramètre doit obligatoirement être directement connecté au routeur, mais il peut englober plusieurs sous-réseaux à la fois (via CIDR) en l'associant à un masque générique.
- **[no] auto-summary**
 - Mode de configuration du protocole de routage
 - Permet d'activer (par défaut) ou de désactiver l'agrégation de routes automatique aux frontières Classful.
- **ip summary-address eigrp {n° AS} {réseau} {masque}**
 - Mode de configuration d'interface
 - Permet de configurer manuellement un agrégat de routes à une frontière Classless.
 - Pour que l'effet de cette commande fonctionne, il faut obligatoirement que l'agrégation de routes automatique soit désactivée (commande **no auto-summary**).
- **variance {multiplicateur}**
 - Mode de configuration du protocole de routage
 - Indique la variance que peut avoir au maximum les routes qui seront incluses dans la table de routage à des fins de partage de charge.
 - Le multiplicateur est un entier pouvant aller de 1 (valeur par défaut) à 128.
- **maximum-paths {nombre}**
 - Mode de configuration du protocole de routage
 - Indique le nombre, allant de 1 (par défaut) à 6, de routes à métrique égale (à plus ou moins la variance) pouvant être mises au maximum dans la table de routage pour une même destination à des fins de partage de charge.

- **bandwidth {BP}**
 - Mode de configuration d'interface
 - Informe les protocoles de routage utilisant la bande passante pour le calcul des métriques de la véritable bande passante de la liaison.
 - La bande passante d'une liaison n'est pas détectée, et a une valeur par défaut de 1544 Kbps (T1) pour les interfaces série haut débit.
 - Le paramètre **BP** est exprimé en Kbps.
- **ip bandwidth-percent eigrp {n° AS} {pourcentage}**
 - Mode de configuration d'interface
 - Spécifie le pourcentage (défaut = 50%) de bande passante qu'EIGRP peut utiliser au maximum sur l'interface spécifiée.
- **timers active-time {disable | temps}**
 - Mode de configuration du protocole de routage
 - Temps d'attente à l'état actif pour une route avant que le routeur ne passe à l'état SIA.
- **passive-interface {type} {numéro}**
 - Mode de configuration du protocole de routage
 - Empêche l'émission et la réception de mises à jour de routage en empêchant la formation d'une relation de voisinage sur l'interface spécifiée.
- **metric weights {TOS} {K1} {K2} {K3} {K4} {K5}**
 - Mode de configuration du protocole de routage
 - Modifie des coefficients entrants en jeu dans le calcul des métriques d'EIGRP.
 - La valeur de **TOS** doit toujours être de 0.
- **ipx routing**
 - Mode de configuration globale
 - Active le routage IPX ainsi que le protocole de routage IPX-RIP.
- **ipx router eigrp {n° AS}**
 - Mode de configuration globale
 - Active l'algorithme de routage pour IPX.
 - Permet de passer dans le mode de configuration du protocole de routage.
- **ipx sap-incremental eigrp {n° AS}**
 - Mode de configuration d'interface
 - Permet d'envoyer des mises à jour incrémentales SAP uniquement lorsqu'un changement a lieu dans la table SAP.
 - Uniquement pour les interfaces LAN, car automatiquement effectué sur les interfaces série.

Pour la visualisation de l'état du protocole EIGRP, nous avons à notre disposition les commandes suivantes :

- **show ip route [eigrp [n° AS]]**
 - Visualise uniquement les routes EIGRP de la table de routage.
- **show ip eigrp neighbors [{type} {numéro} [n° AS]] [detail]**
 - Fournit toutes les informations sur les voisins, l'état de la relation de voisinage ainsi que les interfaces et adresses par lesquels ils communiquent.
- **show ip eigrp topology [all | n° AS | [IP] masque]**
 - Affiche les informations concernant la table de topologie. Il est possible d'afficher les informations pour les destinations connues en fonction du paramètre optionnel (**all** affiche toutes les routes ainsi que tous les chemins alternatifs).
- **show ip eigrp traffic [n° AS]**
 - Donne les informations regroupées sur le trafic total envoyé depuis et vers le processus EIGRP.
- **show ip eigrp interfaces [n° AS] [detail]**
 - Informations relatives aux interfaces participant au processus de routage d'EIGRP. Ceci inclut mais ne se limite pas au nombre de voisins et le SRTT.
- **show ipx route**
 - Visualise la table de routage IPX.

A des fins de dépannage, les commandes **debug** suivantes sont disponibles :

- **debug eigrp packet**
 - Affiche les paquets EIGRP émis et reçus, sachant que le type de message peut être précisé.
- **debug eigrp neighbors**
 - Affiche les paquets Hello émis et reçus par le routeur ainsi que les voisins découverts.
- **debug ip eigrp**
 - Idem que **debug ip eigrp route**
- **debug ip eigrp route**
 - Affiche les changements dynamiques apportés à la table de routage.
- **debug ip eigrp summary**
 - Affiche un résumé des informations concernant EIGRP telles que les voisins, le filtrage et la redistribution.
- **debug eigrp events**
 - Affiche les types de paquets émis et reçus et les statistiques sur les décisions de routage.

4.8. Configuration

4.8.1. Configuration pour IP

La procédure de configuration du protocole EIGRP est la suivante :

- Activer le protocole EIGRP (commande **router eigrp**)
- Indiquer les interfaces devant participer au processus de routage d'EIGRP (commande **network**)
- Optionnel : Spécifier la bande passante réelle de la liaison (commande **bandwidth**)
- Optionnel : Désactiver l'émission/réception des informations de routage vers les interfaces connectées à des réseaux moignons (commande **passive-interface**)
- Optionnel : Meilleure gestion des routes (commandes **maximum-paths**, **variance** et **metric weights**)
- Optionnel : Agrégation de routes manuelle (commandes **no auto-summary** et **ip summary-address**)

4.8.2. Configuration pour IPX

Pour implémenter le protocole EIGRP sur un réseau IPX, il faut procéder comme suit :

- Activer la commutation de paquets IPX (commande **ipx routing**). Ceci active automatiquement IPX-RIP pour tous les réseaux rattachés au routeur.
- Activer IPX-EIGRP (commande **ipx router eigrp**)
- Indiquer les interfaces devant participer au processus de routage d'IPX-EIGRP (commande **network**)
- Optionnel : Indiquer les interfaces ne devant plus participer au processus de routage d'IPX-RIP (commande **no network**)
- Optionnel : Activer les mises à jour incrémentielles SAP d'IPX-EIGRP pour les interfaces LAN (commande **ipx sap-incremental eigrp**)

4.8.3. Préoccupation concernant la bande passante et configuration sur un réseau NBMA

La configuration d'EIGRP sur un réseau NBMA doit avoir une attention plus particulière de la part de l'administrateur réseau. En effet, en fonction du contexte, le fonctionnement par défaut peut parfois poser problème. Ceci concerne plus concrètement la bande passante qu'EIGRP utilise au maximum sur la liaison.

EIGRP peut utiliser par défaut jusqu'à 50% de la bande passante sur chaque liaison. Il se base sur la bande passante indiquée sur chaque interface. Cette valeur n'est pas détectée et a donc une valeur par défaut qui varie en fonction du type d'interface. C'est pourquoi il est très important d'indiquer à EIGRP quelle est la véritable bande passante pour chaque interface. Il existe deux commandes permettant de contrôler la gestion de la bande passante sous EIGRP :

- **bandwidth** : Expliciter la bande passante pour une interface donnée
- **ip bandwidth-percent eigrp** : Indique le pourcentage de la bande passante totale d'une interface qu'EIGRP peut utiliser au maximum. Par défaut, ce pourcentage est fixé à 50%.

Sur les réseaux NBMA, trois types d'interfaces existent, sachant qu'il faut appliquer une politique de gestion de bande passante spécifique pour chaque :

- **Interface point-à-point** : La bande passante indiquée doit correspondre au CIR du circuit virtuel.
- **Sous-interface point-à-point** : Idem que pour une interface point-à-point
- **Sous-interface point-à-multipoint** : La bande passante indiquée doit correspondre au résultat de la multiplication du plus petit CIR des différents circuits virtuels configurés sur cette sous-interface et du nombre de circuits virtuels de cette dernière.