



Formation Cisco CCNA

Support de cours

Auteur: Laboratoire SUPINFO des Technologies Cisco

Version 2.5 – 26 Juin 2006



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 : www.labo-cisco.com - 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. Notions de base en réseau :	6
1.1. Les différents systèmes de numération	6
1.2. Méthodes de conversion de base	7
1.3. La terminologie de base des réseaux	9
1.4. La bande passante numérique et le débit	9
1.5. Présentation du modèle de référence OSI	10
1.6. Le modèle TCP/IP	12
2. Couche 1 - La couche physique	14
2.1. Les notions de base sur les signaux et le bruit dans les systèmes de communication	14
2.2. Notions de base sur le codage de signaux réseau	15
2.3. Les médias de réseau local	18
2.4. Spécifications et raccordement des câbles	21
2.5. Les composants et les équipements de couche 1	22
2.6. Collisions et domaines de collision dans des environnements en couche partagés	23
2.7. Les topologies de base utilisées dans les réseaux	24
3. Couche 2 - La couche liaison de données	27
3.1. Les normes de réseau local	27
3.2. Les sous-couches LLC et MAC	28
3.3. Notions de base de la technologie Token Ring	29
3.4. Notions de base de l'interface FDDI (Fiber Distributed Data Interface)	30
3.5. Notions de base d'Ethernet et d'IEEE 802.3	31
3.6. Les équipements de couche 2	33
4. Couche 3 : La couche réseau	34
4.1. Principe de sélection du chemin	34
4.2. Les équipements de couche 3 : les routeurs	34
4.3. Les communications de réseau à réseau	35
4.4. Les services réseau de la couche 3	36
5. Couche 4 : La couche Transport	37
5.1. La couche transport	37
5.2. TCP et UDP	38
5.3. Les méthodes de connexion TCP	39
6. Couche 5 : La couche Session	41
6.1. Le contrôle du dialogue	41
6.2. La synchronisation du dialogue	41
6.3. La division du dialogue	41
7. Couche 6 : La couche présentation	42
7.1. Fonction et normes de la couche présentation	42
7.2. Cryptage et compression des données	42
8. Couche 7 : La couche application	43
8.1. Principes de la couche application	43
8.2. Le protocole DNS	44

9. La communication LAN.....	45
9.1. CSMA/CD	45
9.2. La commutation LAN.....	46
10. Protocole Spanning-Tree.....	48
10.1. Introduction.....	48
10.2. Théorie concernant Spanning-Tree.....	48
10.3. Théorie concernant Rapid Spanning-Tree	49
10.4. Commandes et configuration de Spanning-Tree.....	50
11. Adressage IP et Subnetting	51
11.1. Principe de l'adressage IP.....	51
11.2. Les sous-réseaux	53
12. Interface utilisateur du routeur.....	55
12.1. Différents modes du routeur	55
12.2. Mode SETUP	56
12.3. Fonctions d'aide du routeur	56
12.4. Utilisation des commandes d'éditions IOS	57
12.5. Utilisation de l'historique des commandes IOS.....	57
13. Composants d'un routeur	58
13.1. Sources de configuration externes	58
13.2. Composants de configuration internes et commandes d'état associées	58
14. Configuration du routeur.....	60
14.1. Fichiers de configuration d'un routeur	60
14.2. Configuration des mots de passe.....	60
14.3. Configuration du nom du routeur et des descriptions.....	61
15. Plate-forme logicielle Cisco IOS.....	62
15.1. Séquence d'amorçage	62
15.2. Caractéristiques fondamentales	62
15.3. Commandes boot system	63
15.4. Manipulation des images logicielles d'IOS	63
16. Adressage IP et interfaces	65
16.1. Adresse IP d'une interface.....	65
16.2. Résolution de nom vers IP statique.....	65
16.3. Service DNS	65
16.4. Spécificités des interfaces WAN	66
17. CDP	67
17.1. CDP.....	67
17.2. Théorie.....	67
17.3. Configuration.....	68
17.4. Visualisation et résolution de problèmes	68

18. VLSM et CIDR	69
18.1. Introduction au routage Classless	69
18.2. CIDR.....	70
18.3. VLSM	71
18.4. Procédure de réalisation.....	72
18.5. Configuration	74
19. VLAN	75
19.1. Concepts	75
19.2. Trunking	77
19.3. VTP	80
20. Le routage.....	82
20.1. Principes fondamentaux.....	82
20.2. Routage statique et dynamique	83
20.3. Routage à vecteur de distance.....	84
20.4. Routage à état des liens.....	85
20.5. Convergence, problème associé et solutions	86
20.6. Contexte des différents algorithmes de routage.....	87
20.7. Configuration initiale du routeur	87
20.8. Protocoles de routage intérieurs et extérieurs	88
21. Protocole RIP	89
21.1. Théorie.....	89
21.2. Configuration	90
22. Protocole RIPv2	92
22.1. Spécifications de RIPv2.....	92
22.2. Configuration	92
23. Protocole IGRP	94
23.1. Théorie.....	94
23.2. Configuration	96
23.3. Vérification	97
24. Protocole OSPF	98
24.1. Caractéristiques.....	98
24.2. Définitions	99
24.3. Fonctionnement dans un réseau ne comportant qu'une aire	100
24.4. Opérations OSPF	102
24.5. Construction de la table de routage.....	103
24.6. Commandes	104
25. Protocole EIGRP	106
25.1. Caractéristiques.....	106
25.2. Termes et définition	107
25.3. Métriques	108
25.4. Protocole Hello	110
25.5. DUAL	112
25.6. Commandes	113

26. ACL	115
26.1. Théorie	115
26.2. ACL standard	117
26.3. ACL étendue	117
26.4. ACL nommée	118
26.5. Mise en place et vérification des ACLs	119
27. NAT et PAT	120
27.1. Adressage privé et public	120
27.2. Translation d'adresses	120
27.3. Configuration	122
28. Protocole DHCP	124
28.1. Introduction	124
28.2. Configuration	127
29. Réseau WAN	129
29.1. Qu'est-ce qu'un réseau WAN	129
29.2. Les différents dispositifs WAN	129
29.3. Normes WAN	129
29.4. Technologies WAN	130
30. Tests de base et résolution de problèmes	132
30.1. Commandes de vérification	132
30.2. Erreurs courantes au niveau des trois premières couches du modèle OSI	133
30.3. Debugging	133
30.4. Procédure de récupération des mots de passe d'un routeur	134
31. Protocole PPP	135
31.1. Étude du protocole	135
31.2. Établissement d'une session	136
31.3. Authentification/configuration	137
32. Technologie RNIS	139
32.1. Technologie	139
32.2. Termes & équipements	140
32.3. Normes	141
32.4. Utilisation/implémentation	143
32.5. Routage à établissement de la connexion à la demande (DDR)	144
32.6. Commandes	145
32.7. Configuration	146
33. Technologie Frame Relay	147
33.1. Technologie	147
33.2. Interface LMI & DLCI	148
33.3. Fonctionnement, table de commutation & processus de transmission	149
33.4. Sous-interfaces Frame Relay	151
33.5. Commandes	153
33.6. Configuration	155

1. Notions de base en réseau :

1.1. Les différents systèmes de numération

1.1.1. Représentation des données pour un système informatique

Un ordinateur pourrait se résumer à un ensemble de commutateurs électriques pouvant prendre deux états :

- En fonction (le courant passe)
- Hors fonction (le courant ne passe pas)

Pour les différentes tâches qu'ils effectuent de nos jours, les ordinateurs utilisent le système de numérotation binaire.

Du fait que les humains fonctionnent avec le système décimal, l'ordinateur doit pouvoir effectuer cette traduction afin de pouvoir traiter les informations des utilisateurs. Ces nombres binaires sont exprimés en *bits*, qui constituent la plus petite unité d'information d'un ordinateur.

Un groupe de 8 bits correspond à un octet, qui représente un caractère de données. Pour un ordinateur, un octet représente également un emplacement de mémoire adressable.

Du fait de la taille des informations contenues dans les ordinateurs actuels, différentes unités de mesure ont été mises en place :

Unité	Définition	Octets	Bits	Exemples
Bit (b)	Chiffre binaire 1 ou 0	1 bit	1 bit	+5 volts ou 0 volts
Octet (o)	8 bits	1 octet	8 bits	01001100 correspond à la lettre L en ASCII
Kilo-octet (Ko)	1 kilo-octet =1 024 octets	1024 octets	8192 bits	mail type : 2ko premiers PC : 64Ko de Ram
Méga-octet (Mo)	1 méga-octet =1024 kilo-octets	1 048 576 octets	8 388 608 bits	disquette = 1,44 Mo CD-ROM = 650 Mo
Giga-octet (Go)	1 giga-octet =1024 méga-octets	1 048 576 kilo-octets	Env. 8 milliards de bits	disque dur type = 4 Go
Téra-octet (To)	1 téra-octet =1024 giga-octets	1 048 576 méga-octets	Env. 8 trillions de bits	quantité théorique de données transmissibles par une fibre optique en 1 seconde

Figure 1.1 Les différentes unités de mesure

1.1.2. Les différents systèmes de numération

L'homme est habitué dès le plus jeune âge à utiliser un système de numérotation pour représenter des valeurs. Ce système comporte 10 symboles : 0 1 2 3 4 5 6 7 8 9 et se nomme « *système de numérotation décimal* ».

Une valeur est de ce fait une notion abstraite pouvant être exprimée selon différents systèmes :

Par exemple, nous savons qu'un ordinateur fonctionne avec des commutateurs électriques pouvant avoir 2 états : en fonction et hors fonction. L'ordinateur va donc utiliser un système de numérotation avec deux symboles : 0 et 1. C'est ce que l'on appelle le système binaire. Il fonctionne de manière analogue au système décimal sauf qu'il n'utilise que 2 symboles.

Exemple : 1011 qui en décimal équivaut à la valeur 11

Autres systèmes, le système hexadécimal, comportant 16 symboles 0 1 2 3 4 5 6 7 8 9 A B C D E F. Les 6 lettres correspondent en décimal à 10 11 12 13 14 15.

Exemple : A2F54B qui équivaut en décimal à la valeur 10679627

Voici les différents systèmes que nous utiliserons ainsi que leur spécificité :

Nom	Symboles utilisés	Référence
binaire	0 1	2
octal	0 1 2 3 4 5 6 7	8
décimal	0 1 2 3 4 5 6 7 8 9	10
hexadécimal	0 1 2 3 4 5 6 7 8 9 A B C D E F	16

Figure 1.2 : les différents systèmes de numérotation

1.2. Méthodes de conversion de base

Le système décimal repose sur les puissances de 10. Chaque symbole composant un nombre décimal représente une puissance de 10 ; chacun ayant pour exposant sa position dans le nombre en partant de la droite ; multiplié par le symbole occupant cette position.

Exemple : $25642 = 2 \times 10^4 + 5 \times 10^3 + 6 \times 10^2 + 4 \times 10^1 + 2 \times 10^0$

Pour convertir une valeur exprimée dans un système en une valeur utilisant le système de numérotation décimal, l'algorithme est le suivant :

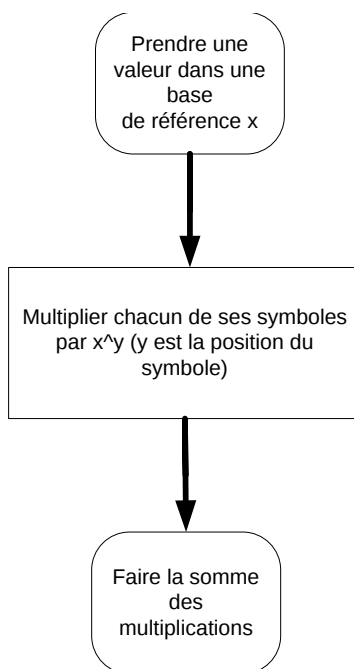


Figure 1.3 : algorithme de conversion d'un système en base x vers le système décimal

:

Exemple de conversion Octal -> Décimal

20165 (octal)

$$2 \times 8^4 + 0 \times 8^3 + 1 \times 8^2 + 6 \times 8^1 + 5 \times 8^0 = 8192 + 0 + 64 + 48 + 5 = 8309(dec)$$

Exemple de conversion Hexadécimal → Décimal

A2F54B (hexadécimal)

$$\begin{aligned}
 &= A(10) \times 16^5 + 2 \times 16^4 + F(15) \times 16^3 + 5 \times 16^2 + 4 \times 16^1 + B(11) \times 16^0 \\
 &= 10485760 + 131072 + 61440 + 1280 + 64 + 11 = 10679627(dec)
 \end{aligned}$$

Exemple de conversion Décimal → Hexadécimal

1036 (dec):

$$\frac{1036}{16} = 64 \text{ reste } 12 (C)$$

$$\frac{64}{16} = 4 \text{ reste } 0$$

$$\frac{4}{16} = 0 \text{ reste } 4$$

$$1036 (dec) = 40C (hex)$$

1.3. La terminologie de base des réseaux

Un réseau est par définition un ensemble d'entités communicant entre elles. Nous allons nous intéresser dans le cadre de ce cours à ce que l'on nomme des réseaux de données ou réseaux informatiques. Ces réseaux sont apparus suite à une demande des entreprises qui recherchaient une méthode pour éviter la duplication des imprimantes et une simplification des communications de données entre des équipements informatiques.

Un réseau de données est donc un ensemble d'entités informatiques communicant ensemble. La première classification de réseau que nous allons faire s'établit sur la base des distances entre les communicants.

Nous allons différencier :

- Les réseaux LAN :
 - o Couvrent une région géographique limitée
 - o Permettent un accès multiple aux médias à large bande
 - o Ils assurent une connectivité continue aux services locaux (Internet, messagerie, ...)
 - o Ils relient physiquement des unités adjacentes
 - Exemple : Une salle de classe
- Les réseaux WAN :
 - o Couvrent une vaste zone géographique
 - o Permettent l'accès par des interfaces séries plus lentes
 - o Assurent une connectivité pouvant être continue ou intermittente
 - o Relient des unités dispersées à une échelle planétaire
 - Exemple : Internet

1.4. La bande passante numérique et le débit

La bande passante d'un réseau représente sa capacité, c'est-à-dire la quantité de données pouvant circuler en une période donnée. Celle-ci se mesure en bits par seconde. Du fait de la capacité des supports réseau actuels, les différentes conventions suivantes sont utilisées :

Unité de bande passante	Abréviation	Équivalence
Bits par seconde	bits/s	1 bit/s = unité fondamentale
Kilobits par seconde	kbits/s	1kbit/s = 1000 bits/s
Mégabits par seconde	Mbits/s	1Mbit/s = 1 000 000 bits/s
Gigabits par seconde	Gbits/s	1Gbit/s = 1 000 000 000 bits/s

Figure 1.5 : unités de mesure de bande passante

A cette notion de bande s'ajoute celle de débit. Le débit est la bande passante réelle, mesurée à un instant précis de la journée. Ce débit est souvent inférieur à la bande passante ; cette dernière représentant le débit maximal du média ; en raison :

- o des unités d'interconnexion de réseaux et de leur charge
- o du type de données transmises
- o de la topologie du réseau
- o du nombre d'utilisateur
- o de l'ordinateur de l'utilisateur et du serveur
- o des coupures d'électricité et autres pannes

De ce fait, le temps de téléchargement d'un fichier peut se mesurer de la manière suivante :

$$\text{Temps de téléchargement (seconde)} = \text{Taille du fichier (bit)} / \text{débit (bit/seconde)}$$

1.5. Présentation du modèle de référence OSI

1.5.1. Le modèle général de communication à couche

La première évolution des réseaux informatiques a été des plus anarchiques, chaque constructeur développant presque sa propre technologie. Le résultat de cela était une quasi-impossibilité de connecter différents réseaux entre eux. Pour palier à cela, l'ISO (Institut de normalisation) décida de mettre en place un modèle de référence théorique décrivant le fonctionnement des communications réseau : le modèle OSI ; à partir des structures réseau prédominantes de l'époque : DecNet et SNA.

Le but de ce modèle est d'analyser la communication en découpant les différentes étapes en 7 couches ; chacune de ces couches remplissant une tâche bien spécifique :

- o Quelles sont les informations qui circulent ?
- o Sous quelle forme circulent-elles ?
- o Quel chemin empruntent-elles ?
- o Quelles règles s'appliquent aux flux d'informations ?

n°	Nom	Description
7	Application	Communication avec les logiciels
6	Présentation	Gestion de la syntaxe
5	Session	Contrôle du dialogue
4	Transport	Qualité de la transmission
3	Réseau	Sélection du chemin
2	Liaison de données	Préparation de l'envoi sur le média
1	Physique	Envoi sur le média physique

Figure 2.1 : les 7 couches du modèle OSI

Les avantages de ce modèle sont :

- o Une division de la communication réseau en éléments plus petits et plus simples pour une meilleure compréhension
- o L'uniformisation des éléments afin de permettre le développement multi constructeur
- o La possibilité de modifier un aspect de la communication réseau sans modifier le reste (Exemple : un nouveau média)

1.5.2. L'encapsulation

Pour communiquer entre les couches et entre les hôtes d'un réseau, OSI a recourt au principe d'encapsulation.

Encapsulation : processus de conditionnement des données consistant à ajouter un en-tête de protocole déterminé avant que les données ne soient transmises à la couche inférieure :

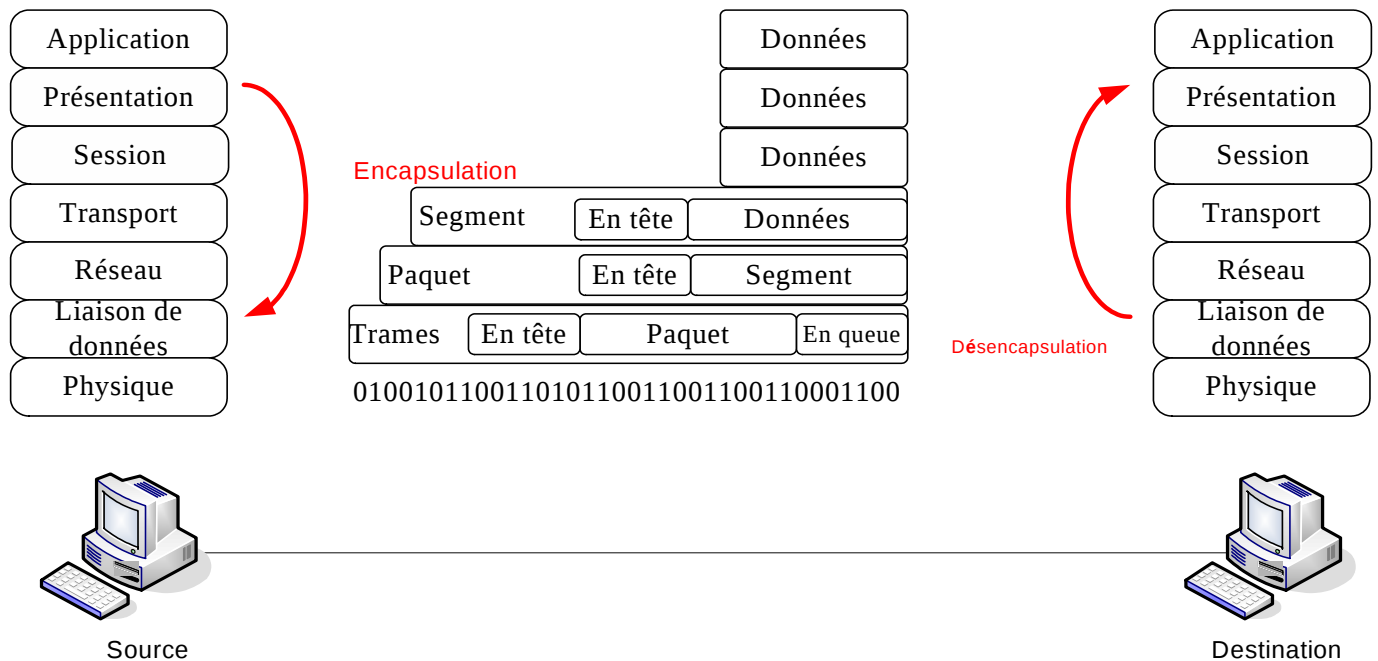


Figure 2.2 : Principe de l'encapsulation

Lorsque 2 hôtes communiquent, on parle de communication d'égal à égal ; c'est-à-dire que la couche n de la source communique avec la couche n du destinataire.

Lorsqu'une couche de la source reçoit des données, elle encapsule ces dernières avec ses informations puis les passe à la couche inférieure. Le mécanisme inverse a lieu au niveau du destinataire où une couche réceptionne les données de la couche inférieure ; enlève les informations la concernant ; puis transmet les informations restantes à la couche supérieure. Les données transitant à la couche n de la source sont donc les mêmes que les données transitant à la couche n du destinataire.

Pour identifier les données lors de leur passage au travers d'une couche, l'appellation « *Unité de données de protocole (PDU)* » est utilisée.

Couche	Désignation
7	Données
6	
5	
4	Segment
3	Paquets
2	Trame
1	Bits

Figure 2.3 : Les PDU des différentes couches

1.6. Le modèle TCP/IP

1.6.1. Présentation de TCP/IP

La forme actuelle de TCP/IP résulte du rôle historique que ce système de protocoles a joué dans le parachèvement de ce qui allait devenir Internet. À l'instar des nombreux développements de ces dernières années, Internet est issu des recherches lancées aux Etats-Unis par le DOD, département de la défense.

Ses fonctions essentielles ne devaient en aucun cas se trouver en un seul point, ce qui le rendrait trop vulnérable. C'est alors que fut mis en place le projet Arpanet (Advanced Research Projects Agency du DOD), qui allait devenir par la suite le système d'interconnexion de réseau qui régit ce que l'on appelle aujourd'hui l'Internet : TCP/IP.

TCP/IP est un modèle comprenant 4 couches :

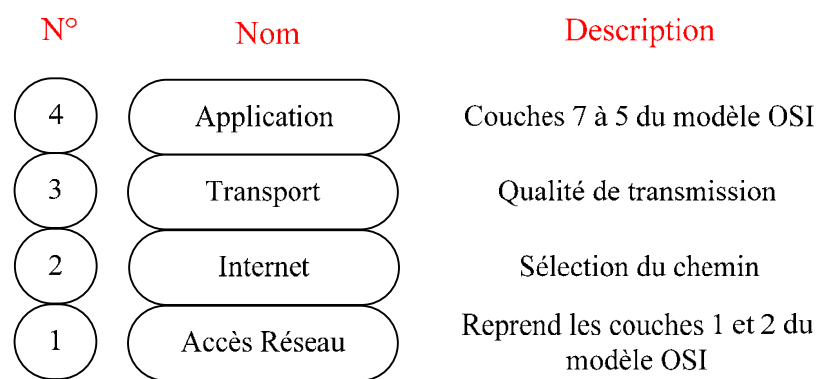


Figure 2.4 : les 4 couches de TCP/IP

1.6.2. Protocole orienté/non orienté connexion

Protocole : Ensemble formel de règles et de conventions qui régit l'échange d'informations entre des unités en réseau

Dans un protocole orienté connexion, TCP/IP établit un dialogue entre la source et le destinataire pendant qu'il prépare les informations de la couche application en segments. Il y a alors un échange de segments de couche 4 afin de préparer une communication et donc une connexion logique pendant un certain temps.

Cette communication faisant appel à un circuit logique temporaire est appelé commutation de paquets, en opposition à la commutation de circuits supposant elle un circuit permanent.

Un protocole non orienté connexion envoie les données sur le réseau sans qu'un circuit ait été établi au préalable.

1.6.3. Comparaison entre OSI et TCP/IP

Ces deux protocoles sont très similaires, dans la mesure où les 2 sont des modèles de communication à couche et utilisent l'encapsulation de données.

On remarque cependant deux différences majeures :

- o TCP/IP regroupe certaines couches du modèle OSI dans des couches plus générales
- o TCP/IP est plus qu'un modèle de conception théorique, c'est sur lui que repose le réseau Internet actuel

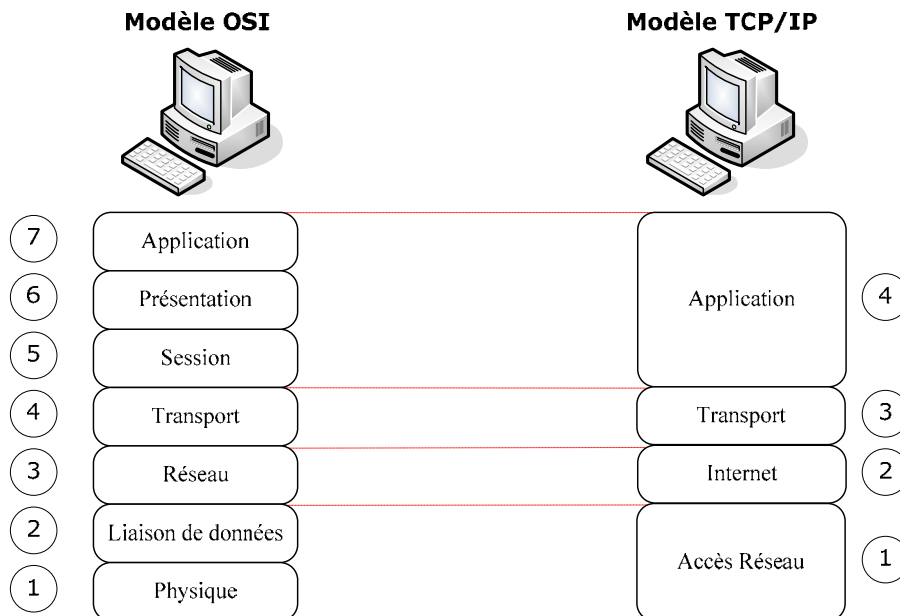


Figure 2.5 : les modèles OSI et TCP/IP

2. Couche 1 - La couche physique

2.1. Les notions de base sur les signaux et le bruit dans les systèmes de communication

2.1.1. La représentation d'un bit dans un média physique

La composante de base de l'information dans les réseaux est le bit. Dans le cas d'un signal électrique, un bit correspond à une impulsion signifiant 0 ou 1.

Exemple :

- o 0 : 0 volt et 1 +5 volts dans le cas d'un signal électrique
- o 0 : faible intensité et 1 forte intensité dans le cas d'un signal optique
- o 0 : courte rafale d'onde et 1 rafale d'onde plus longue dans le cas de transmission sans fil

Mise à la terre de référence : masse électrique permettant d'établir la ligne 0 dans les graphiques de signalisation

2.1.2. Les facteurs pouvant affecter un bit

Il existe différents facteurs pouvant affecter le signal et de ce fait les bits transportés sur le média :

- o **La propagation de signaux réseau :**
 - o Le terme de propagation fait référence au temps que met un bit ; c'est-à-dire une impulsion ; à se déplacer dans le média. Il est impératif que la propagation soit homogène dans le réseau.
- o **L'atténuation du signal réseau :**
 - o Perte de la force du signal. Ce problème est limitable par un bon choix des médias réseau utilisés
- o **La réflexion réseau :**
 - o Retour d'énergie causé par le passage des impulsions dans le média. Si ce retour est trop fort, il peut perturber le signal des impulsions suivantes. Le système binaire ; et donc à 2 états ; peut être perturbé par ces énergies supplémentaires se déplaçant dans le média.
- o **Le bruit :**
 - o Ajout indésirable à un signal. Des sources d'énergie situées à proximité du média fournissent un supplément d'énergie venant perturber le signal.
 - *Diaphonie* : bruit ajouté au signal d'origine d'un conducteur par l'action du champ magnétique provenant d'un autre conducteur
 - *Paradiaphonie* : diaphonie causée par un conducteur interne au câble
 - Le bruit peut être causé par des sources d'alimentations externes, des variations thermiques, des interférences électromagnétiques ou encore des interférences de radio fréquences.
- o **La dispersion :**
 - o Étalement des impulsions dans le temps. Si la dispersion est trop forte, le signal d'un bit peut recouper le signal du précédent ou du suivant. La durée d'une impulsion est fixe, la dispersion correspond à une modification de cette durée au fur et à mesure que le signal se propage dans le média.
- o **La gigue :**

- o Les systèmes numériques sont synchronisés, tout est réglé par des impulsions d'horloge. Si les horloges de la source et du destinataire ne sont pas synchronisées, on obtient alors une *gigue de synchronisation*.
- o **La latence :**
 - o Retard de transmission. Principalement du au déplacement du signal dans le média et à la présence de composants électroniques entre la source et la destination.
- o **Les collisions :**
 - o Se produit lorsque 2 ordinateurs utilisant le même segment de réseau émettent en même temps. Les impulsions se mélangent, détruisant alors les données.

Dès qu'un bit accède au média, il est sujet à tous ces paramètres pouvant perturber la transmission. Dans la mesure où le but n'est pas de transmettre un bit, mais des quantités gigantesques (parfois 1 milliard de bits à la seconde) ; ces paramètres ne sont pas à négliger, car le moindre défaut peut avoir des conséquences importantes sur la qualité de la transmission

2.2. Notions de base sur le codage de signaux réseau

Nous allons donc nous intéresser ici aux méthodes de transmission de bits de façon brute entre l'émetteur et le récepteur.

Tout d'abord, une liaison entre 2 équipements A et B peut être :

- o Simple (unidirectionnelle) : A est toujours l'émetteur et B le récepteur. C'est ce que l'on trouve par exemple entre un banc de mesure et un ordinateur recueillant les données mesurées.
- o Half-duplex (bidirectionnelle à l'alternat) : Le rôle de A et B peut changer, la communication change de sens à tour de rôle (principe talkies-walkies).
- o Full-duplex (bidirectionnelle simultanée) : A et B peuvent émettre et recevoir en même temps (comme dans le cas du téléphone).

La transmission de plusieurs bits peut s'effectuer :

- o En série : les bits sont envoyés les uns derrière les autres de manière synchrone ou asynchrone :
 - o Dans le mode synchrone, l'émetteur et le récepteur se mettent d'accord sur une base de temps (un top d'horloge) qui se répète régulièrement durant tout l'échange. À chaque top d'horloge (ou k tops d'horloge k entier fixé définitivement) un bit est envoyé et le récepteur saura ainsi quand arrivent les bits.
 - o Dans le mode asynchrone, il n'y a pas de négociation préalable, mais chaque caractère envoyé est précédé d'un bit de start et immédiatement suivi d'un bit de stop. Ces deux bits spéciaux servent à caler l'horloge du récepteur pour qu'il échantillonne le signal qu'il reçoit afin d'y décoder les bits qu'il transmet.
- o En parallèle : Les bits d'un même caractère sont envoyés en même temps chacun sur un fil distinct, mais cela pose des problèmes de synchronisation et n'est utilisé que sur de courtes distances (bus par exemple).

2.2.1. Transmission en bande de base

La transmission en bande de base consiste à envoyer directement les suites de bits sur le support à l'aide de signaux carrés constitués par un courant électrique pouvant prendre 2 valeurs (5 Volts ou 0 Volts par exemple). L'émetteur envoie sur la ligne un signal carré du type de celui de la figure ci-dessous pour la séquence de bits 1010 :

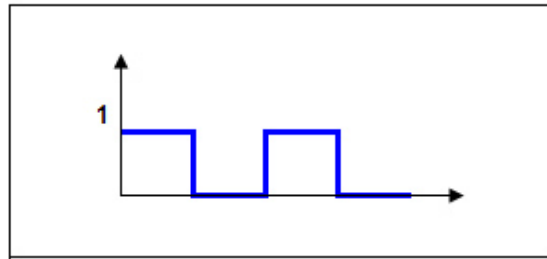


Figure 3.2 : Signal carré de la séquence de bits 1010

Il existe de nombreuses possibilités de coder sur le signal sur un média, voici différents exemples :

- o **Le code tout ou rien** : c'est le plus simple, un courant nul code le 0 et un courant positif indique le 1

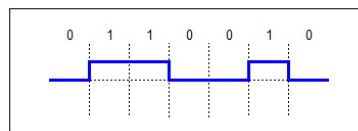


Figure 3.3 : Le code tout ou rien

- o **Le code NRZ** : (non-retour à zéro): pour éviter la difficulté à obtenir un courant nul, on code le 1 par un courant positif et le 0 par un courant négatif.

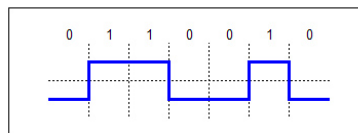


Figure 3.4 : Le code NRZ

- o **Le code bipolaire** : c'est aussi un code tout ou rien dans lequel le 0 est représenté par un courant nul, mais ici le 1 est représenté par un courant alternativement positif ou négatif pour éviter de maintenir des courants continus.

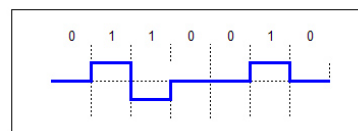


Figure 3.5 : Le code bipolaire

- o **Le code RZ** : le 0 est codé par un courant nul et le 1 par un courant positif qui est annulé au milieu de l'intervalle de temps prévu pour la transmission d'un bit.

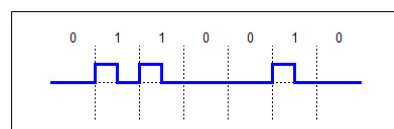


Figure 3.6 : Le code RZ

- o **Le code Manchester** : ici aussi le signal change au milieu de l'intervalle de temps associé à chaque bit. Pour coder un 0 le courant sera négatif sur la première moitié de l'intervalle et positif sur la deuxième moitié, pour coder un 1, c'est l'inverse. Autrement dit, au milieu de l'intervalle il y a une transition de bas en haut pour un 0 et de haut en bas pour un 1.

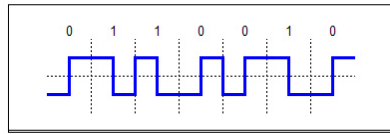


Figure 3.7 : Le code Manchester

- o **Le code Miller** : on diminue le nombre de transitions en effectuant une transition (de haut en bas, ou l'inverse) au milieu de l'intervalle pour coder un 1 et en n'effectuant pas de transition pour un 0 suivi d'un 1. Une transition est effectuée en fin d'intervalle pour un 0 suivi d'un autre 0.

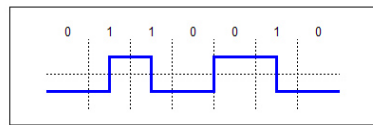


Figure 3.8 : Le code Miller

2.3. Les médias de réseau local

2.3.1. Le câble à paires torsadées non blindées

Fiche technique :

Désignation : UTP

Vitesse : 10 – 100 Mbits/s

Longueur Max. : 100m

Raccordement : Connecteur RJ-45

Impédance : 100 Ohms

Coût : Faible

Le câble UTP est composé de 4 paires de fils torsadés 2 à 2, chacune de ses paires étant isolé des autres. Ce câble compte uniquement sur l'effet d'annulation produit par les paires torsadées pour limiter la dégradation du signal causée par une perturbation électromagnétique et une interférence radioélectrique.

Annulation :

Afin de réduire au maximum la diaphonie entre les paires d'un câble à paires torsadées non blindées, le nombre de torsades des paires de fils doit respecter exactement le nombre de torsades permises par mètre de câble.

Avantages

- o Simple à installer
- o Peu coûteux
- o Petit diamètre (pour installation dans des conduits existants)

Inconvénient :

- o Sensible aux interférences



Figure 3.10 : câble UTP

2.3.2. Le câble à paires torsadées blindées

Fiche technique :

Désignation : STP

Vitesse : 10 – 100 Mbits/s

Longueur Max. : 100m

Raccordement : Connecteur RJ-45

Impédance : 100 Ohms

Coût : Moyennement cher

Le câble à paires torsadées blindées présente tous les avantages et désavantages du câble à paires torsadées non blindées en assurant cependant une plus grande protection contre toute interférence externe au prix certes d'un diamètre plus élevé

Le blindage de ce type de câble doit être mis à la terre lors de son installation, si cela n'est pas effectué correctement, de nombreux problèmes peuvent survenir, car le blindage agit comme une antenne en absorbant les signaux électriques des autres fils du câble et des parasites électriques externes au câble.



Figure 3.11 : câble STP

2.3.3. Le câble coaxial

Fiche technique :

Désignation : Coaxial

Vitesse : 10 – 100 Mbits/s

Longueur Max. : 500m

Raccordement : Connecteur BNC

Impédance : 150 Ohms

Coût : Peu cher

Un câble coaxial est constitué d'un fil de cuivre entouré d'un isolant flexible, lui-même entouré d'une torsade de cuivre ou d'un ruban métallique qui agit comme le second fil du circuit et comme protecteur du conducteur intérieur. Cette deuxième couche ou protection peut aider à réduire les interférences externes. Une gaine de câble enveloppe ce blindage.

Le câble coaxial offre de nombreux avantages du fait de sa capacité à s'étendre sur une plus grande distance et de son coût parmi les plus faibles. C'est une technologie utilisée depuis de nombreuses années pour tous les types de communications de données.

Le câble coaxial existe en plusieurs variantes :

- o Thicknet : Epais et raide à cause de son blindage, il est recommandé pour l'installation de câble fédérateur. Sa gaine est jaune.

Laboratoire Supinfo des Technologies Cisco

Site : www.labo-cisco.com - Mail : labo-cisco@supinfo.com

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

- o Thinnet : D'un diamètre plus réduit, il est plus pratique dans des installations comprenant des courbes. De plus, il est plus économique, mais dispose d'un blindage moins conséquent.
- o Cheapernet : Version économique et de faible diamètre du câble coaxial.



Figure 3.12 : câble Thinnet



Figure 3.13 : câble Thicknet

Il importe d'apporter une attention particulière à la mise à la terre. On doit assurer une solide connexion électrique aux deux extrémités du câble. Manquer à ce principe entraîne des parasites électriques qui causent une interférence au niveau de la transmission du signal du média réseau.

2.3.4. La fibre optique

Fiche technique :

Désignation : FDDI

Vitesse : 100+ Mbits/s

Longueur Max. : 2km en multimode et 3km en monomode

Raccordement : Connecteur multi mode ou monomode

Coût : Cher

Le câble à fibre optique est un support transmettant des impulsions lumineuses. Ce type de média est très coûteux, mais est insensible aux interférences électromagnétiques et peut acheminer des données à un débit très élevé

Le câble à fibre optique utilisé pour les réseaux comprend deux fibres encapsulées dans des enveloppes distinctes. En examinant la coupe transversale, d'un câble optique, il est possible de voir que chaque fibre est entourée de couches de revêtements optiques réfléchissants, un enduit de plastique fait en Kevlar, et que l'ensemble est entouré d'une gaine extérieure assurant la protection de l'ensemble du câble.



Figure 3.14 : fibre optique

2.3.5. Les connecteurs RJ45

Le raccordement 10BaseT standard (le connecteur de point d'extrémité sans prise) est le RJ-45. Il réduit les parasites, la réflexion et les problèmes de stabilité mécanique et ressemble à une prise téléphonique, sauf qu'il compte huit conducteurs au lieu de quatre.

Les connecteurs RJ-45 s'insèrent dans les réceptacles ou les prises RJ-45. Les prises mâles RJ-45 ont huit connecteurs qui s'enclenchent avec la prise RJ-45. De l'autre côté de la prise RJ-45, il y a un bloc où les fils sont séparés et fixés dans des fentes avec l'aide d'un outil semblable à une fourche. Ceci offre un passage de courant en cuivre aux bits.

Laboratoire Supinfo des Technologies Cisco

Site : www.labo-cisco.com - Mail : labo-cisco@supinfo.com

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

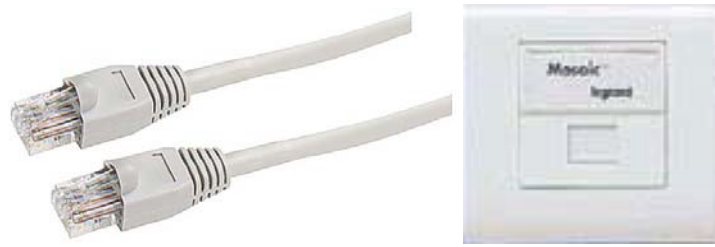


Figure 3.15 : connecteurs RJ45 et prise murale

2.3.6. Les communications sans fil

Les signaux sans fil sont des ondes électromagnétiques qui peuvent circuler dans le vide ou dans des médias tels que l'air. De ce fait, ils ne requièrent aucun média physique.

Pour communiquer, un réseau LAN sans fil utilise :

- o des ondes radio (Ex. : 902MHz)
- o des micro-ondes (Ex. : 2.4GHz)
- o des ondes infrarouges (Ex. : 820 nanomètres)

2.4. Spécifications et raccordement des câbles

Dans le but d'ajouter aux spécifications de l'ISO des normes visant à standardiser les équipements, divers organismes ont mis en place différentes normes. Ces organismes sont :

- o IEEE : Institute of Electrical and Electronics Engineers
- o UL : Underwriters laboratories
- o EIA : Electronic Industries Alliance
- o TIA : Telecommunications Industry Association

2.4.1. Les normes TIA/EIA

Elles définissent les configurations minimales d'installation tout en laissant toute liberté quand au choix du fournisseur d'équipement

Les principales normes sont :

- o TIA/EIA-568-A : norme de câblage pour les télécommunications dans les édifices commerciaux
- o TIA/EIA-569-A : Norme relative aux espaces et aux voies de télécommunications dans les édifices commerciaux

2.4.2. Les spécifications de câblage de la norme TIA/EIA-568-A

Câblage horizontal : câblage situé entre la prise murale et une interconnexion horizontale. Il inclut le média réseau installé horizontalement, la prise ainsi que les terminaisons mécaniques. Il comprend donc le média réseau allant de l'armoire de câblage jusqu'à une zone de travail.

La norme autorise les longueurs suivantes :

- o Longueur maximale d'un câblage horizontal : 90m
- o Longueur maximale des câbles d'interconnexion : 6m
- o Longueur maximale des câbles de raccordement (pour relier les unités réseau au câblage horizontal) : 3m

Laboratoire Supinfo des Technologies Cisco

Site : www.labo-cisco.com - Mail : labo-cisco@supinfo.com

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

De plus, la norme exige la mise à la terre de tous les câbles.

En ce qui concerne le choix du type de câblage, la norme comprend des spécifications définissant les performances des câbles : CAT1, CAT2, CAT3, CAT4 et CAT5. De nos jours, seules les catégories 3, 4 et 5 sont reconnues pour les réseaux locaux.

2.5. Les composants et les équipements de couche 1

De nos jours, les 3 technologies LAN les plus répandues sont :

- o Ethernet
- o Token Ring
- o FDDI

Ethernet étant le plus répandu, nous allons étudier les composants de couche 1 de cette technologie. Le support physique de cette technologie est le câble à paires torsadées.

Composant passif : Qui n'a pas besoin d'une source d'alimentation externe pour fonctionner

Composant actif : Qui nécessite une alimentation externe pour remplir ses fonctions

2.5.1. Les émetteurs-récepteurs

Un émetteur-récepteur (transceiver) convertit un signal en un autre. Il est souvent intégré aux cartes réseau.



Figure 3.17 : émetteurs-récepteurs

2.5.2. Les répéteurs et les concentrateurs

Le répéteur est un composant actif. Son rôle est de régénérer et de re-synchroniser le signal afin de pouvoir étendre la portée des câbles.

Le concentrateur ; ou répéteur multi ports ; reprend le fonctionnement du répéteur en ajoutant une fonctionnalité de connectivité. En effet, il dispose de plusieurs ports ce qui permet d'interconnecter plusieurs équipements réseau. Chaque signal arrivant sur un port est régénéré, re-synchronisé et ré émis au travers de tous les autres ports.



Figure 3.18 : concentrateur 8 ports

Tous ces équipements, passifs ou actifs, créent ou manipulent des bits. Ils ne reconnaissent aucune information dans les bits, ni les adresses, ni les données. Leur fonction se limite donc à déplacer les bits.

2.6. Collisions et domaines de collision dans des environnements en couche partagés

2.6.1. Environnement de média partagé

- o Environnement de médias partagés : Plusieurs hôtes se partagent le même média
- o Environnement de médias partagés étendu : Type d'environnement dans lequel des équipements réseau étendent le réseau afin que celui-ci soit accessible à un plus grand nombre d'utilisateurs.
- o Environnement de réseau point à point : Environnement réseau le plus répandu dans les réseaux commutés. Un équipement est directement connecté à un seul autre équipement.

2.6.2. Les réseaux commutés

Dans les réseaux directement connectés, certains équipements de couche supérieure et/ou une grande distance se situe entre les 2 hôtes. On parle alors de réseaux indirectement connectés dont on distingue 2 types :

- o Réseau à commutation de circuits : Réseau connecté indirectement dans lequel de réels circuits électriques sont maintenus pendant la durée de la communication
- o Réseau à commutation de paquets : Au lieu de dédier une liaison à une connexion exclusive entre 2 hôtes, la source envoie les messages par paquets, chaque paquet contenant suffisamment d'informations pour être acheminé vers l'hôte de destination

2.6.3. Collisions et domaines de collision

Si 2 hôtes du réseau émettent en même temps sur un même segment de réseau, les informations se chevauchent : c'est ce que l'on appelle une collision.

Lorsque cela survient, un hôte le détecte. A ce moment, il envoie un signal de bourrage annonçant le problème à tous les autres. A la réception de ce signal, tous les hôtes arrêtent d'émettre. Chacun calcule alors une valeur aléatoire correspondant au délai précédent une nouvelle tentative d'émission.

L'autre terme pour décrire un environnement de média partagé est « Domaine de collision », à savoir une section de réseau où tous les hôtes partagent le même média.

Des équipements comme le répéteur ou le concentrateur n'effectuant aucun filtrage étendent ce domaine de collision.

Les topologies de base utilisées dans les réseaux

Topologie : décrit la manière dont les équipements réseau sont connectés entre eux. Nous distinguerons les topologies physiques ; décrivant la manière dont les équipements sont reliés par des médias ; des topologies logiques ; décrivant la manière dont les équipements communiquent.

2.6.4. La topologie en bus

- o Perspective Physique : Tous les hôtes sont connectés directement à une liaison :
- o Perspective logique : Tous les hôtes voient tous les signaux provenant de tous les autres équipements

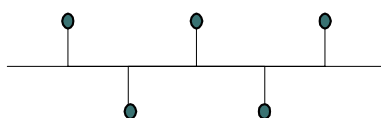


Figure 3.19 : topologie en bus

2.6.5. La topologie en anneau

- o Perspective physique : Les éléments sont chaînés dans un anneau fermé
- o Perspective logique : Chaque hôte communique avec ses voisins pour véhiculer l'information

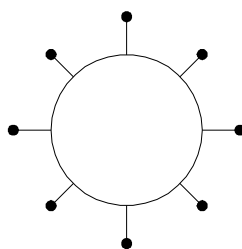


Figure 3.20 : topologie en anneau

Une variante de cette topologie est le double anneau où chaque hôte est connecté à 2 anneaux. Ces deux anneaux ne communiquent pas entre eux. Le deuxième anneau est utilisé comme lien redondant en cas de panne sur le premier.

2.6.6. La topologie en étoile

- o Perspective physique : Cette topologie comporte un nœud central d'où partent toutes les liaisons avec les autres nœuds.
- o Perspective logique : Toutes les informations passent par un seul équipement, par exemple un concentrateur

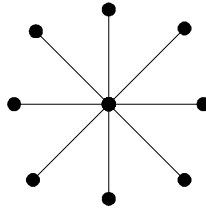


Figure 3.21 : topologie en étoile

2.6.7. La topologie en étoile étendue

Cette topologie est identique à la topologie en étoile si ce n'est que chaque nœud connecté au nœud central est également le centre d'une autre étoile.

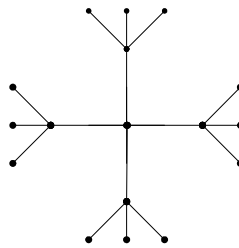


Figure 3.22 : topologie en étoile étendue

2.6.8. La topologie hiérarchique

- o Perspective physique : Cette topologie ressemble à une topologie en étoile sauf qu'elle n'utilise pas de nœud central. Elle utilise un nœud de jonction à partir duquel elle se branche vers d'autres nœuds.
- o Perspective logique : Le flux d'informations est hiérarchique

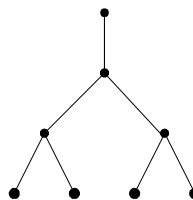


Figure 3.23 : topologie hiérarchique

2.6.9. La topologie complète (maillée)

- o Perspective physique : Chaque nœud est connecté avec tous les autres
- o Perspective logique : Dépend des équipements utilisés

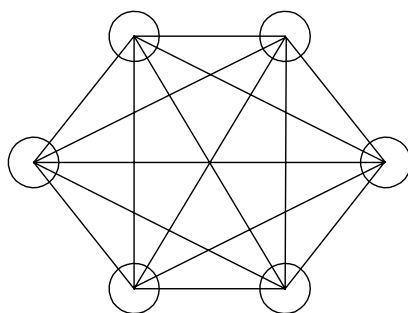


Figure 3.24 : topologie complète

3. Couche 2 - La couche liaison de données

Le modèle OSI comprend 2 couches dites « matérielles » ; en opposition aux couches logicielles. La couche 1 englobe les médias, les signaux ainsi que les bits se déplaçant sur diverses topologies.

La couche Liaison de données a pour fonction de combler tous les manques de la couche physique afin de permettre la communication réseau

3.1. Les normes de réseau local

3.1.1. IEEE et le modèle OSI

Les normes IEEE sont actuellement les normes pré dominantes. Selon l'IEEE, on divise la partie matérielle du modèle OSI en 2 parties :

- o La norme LLC 802.2, ne dépendant pas de la technologie du média utilisé
- o Les éléments spécifiques, tributaires de la technologie, qui intègrent la couche physique du modèle OSI

De plus, cette division sépare la couche Liaison de données en 2 parties :

- o Média Access Control (MAC) : transmission vers le bas jusqu'au média
- o Logical Link Control (LLC) : transmission vers le haut jusqu'à la couche réseau

Couches 1 et 2 du modèle OSI

Spécifications Ethernet

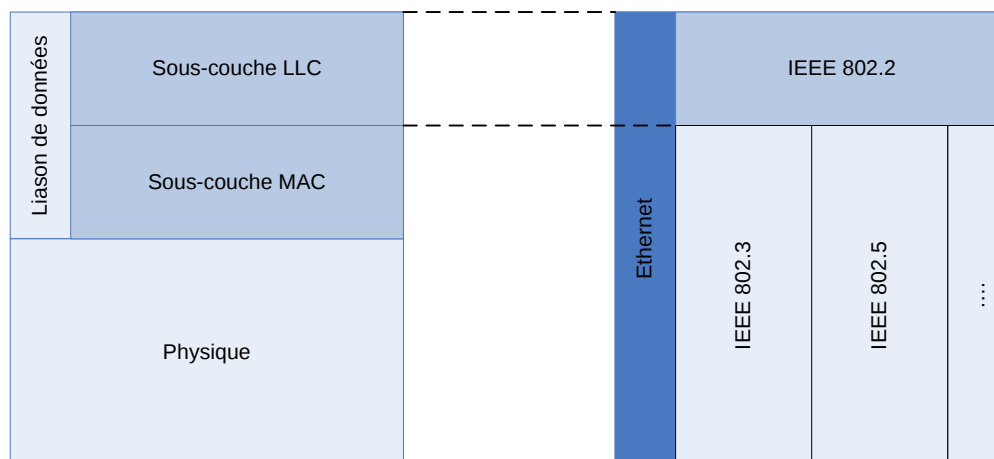


Figure 4.1 : différences entre le modèle OSI et les spécifications de l'IEEE

3.1.2. Les adresses MAC

Une adresse MAC est une adresse matérielle ; c'est-à-dire une adresse unique non modifiable par l'administrateur et stockée sur une mémoire morte (ROM) de la carte réseau.

Les adresses MAC comportent 48bits et sont exprimées sous la forme de 12 chiffres hexadécimaux :

- o 6 chiffres sont administrés par l'IEEE et identifient le fabricant de la carte
- o 6 chiffres forment le numéro de série de la carte

On peut les représenter de 2 manières différentes : par groupe de 4 chiffres séparés par des points ou par groupe de 2 chiffres séparés par des tirets

Exemple : 0000.0c12.3456 OU 00-00-0c-12-34-56

Les LANs de type Ethernet et 802.3 sont des réseaux dits de broadcast, ce qui signifie que tous les hôtes voient toutes les trames. L'adressage MAC est donc un élément important afin de pouvoir déterminer les émetteurs et les destinataires en lisant les trames.

3.2. Les sous-couches LLC et MAC

3.2.1. Le contrôle de lien logique (LLC)

La sous-couche LLC a été créée afin de permettre à une partie de la couche liaison de données de fonctionner indépendamment des technologies existantes.

Cela assure la polyvalence des services fournis aux protocoles de couche réseau situés en amont de cette couche tout en communiquant avec les différentes technologies utilisées pour véhiculer les informations entre la source et la destination.

Le rôle de cette sous-couche est de réceptionner le paquet IP et d'y ajouter les informations de contrôle pour en faciliter l'acheminement jusqu'à la destination. Elle ajoute 2 éléments d'adressage décrit dans la spécification LLC 802.2

La norme IEEE 802.2 définit un certain nombre de champs dans les trames, lesquelles permettent à plusieurs protocoles de couche supérieure de partager une liaison de données physique.

3.2.2. La sous-couche MAC

La sous-couche MAC concerne les protocoles que doit suivre un hôte pour accéder au média. Dans un environnement de média partagé, il permet de déterminer quel ordinateur peut parler. On distingue 2 types de protocoles MAC :

- o Déterministes : chacun son tour
 - o Exemple : Token Ring
- o Non déterministe : premier arrivé premier servi
 - o Exemple : Ethernet

3.3. Notions de base de la technologie Token Ring

Token Ring, mis en place par IBM, a commencé à se développer au début des années 70. C'est aujourd'hui le deuxième type de réseau derrière Ethernet.

Il existe 2 variantes ; ces dernières étant quasi identique : Token Ring IBM et IEEE 802.5 :

	Token Ring IBM	IEEE 802,5
débits :	4 ou 16 Mbits/s	4 ou 16 Mbits/s
stations/ segments	260 (câble à paire torsadée blindée) 72 (câble à paire torsadée non blindée)	250
Topologie	En étoile	non spécifié
Média	Paire torsadée	Non spécifié
Signalisation	Bande de base	Bande de base
Méthode d'accès	Passage de jeton	Passage de jeton
Codage	Manchester différentiel	Manchester différentiel

Figure 4.3 : Caractéristiques des 2 variantes de Token Ring

3.3.1. Principe du MAC Token Ring : le passage de jeton

La topologie physique de Token Ring est en étoile, sa topologie logique en anneau. Dans cet anneau, une petite trame ; le jeton ; circule. Toutes les stations le reçoivent tour à tour. Si une station n'a rien à émettre, elle se contente de récupérer le jeton et de le transmettre à son voisin. Si par contre elle désire émettre des données sur le réseau, elle saisit le jeton, en altère un bit pour en faire son début de trame, puis y ajoute les informations à transmettre avant de transmettre cela à son voisin. Pendant ce temps, aucun jeton ne circule sur le réseau.

La trame circule autour de l'anneau jusqu'au destinataire, qui réceptionne la trame, la copie afin de la traiter puis la remet sur le réseau qu'elle parcourt jusqu'à l'émetteur. Ce dernier s'assure que le destinataire a bien reçu la trame, puis peut soit émettre une nouvelle trame soit ; s'il n'a plus rien à émettre ; remettre le jeton sur le réseau.

Ce principe comporte 2 avantages : il est exempt de toute collision et permet un accès déterministe au média grâce au système de priorité.

3.3.2. Système de priorité d'accès au média

Chaque hôte dispose d'une priorité d'accès au média. Il existe donc une hiérarchie concernant le droit d'accès au média.

Si un hôte A disposant d'une priorité plus haute ; que celui en train d'émettre désire communiquer ; il inscrit son niveau de priorité dans le champ réservation de la trame. Lorsque l'hôte B émetteur récupère sa trame, il inspecte le champ réservation et ; si celui-ci est plus élevé que sa propre priorité ; arrête immédiatement d'émettre et remet le jeton sur le réseau avec pour priorité la valeur du champ réservation.

3.4. Notions de base de l'interface FDDI (Fiber Distributed Data Interface)

3.4.1. Principe d'accès au média de la technologie FDDI

La technologie FDDI utilise le même principe d'accès au média que Token Ring ; à savoir un accès déterministe de type « passage de jeton » ; en ajoutant à celle-ci un second anneau permettant d'assurer un transit fiable si une panne survient sur le premier. Les deux anneaux circulent dans des directions opposées.

3.4.2. Signalisation et médias FDDI

La méthode de codage utilisée par la technologie FDDI est le codage 4B/5B. Les sources de signaux sont des LED ou des Lasers.

Le média utilisé est la fibre optique pour les raisons suivantes :

- aucune émission de signaux électriques pouvant être surveillés
- aucune sensibilité par rapport aux interférences électriques
- débit plus élevé que les médias à base de cuivres

On distingue 2 types de fibre optique :

- monomode : un seul faisceau parcourt la fibre, les lasers sont utilisés comme émetteurs récepteurs.



Figure 4.7 : la fibre monomode

- Multimode : plusieurs faisceaux parcourent la fibre avec des angles différents, selon leur angle de pénétration. Dans le cas d'une fibre multimode, les émetteurs récepteurs utilisés sont des LED.



Figure 4.8 : la fibre multimode

Mode : faisceau de rayons lumineux pénétrant dans la fibre avec un angle particulier.

Il est important de noter que la fibre monomode permet une bande passante plus importante ainsi qu'un parcourt de câbles plus longs (3000m pour du monomode et 2000 pour du multimode). Elle est généralement utilisée pour la connectivité entre 2 immeubles alors que la fibre multimode est réservée à la connectivité au sein d'un bâtiment.

3.5. Notions de base d'Ethernet et d'IEEE 802.3

Conçu à Hawaï dans les années, Ethernet est la technologie la plus répandue dans les réseaux actuels. Au début des années 80 fut mis en place par l'IEEE la norme IEEE 802.3 à partir d'Ethernet.

Ethernet et IEEE 802.3 définissent des technologies semblables :

- Utilisation de CSMA/CD (cf 4.5.2) pour l'accès au média
- Concept de réseaux de broadcast

Il existe cependant quelques différences subtiles, en effet, Ethernet offre des services correspondant aux couches 1 et 2 du modèle OSI alors que IEEE 802.3 définit la couche 1 ainsi que la partie MAC de la couche 2

3.5.1. Structure de trame Ethernet

Trame Ethernet

?	1	6	6	2	46-1500	4
Préambule	Délimiteur de début de trame	Adresse de destination	Adresse d'origine	Type	Données	FCS

Trame IEEE 802.3

?	1	6	6	2	64-1500	4
Préambule	Délimiteur de début de trame	Adresse de destination	Adresse d'origine	Longueur	Données	FCS

Figure 4.9 : Structure de trames Ethernet et IEEE 802.3

- Préambule : composé de 1 et de 0 en alternance, annonce si la trame est de type Ethernet ou 802.3.
- Début de trame : IEEE 802.3 : l'octet séparateur se termine par 2 bits 1 consécutifs servant à synchroniser les portions de réception des trames de toutes les stations.
- Champ d'adresse d'origine : toujours de type unicast
- Champ d'adresse de destination : peut être de type unicast, multicast ou broadcast
- Type (Ethernet) : précise le type de protocole de couche supérieure qui reçoit les données
- Longueur (802.3) : indique le nombre d'octets de données qui suit le champ.
- Données :
 - Ethernet : une fois le traitement de couche 1 et 2 terminé, les données sont transmises au protocole de la couche supérieure indiqué dans le champ type. On peut avoir recours à des octets de remplissage s'il n'y a pas assez de données pour remplir les 64 octets minimaux de la trame
 - IEEE 802.3 : une fois le traitement de couche 1 et 2 terminé, les données sont transmises au protocole de la couche supérieure indiqué dans le champ donnée de la trame on peut aussi ici avoir recours à du remplissage
- FCS : Séquence de contrôle de trame. Cette séquence contient un code de redondance cyclique de 4 octets permettant à l'unité réceptrice de vérifier l'intégrité des données.

3.5.2. MAC Ethernet

Ethernet et 802.3 utilisent un principe d'accès au média non déterministe : CSMA/CD (Carrier Sense Multiple Access / Collision Detect)

Les hôtes se partagent le média, si l'un d'eux désire émettre, il vérifie au préalable que personne n'est en train de le faire, puis commence à émettre (CSMA).

Si cependant 2 hôtes émettent en même temps, il se produit alors une collision. La première station qui détecte une collision envoie alors un signal de bourrage, se traduisant par un arrêt d'émission de tous les hôtes. Les paquets concernés sont alors détruits.

Chaque hôte calcule alors une valeur aléatoire définissant la durée avant de recommencer à émettre, puis le mécanisme de CSMA se remet en fonction.

3.5.3. Signalisation et médias Ethernet

Ethernet utilise un codage de type Manchester.

Il existe actuellement de nombreuses variantes d'Ethernet, la figure ci-dessous nous présente les différents médias et topologie utilisée en fonction du type utilisé :

Type :	Média	Bande passante maximale	longueur de segment maximale	topologie physique	topologie logique
10BASE5	Coaxial Épais	10 Mbits/s	500m	Bus	Bus
10BASE-T	UTP CAT 5	10 Mbits/s	100m	Étoile/ Étoile étendue	Bus
10BASE-FL	Fibre optique Multimode	10 Mbits/s	2000m	Étoile	Bus
100BASE-TX	UTP CAT 5	100 Mbits/s	100m	Étoile + étoile étendue	Bus
100BASE-FX	Fibre optique Multimode	100 Mbits/s	2000m	Étoile	Bus
1000BASE-TX	UTP CAT 5	1 000 Mbits/s	100m	Étoile + étoile étendue	Bus

Figure 4.10 : les différents types d'Ethernet

3.6. Les équipements de couche 2

3.6.1. Les cartes réseau ou NIC

Au niveau de la couche liaison de données, la carte réseau assure le contrôle de lien logique, la désignation ; le verrouillage de trame, l'accès au média ainsi que la signalisation



Figure 4.11 : Cartes réseau PCMCIA et ISA

3.6.2. Les ponts

Les ponts servent à relier des segments réseau en permettant une segmentation des domaines de collisions. Une analyse des trames est effectuée afin d'acheminer les trames sur le bon segment réseau en fonction des adresses MAC.

Il permet de plus de connecter différentes technologies de couche 1 et cumule à ses fonctions celle du répéteur.



Figure 4.12 : pont Ethernet

3.6.3. Les commutateurs

Le commutateur est un pont multi ports. Il permet donc de relier plusieurs segments réseau et d'acheminer les trames sur le bon segment de destination grâce aux informations de couche 2



Figure 4.13 : commutateur Ethernet

4. Couche 3 : La couche réseau

Le rôle de la couche réseau est d'acheminer les données entre l'émetteur et le destinataire au travers de différents réseaux en mettant en place un système d'adressage hiérarchique pour combiner aux manques de l'adressage MAC

Les protocoles de la couche réseau utilisent un système d'adressage garantissant l'unicité des adresses sur le réseau et définissant une méthode d'acheminement des informations entre les réseaux.

4.1. Principe de sélection du chemin

4.1.1. La sélection du chemin

Les méthodes de sélection du chemin permettent aux équipements de couche 3 ; les routeurs ; de déterminer la route à suivre pour acheminer les informations au travers de différents réseaux. Les services de routage utilisent les informations de topologie du réseau pour évaluer les chemins. Ce processus est aussi appelé routage des paquets et prend en compte divers paramètres ; ou métriques ; comme :

- o Densité du trafic
- o Nombre de routeurs à franchir pour joindre la destination
- o Vitesse des liaisons
- o Etc....

4.1.2. L'adressage de la couche réseau

Sur Internet, les ordinateurs communiquent entre eux grâce au protocole TCP/IP qui utilise des numéros de 32 bits, que l'on écrit sous forme de 4 numéros allant de 0 à 255 (4 fois 8 bits), on les note donc sous la forme xxx.xxx.xxx.xxx où chaque xxx représente un entier de 0 à 255.

Ces numéros servent aux ordinateurs du réseau pour se reconnaître, ainsi il ne doit pas exister deux ordinateurs sur le réseau ayant la même adresse IP.

4.1.3. Protocoles routables, non routables

Un protocole routable est un protocole pouvant être acheminé au travers de différents réseaux.

Exemple :

- o IP
- o IPX
- o Appletalk

Par opposition, un protocole non routable ne peut être routé

Exemple :

- o Netbeui

4.2. Les équipements de couche 3 : les routeurs

Routeur : équipement de couche 3 permettant d'interconnecter 2 réseaux ou plus en se basant sur les adresses de couche 3. Le routeur permet également une segmentation des domaines de broadcasts



Figure 5.4 : Routeur Cisco 2600

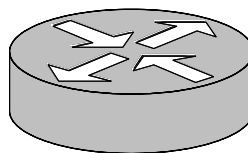


Figure 5.5 : Symbole logique de routeur

4.3. Les communications de réseau à réseau

4.3.1. Le protocole ARP

Le protocole ARP a un rôle phare parmi les protocoles de la couche Internet de la suite TCP/IP, car il permet de connaître l'adresse physique d'une carte réseau correspondant à une adresse IP.

Chaque machine connectée au réseau possède un numéro d'identification de 48 bits. Ce numéro est un numéro unique qui est fixé dès la fabrication de la carte en usine. Toutefois, la communication sur Internet ne se fait pas directement à partir de ce numéro mais à partir de l'adresse IP.

Ainsi, pour faire correspondre les adresses physiques aux adresses logiques, le protocole ARP interroge les machines du réseau pour connaître leur adresse physique, puis crée une table de correspondance entre les adresses logiques et les adresses physiques dans une mémoire cache.

Lorsqu'une machine doit communiquer avec une autre, elle consulte la table de correspondance. Si jamais l'adresse demandée ne se trouve pas dans la table, le protocole ARP émet une requête sur le réseau. L'ensemble des machines du réseau va comparer cette adresse logique à la leur.

Si l'une d'entre-elles s'identifie à cette adresse, la machine va répondre à ARP qui va stocker le couple d'adresses dans la table de correspondance et la communication va alors pouvoir avoir lieu...

4.3.2. Le protocole ICMP

Le protocole ICMP (Internet Control Message Protocol) est un protocole qui permet de gérer les informations relatives aux erreurs aux machines connectées. Etant donné le peu de contrôles que le protocole IP réalise il permet non pas de corriger ces erreurs mais de faire part de ces erreurs aux protocoles des couches voisines. Ainsi, le protocole ICMP est utilisé par tous les routeurs, qui l'utilisent pour reporter une erreur (appelé Delivery Problem).

Toutefois en cas d'erreur sur un datagramme transportant un message ICMP, aucun message d'erreur n'est délivré pour éviter un effet "boule de neige " en cas d'incident sur le réseau.

0	8	16	24	31
Type	Code	Message		

Figure 5.8 : Structure d'un paquet ICMP

4.4. Les services réseau de la couche 3

4.4.1. Les services réseau non orientés connexion (commutation de paquets)

La plupart des services réseaux utilisent de livraison non orientée connexion. Ils traitent chaque paquet séparément. Il se peut que les paquets empruntent des chemins différents et sont rassemblés lorsqu'ils arrivent à destination.

Dans un système non orienté connexion le destinataire n'est pas contacté avant la réception des paquets, comme c'est le cas par exemple pour les services postaux.

Internet est un immense réseau non orienté connexion au sein duquel le protocole IP transporte les paquets. Le protocole TCP (couche 4) y ajoute des services orientés connexion au dessus du protocole IP afin d'assurer une distribution fiable des données.

4.4.2. Les services réseau orientés connexion (commutation de circuits)

Une connexion est établie entre l'émetteur et le destinataire avant le transfert des données. Un exemple de ce système est le système téléphonique.

Tous les paquets sont donc acheminés dans le même circuit physique ou ; plus souvent ; dans le même circuit virtuel.

5. Couche 4 : La couche Transport

5.1. La couche transport

5.1.1. Fonction de la couche transport

Nous avons vu dans les chapitres précédents comment TCP/IP envoie les informations de l'émetteur au destinataire. La couche transport ajoute à ce mécanisme la notion de « qualité de service », à savoir la garantie d'un acheminement fiable des informations au travers du réseau.

5.1.2. Les protocoles de couche 4

La couche 4 du modèle OSI comprend 2 protocoles : TCP et UDP

TCP est un protocole orienté connexion, c'est-à-dire qu'il associe aux transport des informations la notion de qualité en offrant les services suivants :

- fiabilité
- division des messages sortants en segments
- ré assemblage des messages au niveau du destinataire
- ré envoi de toute donnée non reçue

UDP est lui un protocole non orienté connexion, c'est-à-dire qu'il n'offre pas de fonction de contrôle du bon acheminement :

- aucune vérification logicielle de la livraison des messages
- pas de réassemblage des messages entrants
- pas d'accusé de réception
- aucun contrôle de flux

Cependant, UDP offre l'avantage de nécessiter moins de bande passante que TCP. Il peut donc être intéressant d'utiliser ce protocole pour l'envoi de messages ne nécessitant pas de contrôle de qualité.

5.1.3. TCP comme complément d'IP

A IP qui offre un service sans connexion de couche 3 permettant l'acheminement des données au sein d'un réseau s'ajoute TCP ; un protocole de couche 4 ; qui ajoute les capacités de contrôle de flux et de fiabilité de transmission.

Pour faire une analogie avec le système postal, IP serait un exemple d'envoi de courrier ordinaire auquel TCP ajoute le service d'envoi recommandé, garantissant à l'émetteur la remise de la lettre.

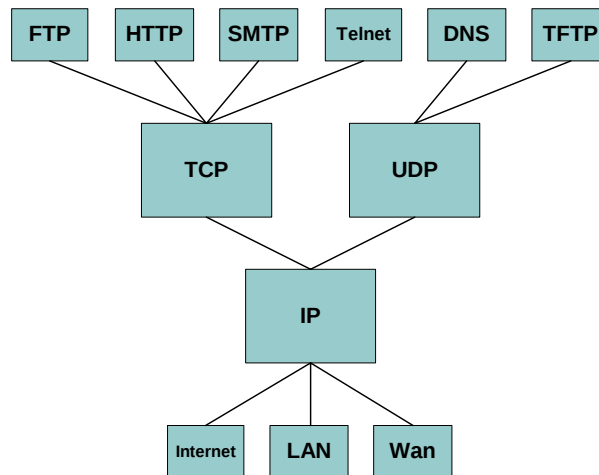


Figure 6.2 : Schéma de protocoles

5.2. TCP et UDP

5.2.1. Les numéros de ports

Afin que plusieurs communications puissent circuler en même temps, TCP et UDP utilisent des numéros de ports. Des conventions ont été établies pour des applications :

Protocole	n° de port	Description
FTP data	20	File Transfer [données par défaut]
FTP	21	File Transfer [contrôle]
ssh	22	SSH
Telnet	23	Telnet
smtp	25	Simple Mail Transfer
domain	53	Domain Name Server
HTTP	80	World Wide Web HTTP
pop3	110	Post Office Protocol - Version 3
sftp	115	Simple File Transfer Protocol
sqlserv	118	SQL Services
nnntp	119	Network News Transfer Protocol
imap2	143	Interactive Mail Access Protocol v2
news	144	NewS
ipx	213	IPX
netware IP	396	Novell Netware sur IP
https	443	Protocole HTTP sécurisé

Figure 6.3 : numéros de ports

Toute application n'ayant pas un numéro de port défini et reconnu se voit attribué un numéro de port aléatoire. Les ports ont été attribués de la manière suivante :

- o 0 → 255 réservés aux applications publiques
- o 255 → 1023 attribué aux entreprises pour les applications commerciales
- o 1023 → + utilisés pour les attributions dynamiques

5.3. Les méthodes de connexion TCP

Un service orienté connexion comporte 3 points importants :

- o Un chemin unique entre les unités d'origine et de destination est déterminé
- o Les données sont transmises de manière séquentielle et arrivent à destination dans l'ordre
- o La connexion est fermée lorsqu'elle n'est plus nécessaire

5.3.1. Connexion ouverte/échange à trois

Les hôtes TCP établissent une connexion en 3 étapes, appelé aussi « connexion ouverte » :

- L'émetteur envoie un paquet avec un numéro de séquence initial (x) avec un bit dans l'en-tête pour indiquer une demande de connexion.

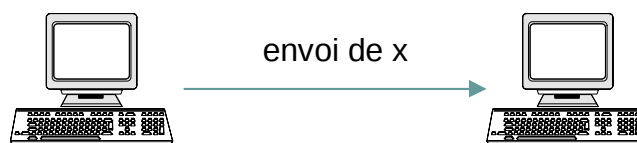


Figure 6.6 : étape n° 1 de la connexion TCP

- Le destinataire le reçoit, consigne le numéro de séquence initial, répond par un accusé de réception «x+1 » et inclut son propre n° de séquence (y).

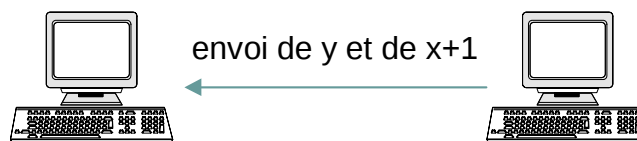


Figure 6.7 : étape n° 2 de la connexion TCP

- L'émetteur reçoit x+1 et renvoie y+1 pour dire au destinataire que la réception s'est bien passée.

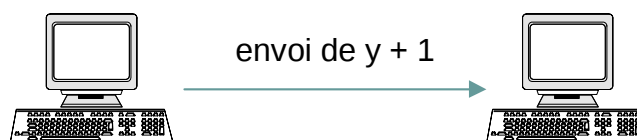


Figure 6.8 : étape n° 3 de la connexion TCP

Quand l'émetteur reçoit x+1, cela signifie que le destinataire a bien reçu tout les paquets ayant pour n° de séquence x et moins et attend la suite.

Il existe également des méthodes garantissant la fiabilité des protocoles.

5.3.2. Positive Acknowledgement Retransmission

La technique Positive Acknowledgement Retransmission ; ou PAR ; consiste à envoyer un paquet, démarrer un compteur puis attendre un accusé de réception avant d'envoyer le suivant.

Si le compteur arrive à expiration avant l'arrivée de l'accusé, les informations sont alors retransmises et un nouveau compteur est déclenché.

Cependant, cette technique est consommatrice de bande passante ; c'est alors qu'intervient le mécanisme de fenêtrage.

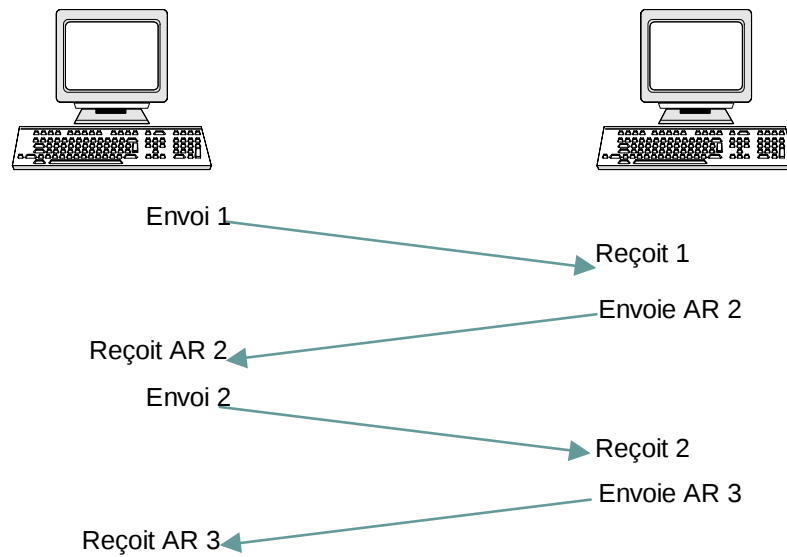


Figure 6.9 : Principe de PAR

5.3.3. Le Fenêtrage

Le Fenêtrage est un mécanisme dans lequel le récepteur envoie un accusé de réception après avoir reçu un certain nombre de données. Si le destinataire n'envoie pas d'accusé, cela signifie pour l'émetteur que les informations ne sont pas parvenues correctement et dans ce cas sont retransmises.

La taille de la fenêtre détermine la quantité de données que l'on peut transmettre avant de recevoir un accusé de réception.

TCP utilise un système d'accusé de réception prévisionnel, ce qui signifie que le numéro d'accusé renvoyé indique la prochaine séquence attendue

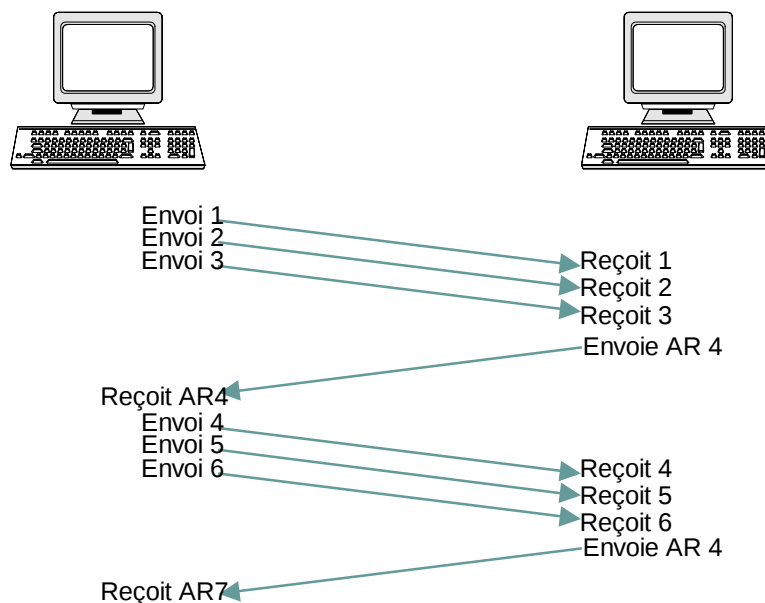


Figure 6.10 : échange TCP fenêtrée avec une fenêtre de 3

6. Couche 5 : La couche Session

Une session est un ensemble de transactions entre deux unités réseau ou plus.

Une analogie pour comprendre la couche session est une communication entre plusieurs individus. Si l'on souhaite que la conversation se déroule correctement, il est impératif de mettre en place diverses règles, afin que les interlocuteurs ne s'interrompent pas par exemple.

Cette notion de contrôle du dialogue est le point essentiel de la couche session.

Le rôle de la couche session est d'ouvrir, gérer et fermer les sessions entre les applications. Cela signifie que c'est elle qui prend en compte :

- le lancement des sessions
- la resynchronisation du dialogue
- l'arrêt des sessions

Elle coordonne donc les applications qui communiquent au travers des hôtes.

Une communication entre ordinateurs suppose de nombreuses conversations courtes (commutation de paquets) avec en plus de cela d'autres communications pour s'assurer de l'efficacité de la communication.

Ces conversations nécessitent que les hôtes jouent à tour de rôles celui de client (demandeur de services) et de serveur (fournisseur de services).

6.1. Le contrôle du dialogue

La couche session décide si la conversation sera de type bidirectionnel simultané ou alterné. Cette décision relève du contrôle du dialogue.

6.2. La synchronisation du dialogue

Cette étape est des plus importante, elle permet aux hôtes communicants dans marquer une pause pour par exemple sauvegarder la communication en cours et re synchroniser le dialogue.

6.3. La division du dialogue

La division du dialogue englobe le lancement, la fin et la gestion ordonnés de la communication.

7. Couche 6 : La couche présentation

7.1. Fonction et normes de la couche présentation

L'un des rôles de la couche présentation est de présenter les données dans un format que le dispositif récepteur est capable de comprendre. La couche présentation joue donc un rôle d'interprète entre les unités qui doivent communiquer par le biais d'un réseau.

La couche 6, la couche présentation, assure trois fonctions principales, à savoir :

- o Le formatage des données (présentation)
- o Le cryptage des données
- o La compression des données

7.2. Cryptage et compression des données

La couche 6 est également responsable du cryptage et de la compression des données.

7.2.1. Le cryptage des données

Le cryptage permet de protéger la confidentialité des informations pendant leur transmission.

Exemple : Les transactions financières, surtout celles qui sont faites avec des cartes de crédit, doivent être cryptées afin de protéger les données sensibles transmises sur Internet.

Une clé de cryptage est utilisée pour crypter les données à la source et pour les décrypter à destination. Un algorithme est donc utilisé pour rendre ces données incompréhensible à quiconque ne disposant pas de la clé.

7.2.2. La compression des données

La couche présentation assure également la compression des fichiers.

La compression applique des algorithmes (formules mathématiques complexes) pour réduire la taille des fichiers. L'algorithme cherche certaines séquences de bits répétitives dans les fichiers et les remplace par un « jeton ».

Le jeton est une séquence de bits raccourcie qui est substituée à la séquence complète.

Exemple : Remplacer « "LaboratoireCisco »" par « Lab »"

8. Couche 7 : La couche application

Le rôle de cette couche est d'interagir avec les applications logicielles. Elle fournit donc des services au module de communication des applications en assurant :

- o L'identification et la vérification de la disponibilité des partenaires de communication voulus
- o La synchronisation des applications qui doivent coopérer
- o L'entente mutuelle sur les procédures de correction d'erreur
- o Le contrôle de l'intégrité des données

Dans le modèle OSI, la couche application est la plus proche du système terminal.

Celle-ci détermine si les ressources nécessaires à la communication entre systèmes sont disponibles. Sans la couche application, il n'y aurait aucun support des communications réseau. Elle ne fournit pas de services aux autres couches du modèle OSI, mais elle collabore avec les processus applicatifs situés en dehors du modèle OSI

Ces processus applicatifs peuvent être des tableurs, des traitements de texte, des logiciels de terminaux bancaires, etc.

De plus, la couche application crée une interface directe avec le reste du modèle OSI par le biais d'applications réseau (navigateur Web, messagerie électronique, protocole FTP, Telnet, etc.) ou une interface indirecte, par le biais d'applications autonomes (comme les traitements de texte, les logiciels de présentation ou les tableurs), avec des logiciels de redirection réseau.

8.1. Principes de la couche application

8.1.1. Les applications réseau directes

La plupart des applications exécutées dans un environnement réseau sont de type client-serveur. Ces applications (logiciels FTP, navigateurs Web ou applications de messagerie électronique) se composent de deux modules, l'un jouant le rôle du client et l'autre, le rôle du serveur.

- o Le module client tourne sur l'ordinateur local : c'est le «demandeur de services ».
- o Le module serveur tourne sur un ordinateur distant et fournit des services en réponse aux demandes du client.

Une application client-serveur répète constamment la boucle d'itération suivante :

- o demande du client
- o réponse du serveur

Ainsi, un navigateur accède à une page Web en envoyant une demande d'adresse Web (URL) à un serveur Web distant. Après avoir localisé la page grâce à l'adresse URL fournie, le serveur Web associé à l'adresse répond à la demande. Ensuite, en fonction des informations reçues du serveur Web, le client pourra demander des pages supplémentaires du même serveur Web ou accéder à une autre page associée à un serveur Web différent.

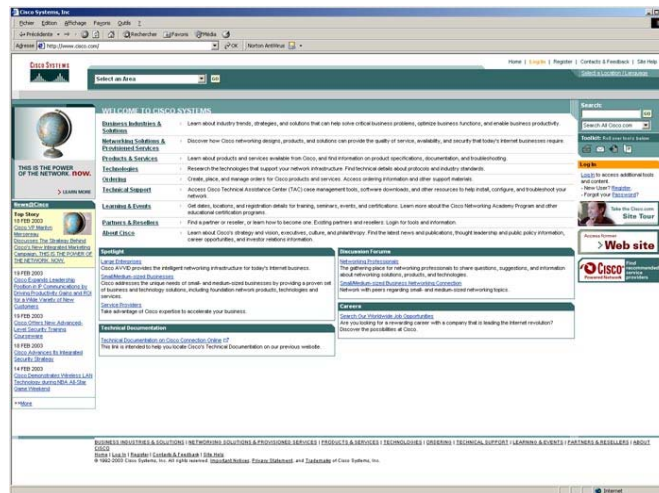


Figure 9.1 : Le Navigateur Microsoft Internet Explorer

8.2. Le protocole DNS

8.2.1. Présentation du protocole DNS

Chaque station possède une adresse IP propre. Cependant, les utilisateurs ne veulent pas travailler avec des adresses IP, mais avec des noms de stations ou des adresses plus explicites comme par exemple <http://www.labo-cisco.com>

Pour répondre à cela, le protocole DNS permet d'associer des noms en langage courant aux adresses numériques.

Résolution de noms de domaines : Corrélation entre les adresses IP et le nom de domaine associé.

8.2.2. Les noms d'hôtes et le « domain name system »

Ce système consiste en une hiérarchie de noms permettant de garantir l'unicité d'un nom dans une structure arborescente.

On appelle nom de domaine, le nom à deux composantes, dont la première est un nom correspondant au nom de l'organisation ou de l'entreprise, le second à la classification de domaine (.fr, .com, ...). Chaque machine d'un domaine est appelée hôte. Le nom d'hôte qui lui est attribué doit être unique dans le domaine considéré (le serveur Web d'un domaine porte généralement le nom WWW).

L'ensemble constitué du nom d'hôte, d'un point, puis du nom de domaine est appelé adresse FQDN (Fully Qualified Domain, soit Domaine Totalement Qualifié). Cette adresse permet de repérer de façon unique une machine. Ainsi, www.cisco.com représente une adresse FQDN.

9. La communication LAN

9.1. CSMA/CD

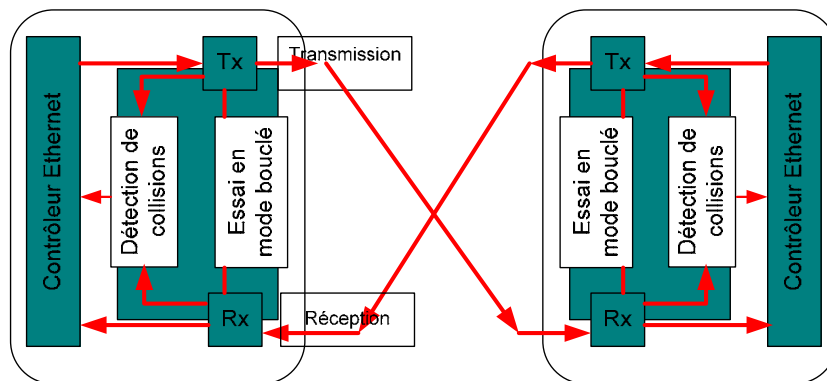
CSMA/CD : *Carrier Sense Multiple Access / Collision Detect*

CSMA/CD est une méthode d'accès au média non déterministe. Le schéma pour l'accès au média est le suivant :

- Les hôtes se partagent un même média ; si l'un d'eux souhaite émettre ; il vérifie au préalable que personne n'est déjà en train de le faire, puis commence à émettre.
- Si cependant 2 hôtes émettent en même temps, il se produit alors une collision :
 - o La première station qui détecte une collision (une modification notable d'amplitude du signal) envoie alors un signal de bourrage, se traduisant par un arrêt d'émission de tous les hôtes.
 - o Les paquets concernés sont alors détruits bits à bits
- Chaque hôte calcule alors une valeur aléatoire définissant la durée avant de recommencer à émettre, puis le mécanisme se remet en fonction.

9.1.1. Les différents types d'Ethernet

- Ethernet Half Duplex :
 - o Transmission (Tx) et réception (Rx) alternée
 - o Utilisation de 50 à 60% de la bande passante en raison des collisions et de la latence
 - o Détection des collisions



Ce mode n'autorise pas l'émission et la réception simultanée. Il n'utilise que 50 à 60% de la bande passante disponible à cause de la latence et des collisions

- Ethernet Full Duplex :
 - o Permet l'émission et la réception simultanée
 - o Nécessite l'utilisation d'un câble contenant 2 paires de fils et d'une connexion commutée entre les 2.
 - o Lors d'une communication, une communication point à point sans collision est créée
 - o Utilisation de 100% de la bande passante.

Note : la topologie physique décrit la disposition physique des éléments alors que la topologie logique définit la méthode de communication des équipements.

Laboratoire Supinfo des Technologies Cisco

Site : www.labo-cisco.com - Mail : labo-cisco@supinfo.com

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

9.2. La commutation LAN

Les réseaux Ethernet sont sujets à divers problèmes affectant les performances du réseau ; telles que :

- Les collisions
- La latence des équipements réseaux
- La remise de données de type broadcast

Afin d'optimiser les performances du réseau, la segmentation est nécessaire.

9.2.1. La segmentation LAN

Le but de la segmentation Lan est d'obtenir une réduction de la taille des domaines de collision afin d'économiser la bande passante disponible.

Il est possible de recourir à trois types de segmentation des domaines de collisions :

- Segmentation par pont :
 - o Segmentation du domaine de collision en 2 grâce au pont, dispositif de couche 2 permettant un filtrage des trames en fonction des adresses MAC des hôtes.
- Segmentation par routeurs :
 - o Segmentation du domaine de broadcast en fonction des adresses réseau de couche 3.
- Segmentation par commutateur :
 - o Segmentation du domaine de collision par la mise en place de chemins commutés entre l'hôte et le destinataire (micro segmentation)

Commuter : *transmettre des données entre 2 ou plusieurs interfaces.*

9.2.2. Fonctionnement d'un commutateur

Le commutateur est un pont multi ports. Il permet de relier plusieurs segments réseau et d'acheminer les trames sur le bon segment de destination grâce aux informations de couche 2.

Un environnement commuté présente les avantages suivants :

- Réduction du nombre de collisions
- Plusieurs communications simultanées
- Liaisons montantes haut débit
- Amélioration de la réponse du réseau
- Hausse de la productivité de l'utilisateur

Les décisions d'acheminements du commutateur sont basées sur les adresses MAC contenues dans les trames circulant sur le réseau :

9.2.3. Les différents types de commutation

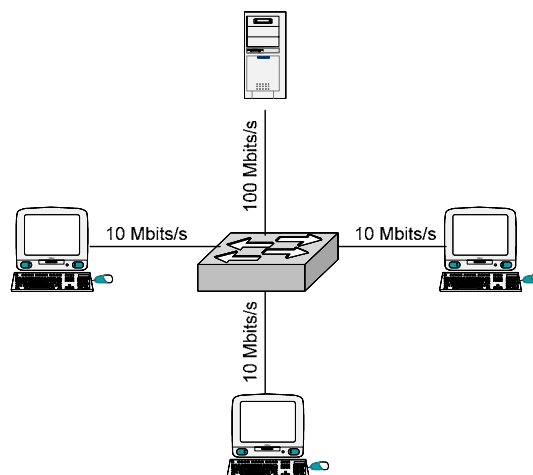
- *Store and forward* : Commutation où le commutateur attend d'avoir reçu toute la trame avant de la commuter. Cette méthode offre une grande vérification d'erreur. Cependant ce traitement augmente la latence réseau
- *Cut Through* : Dès que l'adresse de destination est connue, la trame commence à être commutée. Ce mode est plus rapide que le précédent. Il existe différentes variantes de ce type de commutation :
 - o *Fragment Free* : Filtrage des fragments de collision (inférieur à 64 octets)
 - Les paquets inférieurs à 64 octets sont considérés comme des fragments de collision
 - o *Fast Forward* : Pas de vérification d'erreurs

7 octets	1 octet	6 octets	6 octets	2 octets	Max. 1500 octets	4 octets
Préambule	SFD	Adresse de destination	Adresse d'origine	Longueur	Données	FCS

Fast Forward Fragment Free Store And Forward

La commutation peut aussi être caractérisée en fonction de la bande passante attribuée à chaque port :

- *Commutation symétrique* : Les connexions commutées offrent la même bande passante à chaque port.
- *Commutation asymétrique* : Les connexions commutées offrent des bandes passantes différentes.



9.2.4. La mise en mémoire tampon

La mise en mémoire tampon est utilisée par le commutateur pour stocker et acheminer les paquets aux bons ports ou encore lorsque le port de destination est occupé. On distingue 2 types de méthode :

- *Axée sur les ports* : les paquets sont placés dans des files d'attente spécifiques. Un paquet est acheminé uniquement lorsque tous les autres devant lui ont été envoyés
- *Mémoire partagée* : Les paquets sont stockés dans une mémoire commune. La quantité de mémoire allouée à un port est déterminée par la quantité nécessitée par chaque port (allocation dynamique de mémoire)

10. Protocole Spanning-Tree

10.1. Introduction

Les topologies redondantes sont mises en place pour palier à des liaisons interrompues. En effet, plusieurs chemins peuvent permettre d'accéder au même lien.

Mais si ces chemins redondants ne sont pas correctement gérés, les trames peuvent boucler indéfiniment. Le protocole Spanning-Tree permet d'y remédier.

10.2. Théorie concernant Spanning-Tree

Les commutateurs implémentent le protocole **IEEE 802.1D Spanning-Tree**. Il apporte une réponse au problème de bouclage. Pour ce faire, **STP** (Spanning-Tree Protocol) empêche certains ports de transmettre en mettant les ports dans un état de blocage ou dans un état de transmission, afin qu'il n'y ait qu'un seul chemin possible entre deux segments de LAN.

Un port bloqué ne peut ni recevoir ni émettre et inversement en mode de transmission. En premier lieu, des **BPDUs** (**Bridge Protocol Data Unit**) sont envoyés toutes les 2 secondes sur tous les ports.

Le commutateur qui détient l'identifiant de pont le plus bas (Bridge ID) est élu racine. Le Bridge ID de 8 octets est composé d'une priorité sur 2 octets (32768 par défaut), suivi par l'adresse MAC du port émetteur. Tous les ports du commutateur racine sont placés en état de transmission par le protocole STP.

Le commutateur racine transmet par tous ses ports des BPDUs. Ces messages sont transmis par les commutateurs non racine. A chaque réception de BPDUs, le champ du coût est incrémenté, ce qui permet aux commutateurs non racine de connaître la valeur de l'itinéraire jusqu'à la racine.

Le port de chaque commutateur qui reçoit le BPDUs comportant le coût le plus bas (donc le plus proche du commutateur racine) est élu port racine pour le segment de LAN auquel il est connecté.

Le calcul de la route se base sur la vitesse. Plus elle est grande, plus le coût est bas. Le port par lequel arrivent les BPDUs portant le moindre coût vers la racine est mis en état de transmission. Les autres ports sont mis en état de blocage, pour éliminer toute route redondante et ainsi éviter qu'il y ait des boucles actives.

Les ports prennent d'autres états. Voici un tableau récapitulatif des états appliqués aux ports :

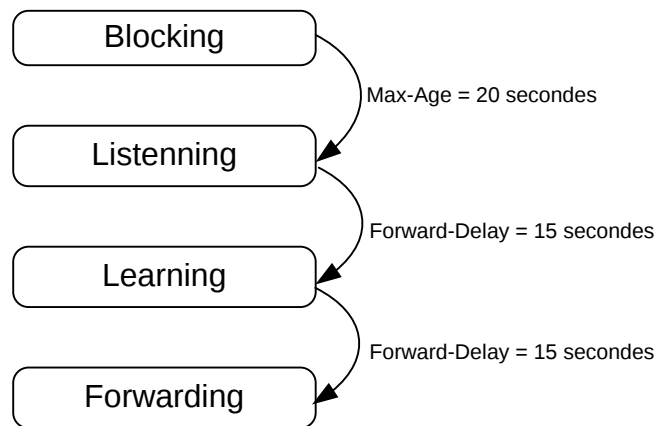
Etat	Description
Transmission	Le port émet et reçoit les trames.
Ecoute	Le port écoute les BPDUs pour s'assurer qu'il n'y ait pas de boucle. Ce processus a une durée de vie de 15 secondes.
Apprentissage	Le port écoute les BPDUs pour découvrir les adresses MAC. Ce processus a une durée de vie de 15 secondes également.
Désactivé	Le port n'est pas utilisé pour des raisons administratives.
Blocage	Le port ne peut ni émettre ni recevoir les trames.

Pour passer d'un état vers un autre, le protocole Spanning-Tree à définit 2 compteurs de temps :

1. Le compteur « Max-Age » d'une durée de 20 secondes.
2. Le compteur « forward-delay » d'une durée de 15 secondes.

Un réseau interconnecté est dit convergent lorsque tous les ports ont pris un état de blocage ou de transmission. Lorsqu'une modification topologique est détectée l'arbre est recalculé et le trafic ne reprend totalement qu'une fois la convergence atteinte.

Le schéma ci-dessous présente les différentes étapes d'une convergence Spanning-Tree:



10.3. Théorie concernant Rapid Spanning-Tree

Le protocole RSTP (Rapid Spanning Tree Protocol) est défini par le standard IEEE 802.1w. Il diffère principalement de STP de part sa convergence plus rapide. En effet, RSTP offre une convergence au minimum 5 fois plus rapide que STP. RSTP prend moins de 10 secondes pour converger.

RSTP et STP partagent certaines similitudes:

- Election d'un commutateur racine suivant le même processus.
- Ils élisent le port racine des commutateurs non racine de la même manière.
- Ils élisent le port désigné pour un segment de LAN de la même façon.
- Ils placent tous les ports dans un état de blocage ou de transmission, à la différence que RSTP utilise l'appellation « Discarding » pour l'état de blocage.

RSTP définit aussi des types de liaisons et de bordures. Les liaisons sont les connections physique entre les commutateurs et les bordures les connections physiques entre un commutateur et un hôte ou un concentrateur.

On distingue:

- Les liaisons point-à-point, c'est-à-dire entre deux commutateurs.
- Les liaisons partagées, c'est-à-dire entre un et plusieurs commutateurs.
- Les bordures point-à-point, entre un hôte et un commutateur.
- Les bordures partagées, entre un concentrateur et un commutateur.

Les ports des liaisons point-à-point et des bordures point-à-point sont immédiatement placés dans l'état de transmission. Ce qui permet d'améliorer la vitesse de convergence des commutateurs.

10.4. Commandes et configuration de Spanning-Tree

- **spanning-tree {identifiant de vlan} root**
 - o Mode de configuration globale
 - o Spécifie le Root-Bridge (par VLAN)
- **spanning-tree {identifiant de vlan} [priority priorité]**
 - o Mode de configuration globale
 - o Change la priorité Spanning-Tree du commutateur (Défaut = 32768)
- **spanning-tree cost {coût}**
 - o Mode de configuration d'interface
 - o Permet de modifier le coût STP.
- **show spanning-tree**
 - o Affiche des informations détaillées sur le protocole STP en cours ainsi que l'état de chaque port.
- **show spanning-tree interface {interface}**
 - o Affiche les informations Spanning-tree du port spécifié.
- **show spanning-tree vlan {vlan id}**
 - o Affiche les informations Spanning-tree du VLAN spécifié.
- **debug spanning-tree**
 - o Affiche les informations de changement topologique STP.

11. Adressage IP et Subnetting

11.1. Principe de l'adressage IP

Comme nous l'avons vu une adresse IP est une adresse 32 bits notée sous forme de 4 nombres entiers séparés par des points. On distingue en fait deux parties dans l'adresse IP:

- o Une partie des nombres à gauche désigne le réseau (on l'appelle netID)
- o Les nombres de droite désignent les ordinateurs de ce réseau (on l'appelle host-ID)

0	8	16	24	31
192	168	12	17	
1 1 0 0 0 0 0 0	1 0 1 0 1 0 0 0	0 0 0 0 1 1 0 0	0 0 0 1 0 0 0 1	

Figure 5.1 : Adresse IP

Les adresses IP ne peuvent communiquer qu'avec des adresses ayant le même numéro de réseau, y compris si des stations se trouvent sur le même segment. C'est ce même numéro qui permet au routeur d'acheminer le paquet au destinataire.

11.1.1. Les classes d'adresses IP

Actuellement l'organisme chargé d'attribuer les adresses IP est l'Internic.

Internic : Internet Network Information center

Les adresses IP sont réparties en plusieurs classes, en fonction des bits qui les composent :

Classe A : Dans une adresse IP de classe A, le premier octet représente le réseau. Le bit de poids fort (le premier bit, celui de gauche) est à zéro, ce qui signifie qu'il y a 27 (00000000 à 01111111) possibilités de réseaux, c'est-à-dire 128.
Toutefois, le réseau 0 (00000000) n'existe pas et le nombre 127 est réservé pour désigner votre machine, les réseaux disponibles en classe A sont donc les réseaux allant de 1.0.0.0 à 126.0.0.

Les trois octets de droite représentent les ordinateurs du réseau, le réseau peut donc contenir :
 $224-2 = 16777214$ ordinateurs.

Une adresse IP de classe A, en binaire, ressemble à ceci : 0 xxxxxxx xxxxxxxx xxxxxxxx
xxxxxxx

Classe B : Dans une adresse IP de classe B, les deux premiers octets représentent le réseau. Les deux premiers bits sont 1 et 0, ce qui signifie qu'il y a 214 (10 000000 00000000 à 10 111111 11111111) possibilités de réseaux, c'est-à-dire 16384.

Les réseaux disponibles en classe B sont donc les réseaux allant de 128.0.0.0 à 191.255.0.0

Les deux octets de droite représentent les ordinateurs du réseau, le réseau peut donc contenir :
 $216-21 = 65534$ ordinateurs.

Une adresse IP de classe B, en binaire, ressemble à ceci : 10 xxxxxx xxxxxxxx xxxxxxxx
xxxxxxx

Classe C : Dans une adresse IP de classe C, les trois premiers octets représentent le réseau. Les trois premiers bits sont 1,1 et 0, ce qui signifie qu'il y a 221 possibilités de réseaux, c'est-à-dire 2097152. Les réseaux disponibles en classe C sont donc les réseaux allant de 192.0.0.0 à 223.255.255.0

L'octet de droite représente les ordinateurs du réseau, le réseau peut donc contenir :
 $2^8 - 2^1 = 254$ ordinateurs.

Une adresse IP de classe C, en binaire, ressemble à ceci : 110 xxxxx xxxxxxxx xxxxxxxx
xxxxxxx

Il arrive fréquemment dans une entreprise qu'un seul ordinateur soit relié à Internet, c'est par son intermédiaire que les autres ordinateurs du réseau accèdent à Internet (on parle généralement de proxy).

Dans ce cas, seul l'ordinateur relié à Internet a besoin de réserver une adresse IP auprès de l'INTERNIC. Toutefois, les autres ordinateurs ont tout de même besoin d'une adresse IP pour pouvoir communiquer ensemble de façon interne.

Ainsi, l'INTERNIC a réservé une poignée d'adresses dans chaque classe pour permettre d'affecter une adresse IP aux ordinateurs d'un réseau local relié à Internet sans risquer de créer de conflits d'adresses IP sur le réseau. Il s'agit des adresses suivantes :

- o 10.0.0.1 à 10.255.255.254
- o 172.16.0.1 à 172.31.255.254
- o 192.168.0.1 à 192.168.255.254

11.1.2. Adresse de réseau et adresses de broadcast

Une adresse réseau est une adresse IP dont tous les bits hôtes sont occupés par des 0 binaires. Cette adresse désigne le réseau lui-même et non pas un hôte précis.

Exemple, dans un réseau de classe A, 113.0.0.0 désigne le réseau comprenant l'hôte 113.1.2.3.

L'adresse de broadcast est une adresse utilisée pour joindre en même temps tous les hôtes d'un réseau. Tous les bits hôtes de celle-ci sont à 1.

Exemple : pour le réseau 192.168.10.0, l'adresse de broadcast est 192.168.10.255

Broadcast : envoi de données à tous les hôtes d'un réseau

Ces adresses ne peuvent donc pas être utilisées pour identifier un hôte sur le réseau

11.2. Les sous-réseaux

Afin d'augmenter les capacités de gestion de trafic dans un réseau, il est possible de subdiviser ce dernier en plusieurs sous-réseaux afin de permettre une segmentation des domaines de broadcast.

Pour cela, on emprunte à la partie hôte des bits que l'on désigne comme champ de sous-réseaux. Le nombre minimal de bits à emprunter est de 2 et le nombre maximal est égal à tout nombre laissant 2 bits à la partie hôte.

11.2.1. Le masque de sous-réseau

Un masque de sous-réseau est une adresse de 32 bits contenant des 1 aux emplacements des bits que l'on désire conserver, et des 0 pour ceux que l'on veut rendre égaux à zéro. Une fois ce masque créé, il suffit de faire un ET entre la valeur que l'on désire masquer et le masque afin de garder intacte la partie que l'on désire et annuler le reste.

Ainsi, un masque réseau se présente sous la forme de 4 octets séparés par des points (comme une adresse IP), il comprend (dans sa notation binaire) des zéros au niveau des bits de l'adresse IP que l'on veut annuler (et des 1 au niveau de ceux que l'on désire conserver).

Ainsi, le réseau associé à l'adresse 34.56.123.12 est 34.0.0.0 (puisque'il s'agit d'une adresse de classe A).

Il suffit donc pour connaître l'adresse du réseau associé à l'adresse IP 34.56.123.12 d'appliquer un masque dont le premier octet ne comporte que des 1 (ce qui donne 255), puis des 0 sur les octets suivants (ce qui donne 0..).

*Le masque est : 11111111.00000000.00000000.00000000
Le masque associé à l'adresse IP 34.208.123.12 est donc 255.0.0.0.*

La valeur binaire de 34.208.123.12 est : 00100010.11010000.01111011.00001100

Un ET entre
00100010.11010000.01111011.00001100
ET
11111111.00000000.00000000.00000000
donne 00100010.00000000.00000000.00000000

C'est-à-dire 34.0.0.0, c'est bien le réseau associé à l'adresse 34.56.123.12

11.2.2. Création de sous-réseaux

Reprenons l'exemple du réseau 34.0.0.0, et supposons que l'on désire que les deux premiers bits du deuxième octet permettent de désigner le réseau.

Le masque à appliquer sera alors : 11111111.11000000.00000000.00000000 c'est-à-dire 255.192.0.0

Si on applique ce masque, à l'adresse 34.208.123.12 on obtient : 34.192.0.0

En réalité, il y a 4 cas de figures possibles pour le résultat du masquage d'une adresse IP d'un ordinateur du réseau 34.0.0.0

- o Soit les deux premiers bits du deuxième octet sont 00, auquel cas le résultat du masquage est 34.0.0.0
- o Soit les deux premiers bits du deuxième octet sont 01, auquel cas le résultat du masquage est 34.64.0.0
- o Soit les deux premiers bits du deuxième octet sont 10, auquel cas le résultat du masquage est 34.128.0.0
- o Soit les deux premiers bits du deuxième octet sont 11, auquel cas le résultat du masquage est 34.192.0.0

Ce masquage divise donc un réseau de classe A (pouvant admettre 16777214 ordinateurs) en 4 sous-réseaux pouvant admettre 222 ordinateurs, c'est-à-dire 4194304 ordinateurs.

Au passage on remarque que le nombre d'ordinateurs possibles dans les deux cas est au total de 16777214 ordinateurs ($4 \times 4194304 - 2 = 16777214$)

Le nombre de sous-réseaux dépend du nombre de bits que l'on attribue en plus au réseau (ici 2). Le nombre de sous-réseaux est donc :

Nombre de bits	Nombre de SR
1	2
2	4
3	8
4	16
5	32
6	64
7	128
8	256

Figure 5.3 : bits de masque de sous-réseaux

Lors d'une division d'une plage d'adresses en plusieurs sous-réseaux, il est important de constater que le premier sous-réseau est interdit, car son adresse réseau est la même que celle du réseau initial, de même que le dernier réseau est interdit, car son adresse de broadcast est la même que celle du réseau initial. Il y a donc moins d'adresses disponibles étant données qu'à chaque division certaines adresses ne sont plus utilisables

12. Interface utilisateur du routeur

12.1. Différents modes du routeur

Le programme d'exécution des commandes, ou EXEC, est l'un des composants de la plate-forme logicielle Cisco IOS. EXEC reçoit et exécute les commandes entrées à l'intention du routeur.

Il existe une multitude de modes différents accessibles sur un routeur Cisco :

- **Mode utilisateur** : Mode lecture qui permet à l'utilisateur de consulter des informations sur le routeur, mais ne lui permet pas d'effectuer des modifications. Dans ce mode, on dispose uniquement de commandes de visualisations d'état de fonctionnement du routeur. C'est dans ce mode que l'on arrive lorsque l'on se connecte au routeur.
- **Mode privilégié** : Mode lecture avec pouvoir. On dispose d'une panoplie complète de commandes pour visualiser l'état de fonctionnement du routeur, ainsi que pour importer/exporter et sauvegarder des fichiers de configurations et des images d'IOS.
- **Mode de configuration globale** : Ce mode permet d'utiliser toutes les commandes de configuration ayant une portée globale à tout le routeur.
- **Modes de configuration spécifiques** : On ne dispose dans chaque mode spécifique que des commandes ayant une portée localisée au composant du routeur spécifié par ce mode.
- **Mode SETUP** : Mode affichant un dialogue interactif à l'écran de la console, grâce auquel l'utilisateur néophyte peut créer une configuration élémentaire initiale.
- **Mode RXBoot** : Mode de maintenance permettant notamment de récupérer des mots de passe perdus.

On peut facilement identifier le mode actuel dans lequel on est en repérant l'invite de commande que nous fournit le routeur :

Mode	Invite de commande
Utilisateur	Router >
Privilégié	Router #
Configuration globale	Router (config) #
Interface	Router (config-if) #
Ligne	Router (config-line) #
Routeur	Router (config-router) #

Nous allons maintenant voir les commandes permettant de naviguer dans les différents modes du routeur :

Nous verrons comment atteindre les modes de configuration spécifiques dans les chapitres suivants.

Il existe aussi deux commandes permettant d'arrêter l'exécution d'une commande en cours :

- **Ctrl-Maj-6**
- **Ctrl-C**

12.2. Mode SETUP

Le mode SETUP constitue une des routines de la configuration initiale. L'objectif principal du mode SETUP est de créer rapidement une configuration minimale, à savoir :

- Nom d'hôte du routeur.
- Mots de passe du mode privilégié.
- Mot de passe des lignes VTY.
- Client SNMP.
- L'adresse IP pour une interface.

Les caractéristiques du mode SETUP sont les suivantes :

- Il est lancé manuellement grâce à la commande **setup**.
- Pour la plupart des invites du dialogue de configuration système, les réponses par défaut apparaissent entre crochets [] à la suite de la question.
- Il suffit d'appuyer sur la touche **Entrée** pour accepter ces valeurs par défaut.
- Si le système a déjà été configuré, les valeurs par défaut affichées sont celles de la configuration actuelle.
- Si on configure le système pour la première fois, il s'agit des valeurs par défaut définies en usine.
- Si aucune valeur par défaut n'a été définie en usine, comme dans le cas des mots de passe, aucune valeur n'est affichée après le point d'interrogation (réponse par défaut).
- Pendant le processus de configuration, on peut appuyer à tout moment sur les touches **Ctrl+C** pour mettre fin au processus et recommencer.
- Une fois la configuration terminée, toutes les interfaces sont désactivées.

Lorsque l'on a terminé le processus de configuration en mode SETUP, la configuration que l'on vient de créer est affichée. Le système nous demande alors si on veut utiliser cette configuration.

12.3. Fonctions d'aide du routeur

Le principe d'aide pour les commandes sur la plate-forme logicielle IOS est très simple et est constitué de trois choses :

- **Le caractère ?** : Ce caractère permet d'obtenir les différentes possibilités disponibles. En tant que commande à lui seul, ce caractère indique au routeur de nous fournir une liste complète des commandes accessibles depuis le mode dans lequel on se trouve. Si on entre une partie d'un mot d'une commande, cette fonctionnalité nous affiche les commandes probables commençant par cette partie de mot. Après cela, ce caractère nous indique les autres mots possibles pour ce début de commande. Ce caractère est très souvent utilisé en complément de la **touche de tabulation** pour réaliser certaines commandes parfois complexes.
- **Le caractère ^** : Celui-ci nous indique à quel endroit se trouve une erreur dans une commande erronée. Dans ce cas, il suffit juste de retaper la commande jusqu'à ce caractère, puis d'utiliser le caractère ? pour obtenir la liste des possibilités pour cette commande.
- **La touche de tabulation** : Cette touche est très couramment utilisée en environnement IOS car, à l'instar de certains Shell UNIX, elle effectue une complétion maximale par rapport aux différentes possibilités.

12.4. Utilisation des commandes d'éditations IOS

L'interface utilisateur offre un mode d'édition avancée nous permettant de modifier une commande au cours de la frappe. Voici un tableau résumant ces combinaisons de touche :

Commande	Description
Ctrl-A	Revient au début de la ligne de commande
Échap-B	Reculé d'un mot
Ctrl-B ou flèche vers la gauche	Reculé d'un caractère
Ctrl-E	Va à la fin de la ligne de commande
Ctrl-F ou flèche vers la droite	Avance d'un caractère
Échap-F	Avance d'un mot

Il existe un autre point à voir. Il ne s'agit pas d'une commande en lui-même, mais plutôt d'un petit système d'information pratique. Il s'agit du **caractère \$** qui peut apparaître en début de ligne écran lorsque la commande en elle-même fait plus d'une ligne écran.

12.5. Utilisation de l'historique des commandes IOS

L'interface utilisateur fournit un historique des commandes entrées. Cette fonction est particulièrement utile pour rappeler des commandes ou des entrées longues ou complexes. La fonction d'historique des commandes vous permet d'accomplir les tâches suivantes :

- Réglage de la capacité du tampon d'historique des commandes.
- Rappel des commandes.
- Désactivation de la fonction d'historique des commandes.

Par défaut, la fonction d'historique des commandes est active et le système enregistre 10 lignes de commandes dans son tampon.

Ce tableau nous indique les différentes commandes d'historique que nous avons à notre disposition :

Commande	Description
Ctrl-P ou flèche vers le haut	Rappel de la commande précédente
Ctrl-N ou flèche vers le bas	Rappel de la commande la plus récente
show history	Affiche le listing des commandes en mémoire
terminal history size {taille}	Définit la taille de la mémoire de commandes (valeur maximale de 256)
terminal no editing	Désactive les fonctions d'éditations avancées
terminal editing	Réactive les fonctions d'éditations avancées

Les quatre dernières commandes sont utilisables dans les modes utilisateur et privilégié uniquement.

13. Composants d'un routeur

13.1. Sources de configuration externes

On va présenter les différents composants du routeur qui jouent un rôle essentiel dans le processus de configuration. Un routeur peut être configuré à partir des sources externes suivantes :

- **La ligne console** offre un accès direct au routeur via un câble console.
- **La ligne auxiliaire** permet de connecter un terminal distant au routeur via une ligne RTC par le biais de modems interposés.
- **Les 5 lignes VTY** (terminaux virtuels de 0 à 4). Ces lignes nous permettent d'accéder au routeur par l'intermédiaire de sessions Telnet.
- **Un serveur TFTP** sur le réseau, sur lequel on peut exporter et/ou importer des configurations ainsi que des images d'IOS.
- **Un navigateur Web** en activant le serveur http sur le routeur. Cette activation se fait par l'intermédiaire de la commande **ip http server** dans le mode de configuration globale.

Nous allons maintenant expliquer les différences entre port, ligne et interface :

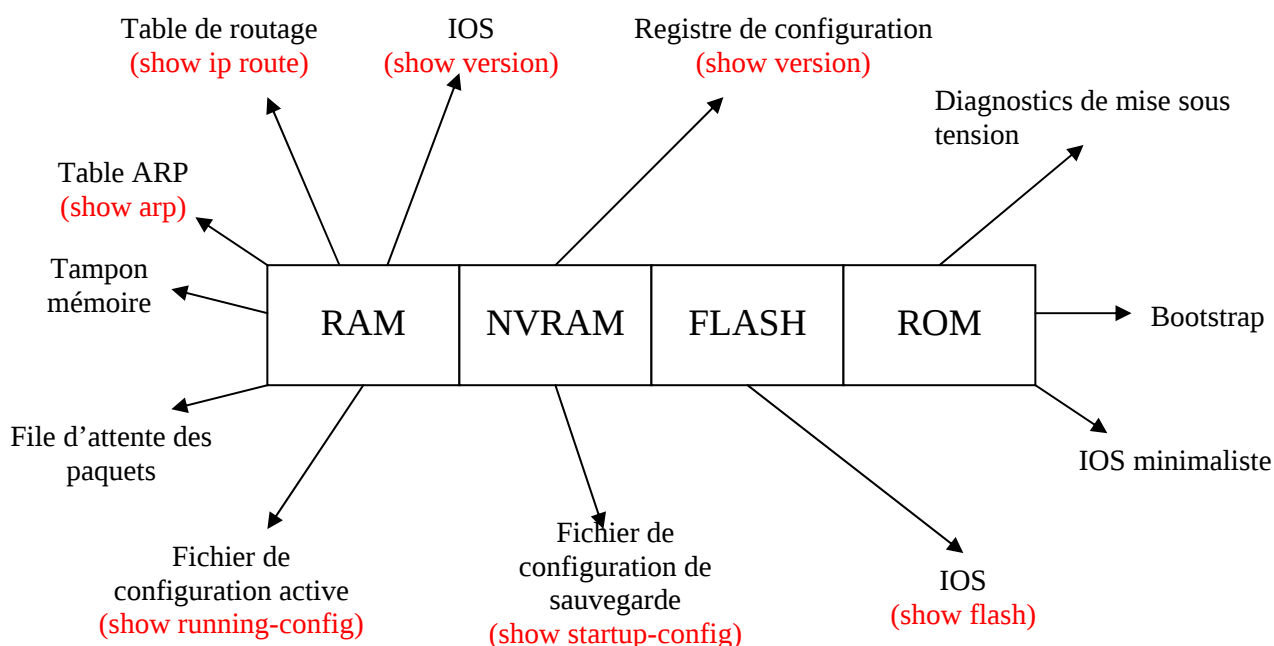
- **Port** : Il s'agit de la partie physique uniquement (Ports RJ45, AUI, Serial).
- **Ligne** : Elles servent uniquement à avoir un accès au routeur afin de pouvoir l'administrer (Lignes console, auxiliaire et VTY).
- **Interface** : Ce sont elles qui interviennent dans le processus d'acheminement de l'information (Paquets). Elles seules possèdent des adresses de couche 2 et 3 (Interfaces Ethernet, Serial).

13.2. Composants de configuration internes et commandes d'état associées

Les composants de configuration internes du routeur sont les suivants :

- **RAM** : C'est la mémoire de travail principal du routeur. Le contenu de cette mémoire est perdu lors de la mise hors tension ou du redémarrage. Sa taille est généralement de 6 ou 8 Mo pour un routeur 25xx.
- **NVRAM (Non-Volatile RAM)** : Elle est relativement lente et est de taille restreinte (environ 32 Ko). Son contenu est conservé lors de la mise hors tension ou du redémarrage.
- **Flash** : Mémoire morte effaçable et reprogrammable (EEPROM). Son contenu est conservé lors de la mise hors tension et du redémarrage. On dispose par défaut de 8Mo de mémoire flash. Elle est l'équivalente du disque dur d'un PC.
- **ROM** : Le contenu de cette mémoire est inaltérable. Le seul moyen de modifier le contenu est de remplacer des puces enfichables sur la carte mère. Sa taille est d'environ 4 Mo. Cette mémoire est l'équivalent du BIOS d'un PC.
- **Interfaces** : Connexions réseau situées sur la carte mère ou sur des modules d'interface distincts, grâce auxquelles les paquets entrent dans le routeur et le quittent.

Les routeurs Cisco proposent plusieurs commandes qui nous permettent d'établir si le routeur fonctionne correctement ou s'il existe des problèmes. Les principales commandes de visualisation d'état sont les suivantes :



- **show version** : Affiche la configuration matérielle système, la version d'IOS, le nom et la source des fichiers de configuration et l'image d'amorçage, ainsi que la valeur du registre de configuration.
- **show processes** : Affiche des informations sur les processus actifs.
- **show protocols** : Affiche le nom et l'état de tous les protocoles configurés de couche 3.
- **show memory** : Affiche des statistiques sur la mémoire du routeur, y compris sur la mémoire disponible.
- **show stacks** : Contrôle l'utilisation de la pile par les processus et les routines d'interruption et affiche le motif du dernier redémarrage système.
- **show buffers** : Fournit des statistiques sur les tampons du routeur.
- **show arp** : Affiche les entrées ARP connues.
- **show flash** : Affiche des informations sur la mémoire flash, telles que la quantité d'espace libre et le nom des fichiers présents dans cette mémoire.
- **show running-config** : Affiche le fichier de la configuration active.
- **show startup-config** : Affiche le fichier de la configuration de sauvegarde.
- **show interfaces [{type} {numéro}]** : Affiche les informations de configuration ainsi que des statistiques de trafic pour chaque interface configurée sur le routeur. Il est possible d'afficher les informations pour une seule interface.
- **clear counters [{type} {numéro}]** : Permet de mettre à zéro toutes les statistiques des interfaces du routeur. Il est possible d'effectuer cette opération sur une seule interface, en indiquant en paramètre l'interface désirée.
- **show ip route** : Affiche la table de routage IP. Cette commande indique de quelle manière chaque entrée de la table a été apprise (statiquement ou par quel protocole de routage).
- **show ip protocols** : Affiche les valeurs des compteurs de routage et les informations de réseau associées à l'ensemble du routeur. Cette commande nous indique les différents réseaux avec lesquels le protocole de routage est configuré pour communiquer, ainsi que la distance administrative de ce dernier.

14. Configuration du routeur

14.1. Fichiers de configuration d'un routeur

Les informations contenues dans un fichier de configuration sont les suivantes :

- Tout d'abord, des informations génériques concernant la version d'IOS avec laquelle le fichier de configuration est prévu pour fonctionner.
- Le nom du routeur ainsi que le mot de passe du mode privilégié.
- Les entrées statiques de résolution de nom vers IP.
- Puis chaque interface avec sa configuration spécifique.
- Toutes les informations de routage.
- Finalement, chaque ligne et sa configuration spécifique.

Les différentes commandes pour les versions 11.x ou ultérieures de la plate-forme logicielle Cisco IOS associées aux fichiers de configuration sont :

- **configure terminal** : Permet de configurer manuellement un routeur à partir d'une console.
- **configure memory** : Charge les informations de configuration à partir de la mémoire NVRAM.
- **copy tftp running-config** : Charge les informations de configuration à partir d'un serveur de réseau TFTP dans la mémoire RAM.
- **show running-config** : Affiche la configuration en cours dans la mémoire RAM.
- **copy running-config startup-config** : Copie la configuration en cours à partir de la mémoire RAM pour la stocker dans la mémoire NVRAM.
- **copy running-config tftp** : Copie la configuration en cours à partir de la mémoire RAM pour la stocker sur un serveur de réseau TFTP.
- **show startup-config** : Affiche la configuration enregistrée, qui représente le contenu de la mémoire NVRAM.
- **erase startup-config** : Efface le contenu de la mémoire NVRAM.

14.2. Configuration des mots de passe

On peut protéger notre système à l'aide de mots de passe pour en restreindre l'accès. Une protection par mot de passe peut être installée pour chaque ligne et se réalise de la manière suivante :

- **line {console | aux | vty} {numéro}** : Permet de passer dans le mode de configuration spécifique à la ligne précisée.
- **password {mot de passe}** : Affecte le mot de passe souhaité à la ligne en cours de configuration.

Quelques informations importantes sont à mettre en évidence concernant les lignes VTY :

- La commande **line vty 0 4** permet d'entrer dans le mode de configuration de toutes les lignes VTY. Ce qui signifie que la commande **password** affectera le même mot de passe pour toutes les lignes VTY de notre routeur.
- Une ligne VTY est active uniquement si un mot de passe est configuré sur cette dernière.

Les mots de passe pour les lignes console et auxiliaire ne sont pris en compte qu'après redémarrage du routeur.

On peut restreindre aussi l'accès au mode privilégié en utilisant au moins une des commandes suivantes :

- **enable password {mot de passe}** : Limite l'accès au mode privilégié.
- **enable secret {mot de passe}** : Idem que **enable password**, mais utilise un processus de cryptage propriétaire de Cisco pour modifier la chaîne de caractère du mot de passe.

Le mot de passe **enable secret** est prioritaire au mot de passe **enable password**, ce qui signifie que, si on utilise ces deux commandes en simultanément, il nous faudra indiquer le mot de passe **enable secret** pour accéder au mode privilégié.

De plus, on peut protéger l'affichage des mots de passe écrits en clair dans le fichier de configuration à l'aide de la commande **service password-encryption**. Cette commande utilise pour cela un algorithme propriétaire Cisco.

14.3. Configuration du nom du routeur et des descriptions.

On va étudier comment configurer :

- Le nom d'hôte du routeur.
- Une bannière de connexion.
- Une description pour chaque interface.

L'attribution d'un nom à notre routeur est une des premières tâches de base à exécuter :

- Il faut utiliser la commande **hostname {nom d'hôte}** à l'invite du mode de configuration globale.
- Le nom du routeur est considéré comme le nom d'hôte
- C'est le nom affiché par l'invite du système.
- Le nom par défaut est **Router**.

On peut aussi configurer une bannière de connexion. Cette bannière s'affiche lors de la connexion et permet de transmettre un message aux utilisateurs du routeur (pour les avertir par exemple d'un arrêt imminent du routeur).

- Pour définir ce message, il faut utiliser la commande **banner motd** dans le mode de configuration globale.
- Il faut encapsuler le message entre deux signes dièse (#), afin d'indiquer au routeur le début et la fin du message.

Enfin, on peut indiquer une description pour chaque interface du routeur. Ceci est très utile pour ceux qui seraient censés travailler sur ce routeur et qui ne connaissent pas forcément à quoi peut être attribué cette interface. Pour cela, il faut :

- Passer dans le mode de configuration de l'interface souhaitée avec la commande **interface {type} {numéro}** depuis le mode de configuration globale.
- Puis d'utiliser la commande **description {texte}**. Ce texte ne pourra pas excéder 80 caractères.

15. Plate-forme logicielle Cisco IOS

15.1. Séquence d'amorçage

Après le test de mise sous tension, les étapes suivantes se déroulent pendant l'initialisation du routeur :

- **Étape 1** : Le bootstrap générique, en mémoire ROM, s'exécute sur le processeur. Le bootstrap est une opération simple et prédéfinie qui charge des instructions. Celles-ci chargent à leur tour d'autres instructions en mémoire ou activent d'autres modes de configuration.
- **Étape 2** : Le système d'exploitation peut être installé à plusieurs endroits. Son emplacement est précisé dans le champ d'amorçage du registre de configuration. Si le champ indique une mémoire flash ou un serveur TFTP, les commandes **boot system** du fichier de configuration précisent l'emplacement exact de l'image.
- **Étape 3** : L'image du système d'exploitation est chargée. Une fois chargé et en fonction, le système d'exploitation recherche les composants matériels et logiciels, puis il affiche les résultats sur la console.
- **Étape 4** : Le fichier de configuration stocké dans la mémoire NVRAM est chargé dans la mémoire principale, puis il est exécuté ligne par ligne. Ces commandes de configuration lancent les processus de routage, fournissent les adresses aux interfaces, définissent les caractéristiques des médias, etc.
- **Étape 5** : Si la mémoire NVRAM ne contient pas de fichier de configuration valide, le système d'exploitation exécute une routine de configuration initiale interactive appelée dialogue de configuration système ou mode SETUP.

La commande **reload** permet de redémarrer à chaud le routeur.

15.2. Caractéristiques fondamentales

Le processus de chargement de la plate-forme logicielle IOS se fait dans cet ordre :

- Le bootstrap identifie la valeur du champ d'amorçage du registre de configuration.
- La valeur du champ d'amorçage indique l'emplacement des commandes **boot system**.
- Si le routeur ne trouve pas de commandes **boot system** dans l'emplacement spécifié, alors il prendra celles par défaut contenues dans la mémoire ROM.
- Les commandes **boot system** précisent l'emplacement de l'image d'IOS ainsi que l'ordre de recherche de cet emplacement.
- Si ces commandes **boot system** indiquées ne permettent pas de trouver une image d'IOS valide, alors le routeur prendra la séquence d'amorçage par défaut précisée en mémoire ROM.

La séquence d'amorçage par défaut est :

- Recherche de l'image d'IOS installée par défaut en mémoire flash.
- Puis recherche d'un fichier sur le serveur TFTP dont l'IP est 255.255.255.255.
- Finalement, en dernier recours, le routeur prendra l'image présente en mémoire ROM.

L'ordre dans lequel le routeur cherche les commandes **boot system** dépend de la valeur indiquée dans le champ d'amorçage du registre de configuration. On peut modifier la valeur par défaut grâce à la commande du mode de configuration globale **config-register {valeur}**. Il faut utiliser un nombre hexadécimal comme argument à cette commande.

Le registre de configuration est un registre de 16 bits qui se trouve dans la mémoire NVRAM. Les 4 bits inférieurs constituent le champ d'amorçage. Le tableau suivant nous indique les différentes valeurs possibles pour ce champ d'amorçage, ainsi que leur signification :

Valeur	Description
0x---0	Utiliser le mode moniteur de mémoire ROM (démarrer manuellement à l'aide de la commande b)
0x---1	Démarrer automatiquement à partir de la mémoire ROM
0x---2 à 0x---F	Rechercher les commandes boot system dans la mémoire NVRAM (0x---2 par défaut)

La commande **show version** affiche :

- La version et le numéro de révision de la plate-forme logicielle Cisco IOS en exploitation sur le routeur.
- La valeur du registre de configuration
- Le nom du fichier de l'image IOS qui a été chargée ainsi que sa provenance.
- Des informations diverses sur les tailles des mémoires installées sur le routeur.

15.3. Commandes boot system

La commande **boot system** (aussi appelée donnée ou option bootstrap) peut désigner trois types d'emplacements pour la plate-forme logicielle Cisco IOS :

- **boot system flash {nom du fichier}** : Cette commande nous permet de spécifier le fichier présent en mémoire flash qui va être chargé au démarrage.
- **boot system tftp {nom du fichier} {IP du serveur TFTP}** : Dans le cas où la mémoire flash serait endommagée, le chargement d'une image système à partir d'un serveur TFTP représente une solution de secours. C'est aussi la méthode la plus utilisée pour mettre à jour un routeur vers une nouvelle version d'IOS.
- **boot system rom** : Si la mémoire flash est endommagée et que le serveur TFTP ne réussit pas à charger l'image, un amorçage à partir de la mémoire ROM est la dernière option bootstrap dont dispose la plate-forme logicielle.

La commande **show flash** permet de visualiser l'état de la mémoire flash. Elle est très importante car elle nous donne l'espace mémoire libre. En effet, si l'on veut copier une image de la plate-forme logicielle dans la mémoire flash, il faut s'assurer que l'on dispose de suffisamment de place pour réaliser l'opération.

15.4. Manipulation des images logicielles d'IOS

La convention d'attribution de noms pour la version 11.2 de la plate-forme logicielle Cisco IOS comprend trois parties :

- La plate-forme sur laquelle l'image est exécutée.
- Une lettre ou des séries de lettres identifiant les fonctions et les capacités spéciales compatibles avec l'image.
- Un caractère qui indique l'endroit où l'image est exécutée et si elle a été compressée (exemple : **1** = relogeable, non compressée ; **m** = RAM, non compressée ; **mz** = RAM, compression zip).

Nous avons la possibilité d'exporter une image vers un serveur TFTP ainsi que d'en importer une vers la mémoire flash. Avant de réaliser l'une ou l'autre des opérations précédemment citées, il est bon d'utiliser la commande **show flash** afin de connaître le nom de l'image actuellement présente dans la mémoire flash.

L'exportation d'une image logicielle d'IOS (de la mémoire flash vers le serveur TFTP) s'effectue de la façon suivante :

- **copy flash tftp** : Cette commande lance la procédure d'exportation.
- IOS nous invite ensuite à indiquer l'adresse IP du serveur TFTP (par défaut 255.255.255.255).
- Il nous incombe ensuite à entrer le nom du fichier qui sera écrit sur ce serveur TFTP.
- Enfin, un commentaire nous indique si l'opération s'est correctement déroulée.
- L'importation d'une image d'IOS est très similaire, bien qu'il y ait un peu plus d'informations :
- **copy tftp flash** : Lance la procédure d'importation en mémoire flash.
- On nous invite ensuite à indiquer l'adresse IP du serveur TFTP source.
- Il faut maintenant indiquer le nom du fichier sur le serveur TFTP.
- Il nous faut maintenant confirmer la copie de cette image dans la mémoire flash.
- IOS nous donne à ce moment la quantité d'espace disponible dans la mémoire flash. Ceci est très important pour prendre une décision sur la question suivante.
- Cette question correspond au formatage de la mémoire flash. Ceci peut s'avérer utile si on ne dispose pas d'assez d'espace libre. On peut l'effectuer aussi si on ne veut qu'une seule image d'IOS dans la mémoire flash.
- Viennent ensuite les phases de formatage, au cas où l'on a confirmé la question précédente, et de chargement de l'image logicielle d'IOS dans la mémoire flash. Un checksum est effectué pour vérifier l'intégrité de l'image.

16. Adressage IP et interfaces

16.1. Adresse IP d'une interface

L'adressage IP des interfaces est la deuxième phase dans le processus de configuration d'un routeur. Pour attribuer une adresse IP à une interface, il faut être dans le mode de configuration de cette dernière.

Pour cela, on utilise la commande **interface {type} {numéro}** depuis le mode de configuration globale.

La commande **ip address {IP de l'interface} {masque de sous-réseau}** nous permet de réaliser cette attribution.

On se souvient qu'un masque de sous-réseau peut être écrit de plusieurs manières différentes :

- Notation entière avec des points de séparation (option par défaut). Exemple : 255.255.255.0
- Nombre de bits. Exemple : /24
- Notation hexadécimale. Exemple : 0xFFFFF00

Il existe donc une commande nous permettant de spécifier le format du masque de sous-réseau. Il s'agit de la commande **term ip netmask-format {format}**.

Il reste maintenant à activer l'interface. La commande **no shutdown** permet d'effectuer cette opération.

16.2. Résolution de nom vers IP statique

Tout comme sur les systèmes d'exploitation les plus connus, il est possible d'avoir une correspondance IP / nom d'hôte (exemple : fichier hosts sous UNIX). Pour cela, nous avons à notre disposition la commande **ip host {nom} [tcp-port-number] {adresse IP} [{2° adresse IP} ...]**, accessible depuis le mode de configuration globale (NB : les paramètres entre crochets sont optionnels).

Il est tout à fait possible d'attribuer plusieurs adresses IP pour un même nom d'hôte. Il est en effet possible pour un seul et même hôte d'avoir plusieurs adresses.

Le paramètre **tcp-port-number** permet d'explicitement le port TCP à utiliser lors de l'exécution de la commande **telnet** ou **connect**. Implicitement, sa valeur est 23.

16.3. Service DNS

La commande **ip name-server {adresse serveur DNS} [{adresse d'un 2° serveur} ...]** permet de définir les machines qui peuvent fournir le service de noms. On peut spécifier jusqu'à 6 adresses IP en tant que serveurs de noms.

Pour mettre en correspondance les noms de domaine avec des adresses IP, on doit identifier les noms de machine, spécifier un serveur de noms et activer le service DNS.

La commande **ip domain-lookup** du mode de configuration globale permet d'activer le service DNS sur un routeur Cisco. Ce service est par ailleurs activé par défaut. Lorsque ce service est désactivé, le routeur ne génère pas ou ne transmet pas les paquets de broadcast générés par le système de noms.

La commande **show hosts** permet d'afficher la liste des noms de machine ainsi que les adresses associées, la méthode d'apprentissage ainsi que la pertinence de chaque entrée. Cette commande nous livre un maximum d'informations dont les significations sont expliquées dans le tableau suivant :

Information	Description
Host	Noms des machines connues
Flag	Description de la méthode utilisée pour apprendre les informations et pour juger de leur pertinence actuelle
perm	Configuré manuellement dans une table d'hôtes
temp	Acquis par le biais d'un serveur DNS
OK	Entrée en cours
EX	Entrée obsolète, expirée
Age	Temps (en heures) écoulé depuis que le logiciel a consulté l'entrée
Type	Champ de protocole
Address(es)	Adresses logiques associées au nom de machine

16.4. Spécificités des interfaces WAN

Les connexions WAN nécessitent une configuration supplémentaire par rapport à celles en LAN. En effet, il faut indiquer à quelle vitesse va fonctionner notre liaison.

Ceci se fait par l'intermédiaire de la commande **clock rate {valeur}**. Cette configuration doit être effectuée uniquement sur la partie ETCD de la liaison.

17. CDP

17.1. CDP

Le protocole CDP (Cisco Discovery Protocol) est un protocole propriétaire Cisco permettant la découverte des voisins.

Il permet d'obtenir des informations sur les dispositifs connectés au routeur local. Ce protocole devient très utile lorsque l'on n'a aucun moyen (visuellement ou par accès de configuration) pour analyser la topologie réseau.

17.2. Théorie

Le protocole CDP permet principalement de connaître les plateformes et les protocoles utilisés par les dispositifs voisins (c'est-à-dire directement connectés).

Voici les différentes caractéristiques du protocole CDP :

- Existe depuis IOS 10.3
- Actif par défaut
- Fonctionne au niveau de la couche 2 (permet donc d'obtenir des informations sur les voisins même si les protocoles de couche 3 sont différents ou non configurés)
- Trames CDP multicast envoyées toutes les 60 secondes

CDP peut fournir ces informations :

Information	Description
ID de dispositif	Nom d'hôte et nom de domaine du voisin
Liste d'adresses	Une adresse pour chaque protocole routé du voisin
Identifiant de port	Interface du voisin utilisée pour se connecter au routeur local
Liste de capacités	Fonction du dispositif voisin (routeur, pont, commutateur, etc.)
Version d'IOS	Version d'IOS du voisin
Plateforme	Type de dispositif (Cisco 2620XM, Catalyst 2950, etc.)

17.3. Configuration

La configuration de CDP est très simple, et se résume à ces commandes :

- **[no] cdp run**
 - o Mode de configuration globale
 - o Active/désactive le protocole CDP pour tout le routeur
 - o Actif par défaut
- **[no] cdp enable**
 - o Mode de configuration d'interface
 - o Active/désactive le protocole CDP pour cette interface
 - o Actif par défaut sur toutes les interfaces fonctionnelles
- **cdp timer {temps}**
 - o Mode de configuration globale
 - o Spécifie l'intervalle de temps en secondes pour l'émission des trames CDP
 - o Temps par défaut : 60 secondes
- **cdp holdtime {temps}**
 - o Mode de configuration globale
 - o Spécifie le temps en secondes avant suppression d'une information non rafraîchie
 - o Temps par défaut : 180 secondes

17.4. Visualisation et résolution de problèmes

Voici les commandes utilisées pour afficher les informations obtenues grâce à CDP :

- **show cdp** : Affiche les compteurs de temps pour CDP
- **show cdp interface [{type} {numéro}]** : Affiche les interfaces sur lesquelles CDP est activé
- **show cdp entry {nom | *}** : Affiche les informations d'un ou des voisins
- **show cdp neighbors [detail]** : Affiche la liste des voisins CDP ainsi que les informations les concernant
- **show cdp traffic** : Affiche les compteurs de trafic CDP
- **clear cdp counters** : Réinitialise les compteurs de trafic CDP
- **clear cdp table** : Vide la table d'informations CDP

18. VLSM et CIDR

18.1. Introduction au routage Classless

Au début des années 90, Internet subissait une croissance exponentielle annonçant un épuisement des adresses IPv4, notamment celles de classe B.

Cette pénurie d'adresse est principalement due au découpage fixe de l'espace d'adressage total IPv4 en classes (classe A, classe B, classe C) qui fige le nombre de réseaux possibles et le nombre d'hôtes maximum par réseau.

En effet, lorsque l'on utilise un **adressage classful**, les masques de sous-réseaux ne sont pas envoyés sur le réseau. Les équipements réseaux utilisent donc des masques de sous-réseaux par défaut qui sont les suivants :

- Classe A : 255.0.0.0 ou /8
- Classe B : 255.255.0.0 ou /16
- Classe C : 255.255.255.0 ou /24

Il est dans ce cas impossible de créer des sous-réseaux et de former des groupes d'utilisateur de différentes tailles au sein d'un réseau d'entreprise.

Ce problème est résolu avec l'utilisation d'un **adressage classless** (sans classe) qui permet d'envoyer le masque de sous-réseau utilisé aux autres équipements et de ce fait, de créer des sous-réseaux de taille variable.

Le CIDR et le VLSM sont des exemples de procédures utilisant un adressage classless. Bien que complémentaires, celles-ci sont différentes. Le VLSM peut d'ailleurs être vu comme une extension du CIDR au niveau d'une organisation.

Le VLSM permet en effet d'éviter le gaspillage d'adresse au sein d'une organisation en utilisant des masques de taille variable, tandis que le CIDR permet de diminuer significativement le nombre d'entrées des tables de routage en utilisant des agrégations de routes.

Il existe cependant des règles à suivre concernant la création et l'utilisation de sous-réseaux. Ces règles sont régies par les RFC 950 (règle du 2ⁿ-2) et RFC 1878 (règles du 2ⁿ-1 et du 2ⁿ) :

- **Règle du 2ⁿ - 2** → impossible d'utiliser le premier sous-réseau ainsi que le dernier sous-réseau
- **Règle du 2ⁿ - 1** → impossible d'utiliser le premier sous-réseau
- **Règle du 2ⁿ** → utilisation de tous les sous-réseaux

L'utilisation d'une de ces règles par rapport à une autre dépend uniquement des capacités techniques des équipements. De nos jours la majorité des réseaux utilisent la règle du 2ⁿ puisqu'elle permet de limiter au maximum le gaspillage d'adresses IP.

18.2. CIDR

L'expansion d'Internet a entraîné l'augmentation de la taille des tables de routage sur de nombreux routeurs, notamment les routeurs des fournisseurs d'accès à Internet.

Pour alléger de manière considérable ces tables de routage, une solution permettant d'agréger plusieurs routes en une seule a dû être mise en place : c'est le principe du **CIDR** (Classless Inter-Domain Routing).

Pour ce faire, une comparaison binaire de l'ensemble des adresses à agréger est nécessaire. Il faut en effet arriver à déterminer les bits de la partie réseau qui sont en commun dans toutes ces adresses et mettre à zéro tous les bits restant.

De cette manière une délimitation entre la partie réseau commune et le reste de l'adresse sera effectuée. Celle-ci permettra de déterminer l'adresse agrégée ainsi que le nouveau masque de sous-réseau à utiliser.

L'exemple suivant illustre l'utilisation d'une agrégation de quatre adresses réseaux en une seule adresse. Il faut en effet agréger les 4 réseaux ci-dessous :

- 10.3.4.0 255.255.255.0 (ou /24)
- 10.3.5.0 255.255.255.0 (ou /24)
- 10.3.6.0 255.255.255.0 (ou /24)
- 10.3.7.0 255.255.255.0 (ou /24)

Processus d'agrégation (ou summarization) de routes en une seule :

<u>Adresses réseaux :</u>	10.3.4.0 - 00001010 . 00000011 . 00000100 . 00000000
	10.3.5.0 - 00001010 . 00000011 . 00000101 . 00000000
	10.3.6.0 - 00001010 . 00000011 . 00000110 . 00000000
	10.3.7.0 - 00001010 . 00000011 . 00000111 . 00000000
<u>Nouveau masque :</u>	255.255.252.0 - 11111111 . 11111111 . 11111100 . 00000000
<u>Nouvelle route agrégée :</u>	10.3.4.0 255.255.252.0 (ou /22)

Cependant l'emploi de CIDR n'est possible que si :

- Le protocole de routage utilisé transporte les préfixes étendus dans ses mises à jour.
- Les routeurs implémentent un algorithme de la correspondance la plus longue.
- Un plan d'adressage hiérarchique est appliqué pour l'assignation des adresses afin que l'agrégation puisse être effectuée.
- Les hôtes et les routeurs supportent le routage classless.

18.3. VLSM

L'utilisation du masque de sous-réseau à taille variable (**Variable Length Subnet Mask**) permet à un réseau classless d'utiliser différents masques de sous-réseaux au sein d'une organisation et d'obtenir par conséquent des sous-réseaux plus appropriés aux besoins.

Cependant, certaines conditions sont requises pour utiliser le VLSM :

- Il est nécessaire d'employer un protocole de routage supportant le VLSM. **RIPv.2, OSPF, IS-IS, EIGRP, BGP** ainsi que le **routage statique** supportent VLSM. Les protocoles de routage classless, contrairement aux protocoles de routage classful (RIPv.1, IGRP), transmettent dans leurs mises à jour de routage, le masque de sous-réseau pour chaque route.
- Les routeurs doivent implémenter un algorithme de la correspondance la plus longue. En effet, les routes qui ont le préfixe le plus élevé sont les plus précises. Les routeurs dans leurs décisions d'acheminement doivent être capables de déterminer la route la plus adaptée aux paquets traités.
- Un plan d'adressage hiérarchique doit être appliqué pour l'assignation des adresses afin que l'agrégation puisse être effectuée.

VLSM repose sur l'agrégation. C'est-à-dire que plusieurs adresses de sous-réseaux sont résumées en une seule adresse. L'agrégation est simple, l'on retient simplement la partie commune à toutes les adresses des sous-réseaux.

Pour conceptualiser un réseau conforme VLSM, il faut:

- 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, les grandes agglomérations, avec pour chaque agglomération, les villes, le nombre de bâtiments dans chaque ville, le nombre d'étages par bâtiment et le nombre d'utilisateur par étage.
- Pour chacune de ces entités, réserver le nombre de bits nécessaire en prévoyant l'évolutivité du réseau.
- Calculer le masque de sous-réseau à chaque niveau de l'organisation.

18.4. Procédure de réalisation

Les procédures de réalisation de plan d'adressage avec du VLSM symétrique puis asymétrique sont expliquées. Néanmoins, il faut savoir que le VLSM symétrique n'est qu'une étude de cas scolaire et que le VLSM asymétrique est ce qui est réellement utilisé dans la réalité.

18.4.1. VLSM Asymétrique

Le VLSM Asymétrique, ou plus simplement, VLSM, correspond à une topologie d'entreprise où les différents niveaux hiérarchiques et les instances ne sont pas similaires (nombre, taille etc.)

Procédure :

- **Etape 1 : Identifier le besoin :**

Dessiner la topologie, identifier les besoins à chaque niveau hiérarchique.

- **Etape 2 : Recensement :**

Connaître le nombre d'utilisateurs pour chaque sous-réseau (puisque'ils peuvent être différents à chaque niveau maintenant), ce qui revient à connaître la taille de chaque sous-réseau (ne pas oublier qu'on ne peut pas utiliser la première ni la dernière adresse et qu'il faut une adresse IP pour la passerelle).

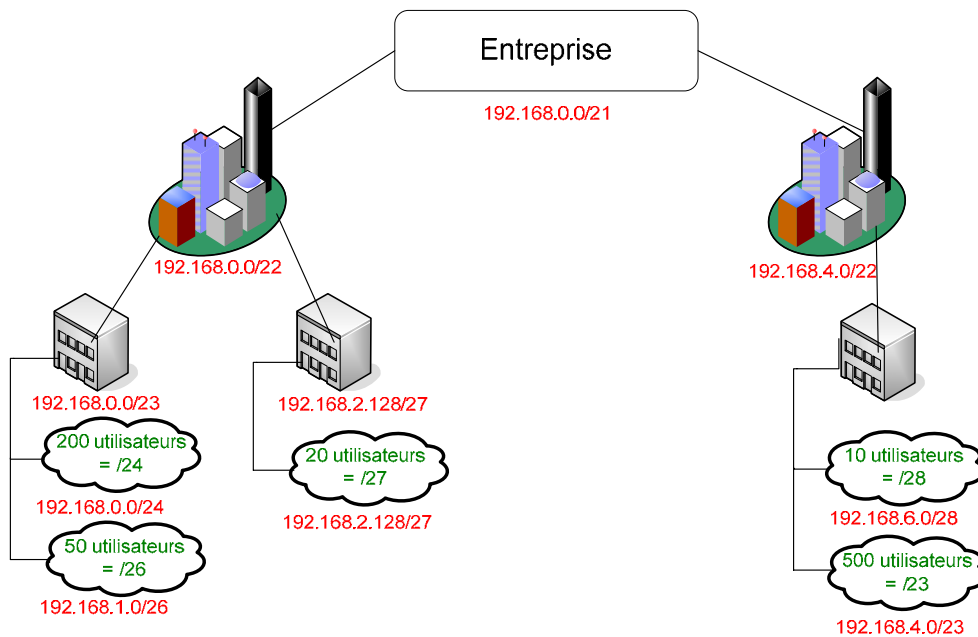
Si le nombre d'utilisateur n'est pas connu à chaque niveau de la hiérarchie, on peut suivre un processus descendant ('top down') : répartir équitablement le nombre d'utilisateur pour un niveau hiérarchique supérieur vers le niveau directement inférieur.

- **Etape 3 : Classe d'adresse utilisée :**

Déterminer la classe d'adresse ou l'agrégat d'adresses (le choix dépendant du contexte), en additionnant tous les bits nécessaires pour identifier chaque niveau hiérarchique de l'entreprise.

- **Etape 4 :**

En suivant un processus remontant récursif maintenant, on va agréger les différentes instances d'un niveau pour obtenir l'identifiant réseau du niveau hiérarchique directement supérieur jusqu'à obtenir l'adresse agrégée de toute l'entreprise.



Etape 1 : Une entreprise dans deux villes. Deux bâtiments dans la première ville, un seul bâtiment dans la deuxième ville. Tous les bâtiments ont deux étages sauf un qui en a qu'un seul. Le nombre d'utilisateur varie d'un étage à l'autre.

Etape 2 : Recensement (en vert). Ne pas oublier l'adresse pour le broadcast, l'adresse pour le réseau et l'adresse pour la passerelle.

Etape 3 : Dans ce contexte, on peut découper une classe B (beaucoup de gaspillage) ou agréger plusieurs classe C. On choisira une classe C

Etape 4 : En remontant, on adresse chaque étage, chaque bâtiment etc. (en rouge)

18.5. Configuration

Lorsque la règle du 2^n-1 est appliquée, il est convenu de ne pas utiliser le premier sous-réseau pour éviter toute confusion. En effet, l'adresse réseau du premier sous-réseau correspond à l'adresse réseau de toute la plage d'adresse.

Pour limiter le gaspillage d'adresse, en utilisant la règle du 2^n , il suffit d'utiliser la commande **ip subnet-zero** qui permet l'utilisation du premier sous-réseau calculé. Cette fonctionnalité est active par défaut depuis la version 12.0 de l'IOS.

- **ip subnet-zero**
 - o Mode de configuration globale
 - o Permet d'utiliser le premier sous-réseau (2^n)

Par ailleurs, la commande **ip classless** active la prise en charge des informations ne respectant pas le découpage d'adresses en classes. C'est-à-dire qu'elle permet d'activer le support des masques de sous-réseau et d'une route par défaut. Cette commande est active par défaut.

- **ip classless**
 - o Mode de configuration globale
 - o Permet d'activer le support des masques de sous-réseau et d'une route par défaut

Lors de l'emploi du VLSM, il faut avant tout s'assurer du bon calcul des masques de sous-réseaux. Une fois cette étape effectuée nous pouvons configurer les interfaces.

- **interface {type} {numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'interface
- **ip address {IP} {masque}**
 - o Mode de configuration d'interface
 - o Permet d'attribuer une adresse IP à cette interface

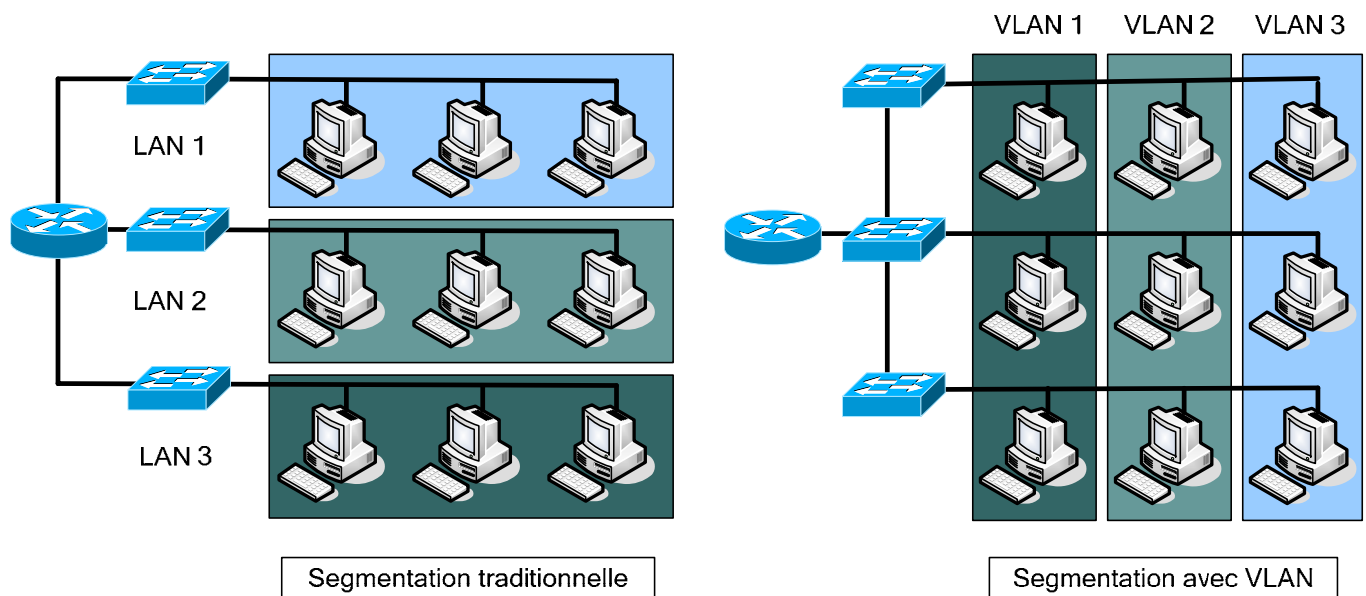
19. VLAN

19.1. Concepts

Un LAN virtuel est un ensemble d'unités regroupées en domaine de broadcast quelque soit l'emplacement de leur segment physique.

Les principales différences entre la commutation traditionnelle et les LAN virtuels sont:

- Les LAN virtuels fonctionnent au niveau des couches 2 et 3 du modèle OSI.
- La communication inter LAN virtuels est assurée par le routage de couche 3.
- Les LAN virtuels fournissent une méthode de contrôle des broadcasts.
- Les LAN virtuels permettent d'effectuer une segmentation selon certains critères:
 - Des collègues travaillant dans le même service.
 - Une équipe partageant le même applicatif.
- Les LAN virtuels peuvent assurer la sécurité des réseaux en définissant quels nœuds réseaux peuvent communiquer entre eux.



Il est donc possible de segmenter le réseau en plusieurs domaines de broadcast afin d'en améliorer les performances.

On distingue 2 méthodes de création pour les LAN virtuels :

- **LAN statique** : ces VLAN sont dits accès sur les ports. L'appartenance à un VLAN est en effet fonction du port sur lequel est connecté un utilisateur (corrélation de couche 1 : port <-> VLAN). La configuration des commutateurs se fait donc en attribuant un port à un VLAN.
- **LAN dynamique** : dans cette configuration, l'appartenance à un VLAN est déterminée par une information de couche supérieure : 2 ou plus (corrélation de couche >=2 <-> VLAN). Typiquement, on peut baser l'appartenance à un VLAN en fonction de l'adresse MAC de l'utilisateur. Cette configuration nécessite un logiciel d'administration réseau (ex : CiscoWorks 2000) basé sur un serveur. Lors de la connexion d'un hôte au commutateur, ce dernier enverra une requête au serveur lui indiquant, par exemple, l'adresse MAC du nouvel hôte connecté. Le serveur, grâce à une base de données liant MAC et VLAN (remplie par l'administrateur), renverra alors le VLAN d'appartenance au commutateur.

19.1.1. Commandes générales

- **vlan database**
 - o Mode privilégié
 - o Permet d'accéder au mode de configuration de VLAN.
- **vlan vlan_id [name { nom du vlan }]**
 - o Mode de configuration des VLAN (vlan database)
 - o Permet de créer et nommer les VLANs.
- **switchport mode {access | dynamic {auto | desirable} | trunk}**
 - o Mode de configuration d'interface
 - o Permet de configurer une interface pour le trunking ou pour un VLAN.
- **switchport access vlan vlan-id**
 - o Mode de configuration d'interface
 - o Permet de configurer un VLAN statique sur une interface.

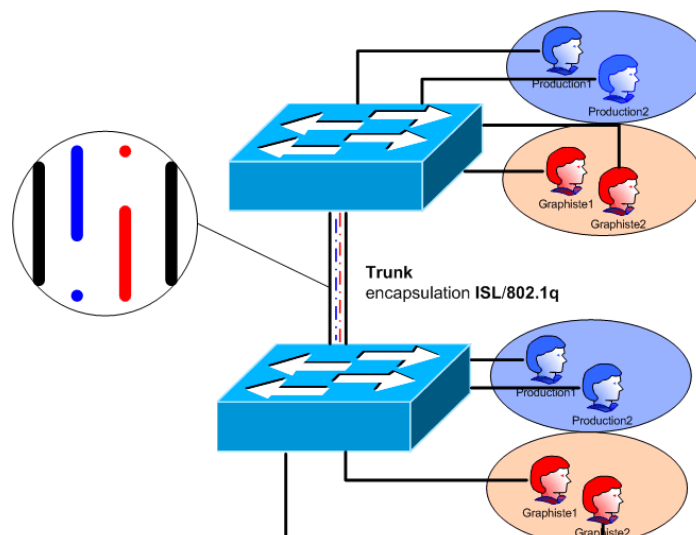
19.1.2. Commandes show associées

- **show interfaces [interface-id | vlan vlan-id] [switchport | trunk]**
 - o Affiche les statuts du trunking.
- **show vlan [brief | id vlan-id | name vlan-name | summary]**
 - o Liste les informations sur le VLAN.
- **show vlan [vlan]**
 - o Affiche des informations sur le VLAN.
- **show spanning-tree vlan vlan-id**
 - o Affiche les informations spanning-tree pour le VLAN spécifié.

19.1.3. Configuration

- **Configurer un VLAN statique**
 - o Entrez dans le mode de configuration de VLAN à l'aide de la commande **vlan database**.
 - o Créez le VLAN avec la commande **vlan {vlan number}**.
 - o Entrez dans le mode de l'interface que vous souhaitez associer au VLAN.
 - o Spécifiez le mode du port pour un VLAN : **switchport mode access**.
 - o Spécifiez le VLAN avec la commande **switchport access vlan vlan-id**.
- **Sauvegarder la configuration VLAN**
 - o Les configurations de VLAN sont automatiquement sauvegardées dans la flash dans le fichier **vlan.dat**.

19.2. Trunking



Le trunking permet, dans des réseaux comportant plusieurs commutateurs, de transmettre à un autre commutateur via un seul port, le trafic de plusieurs VLAN (dont les membres sont dispatchés sur plusieurs commutateurs). Le problème étant que différents trafics isolés (de différents VLAN) doivent emprunter un seul câble.

On a donc plusieurs trafics logiques sur une liaison physique : on appelle cette notion un trunk. Afin d'identifier l'appartenance des trames aux VLAN, on utilise un système d'étiquetage (ou encapsulation) sur ce lien.

Il en existe deux protocoles :

- **ISL** (Inter Switch Link) qui est un protocole propriétaire Cisco.
- **802.1q** qui est un standard de l'IEEE.

19.2.1. Protocole ISL

Cisco avait développé bien avant l'IEEE son protocole ISL. Comme ISL est un protocole propriétaire Cisco, il ne peut être appliqué qu'à des commutateurs Cisco.

Avec l'emploi d'ISL, la trame originelle est encapsulée entre un en-tête de 26 octets et un en-queue de 4 octets.

Trame ISL

En-tête ISL 26 octets	Trame Ethernet encapsulée	FCS 4 octets
--------------------------	------------------------------	-----------------

Composition de l'en-tête ISL

DA 40 bits	Type 4 bits	Util. 4 bits	SA 48 bits	LEN 16 bits	AAAA03 24 bits	HSA 24 bits	VLAN 16 bits	BPDU 1 bit	INDEX 16 bits	RES 16 bits
---------------	----------------	-----------------	---------------	----------------	-------------------	----------------	-----------------	---------------	------------------	----------------

- DA : Adresse multicast de destination qui prend la valeur 0x01-00-0C-00-00 ou 0x03-00-0C-00-00.
- Type : Indique le type de trame (Ethernet, Token Ring, etc.).
- Util : Indique la priorité de traitement de la trame.
- SA : Adresse MAC source.
- LEN : Longueur de la trame encapsulé moins les 18 bits des champs DA, Type, Util., SA, LEN et FCS.
- AAAA03 : Champ SNAP d'une valeur fixe 0xAAAA03.
- HSA : Contient la portion constructrice de l'adresse MAC source.
- VLAN : Identifiant de VLAN.
- BPDU : Utilisé par l'algorithme Spanning Tree pour déterminer les informations topologiques.
- INDEX : Employé à des fins diagnostiques uniquement.
- RES : Utilisé quand une trame Token Ring ou FDDI est encapsulé dans une trame ISL.

19.2.2. Protocole 802.1q

Contrairement à ISL le protocole développé par L'IEEE 802.1q n'encapsule pas la trame Ethernet originale, mais insère un en-tête additionnel de 4 octets qui contient un champ d'identification du VLAN.

Le champ de contrôle de trame (FCS) doit être recalculé à cause de l'ajout de l'en-tête additionnel.

Trame Ethernet avec 802.1q.

Dest	Src	Etype	Tag	Long/Type Ether	Données	FCS
------	-----	-------	-----	-----------------	---------	-----

En-tête Tag.

Priorité	ID VLAN
----------	---------

19.2.3. Comparaison entre ISL et IEEE 802.1q

ISL	IEEE 802.1q
Encapsule la trame d'origine.	Ajoute un en-tête additionnel à la trame d'origine.
Comporte un champ d'identification de VLAN de 12 bits.	
Utilisation de PVST (Per VLAN Spanning Tree) pour obtenir un arbre STP par VLAN.	

19.2.4. Commandes associées

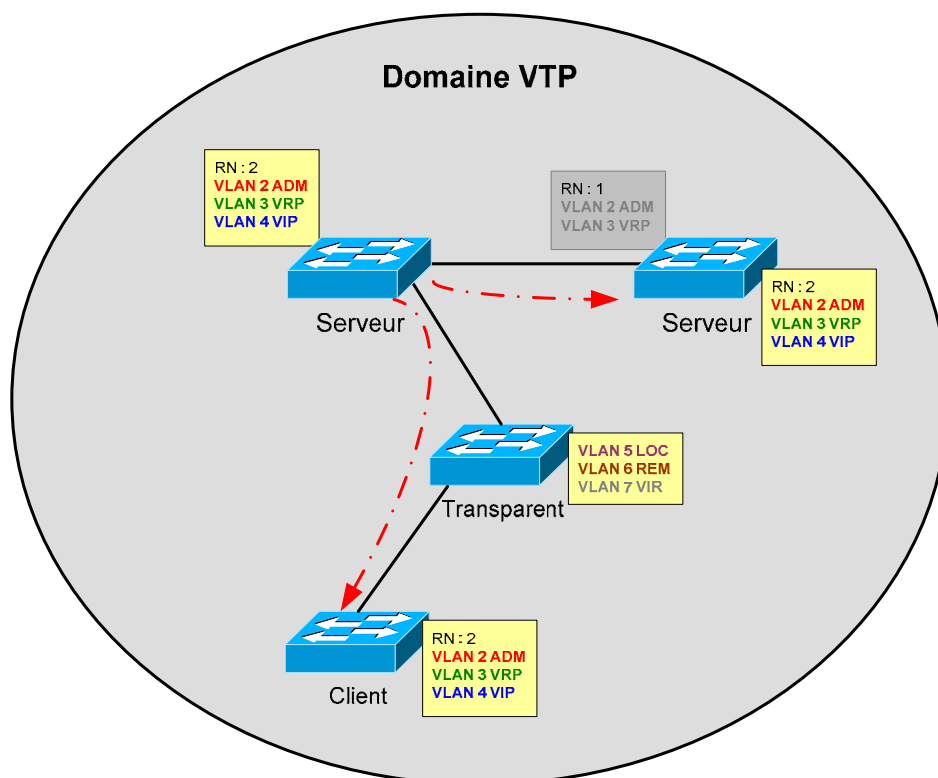
- **switchport mode trunk**
 - Mode de configuration d'interface
 - Active le mode « Trunking » sur l'interface
- **switchport trunk [allowed | encapsulation | native | pruning]**
 - Mode de configuration d'interface
 - Autorise la transport du trafic de certain VLAN sur le lien « trunk »
 - Spécifie le type d'encapsulation (ISL, 802.1q)
 - Permet d'activer le « Pruning » des VLAN
- **show port capabilities [numéro/sous-numéro]**
 - Mode privilégié
 - Affiche les fonctionnalités supportées par l'interface.
- **show interfaces [N° Module/ N° Port] trunk**
 - Mode de configuration d'interface
 - Permet de vérifier la configuration du « trunking » sur cette interface.

19.3. VTP

19.3.1. Théorie sur le protocole VTP

VTP (VLAN Trunking Protocol), protocole propriétaire Cisco permet, aux commutateurs et routeurs qui l'implémentent, d'échanger des informations de configuration des VLAN.

Il permet donc de redistribuer une configuration à d'autres commutateurs, évitant par la même occasion à l'administrateur de faire des erreurs, en se trompant par exemple de nom de VLAN. VTP diffuse ses mises à jour au sein du domaine VTP toutes les 5 min ou lorsqu'une modification a lieu.



Les mises à jour VTP comportent :

- Un numéro de révision (**Revision Number**) qui est incrémenté à chaque nouvelle diffusion. Cela permet aux commutateurs de savoir s'ils sont à jour.
- Les noms et numéro de VLAN.

Dans un domaine VTP, on distingue une hiérarchie comprenant trois modes de fonctionnement :

- VTP **serveur**
- VTP **client**
- VTP **transparent**

Les commutateurs qui font office de serveur VTP peuvent créer, modifier, supprimer les VLAN et d'autres paramètres de configuration. Ce sont eux qui transmettront cette configuration aux commutateurs en mode client (ou serveur) dans leur domaine VTP.

Les commutateurs fonctionnant en mode client ne peuvent que recevoir et transmettre les mises à jour de configuration.

Le mode transparent, lui, permet aux commutateurs de ne pas tenir compte des mises à jour VTP. Ils sont autonomes dans le domaine VTP et ne peuvent configurer que leurs VLAN (connectés localement). Cependant, ils transmettent aux autres commutateurs les mises à jour qu'ils reçoivent.

Les commutateurs en mode serveur et client mettent à jour leur base de données VLAN, si et seulement si, ils reçoivent une mise à jour VTP concernant leur domaine et contenant un numéro de révision supérieur à celui déjà présent dans leur base.

Fonction	Mode Serveur	Mode Client	Mode Transparent
Envoi de messages VTP	OUI	NON	NON
Réception des messages VTP ; Synchronisation VLAN	OUI	OUI	NON
Transmission des messages VTP reçus	OUI	OUI	OUI
Sauvegarde de configuration VLAN (en NVRAM ou Flash)	OUI	NON	OUI
Edition des VLANs (création, modification, suppression)	OUI	NON	OUI

Lorsqu'un hôte d'un VLAN envoie un broadcast, celui-ci est transmis à tous les commutateurs du domaine VTP. Il peut arriver que dans ce domaine, des commutateurs n'aient pas le VLAN concerné sur un de leur port.

Ce broadcast leur est alors destiné sans aucune utilité. Le **VTP pruning** empêche la propagation de ces trafics de broadcast aux commutateurs qui ne sont pas concernés.

19.3.2. Commandes associées

- **vlan database**
 - o Mode privilégié
 - o Permet d'accéder au mode de configuration de VLAN.
- **vlan vlan_id [name { nom du vlan }]**
 - o Mode de configuration de VLAN
 - o Permet de créer et nommer les VLANs.
- **vtp domain nom de domaine { password mot de passe | pruning | v2-mode | {server | client | transparent}}**
 - o Mode de configuration de VLAN
 - o Spécifie les paramètres VTP.
- **show vtp status**
 - o Mode privilégié
 - o Affiche la configuration VTP et le statut du processus.

20. Le routage

20.1. Principes fondamentaux

La couche réseau fournit un acheminement de bout en bout et au mieux des paquets à travers les réseaux interconnectés.

Le principe de base du routage est synthétisé par les deux étapes suivantes :

- **Détermination du chemin** : La couche réseau utilise une table de routage pour déterminer quel est le meilleur chemin à emprunter pour atteindre le réseau de destination. Le principe de métrique est utilisé afin d'offrir une mesure de qualité pour un chemin.
- **Commutation** : La fonction de commutation permet à un routeur d'accepter un paquet d'une interface et de le transmettre par le biais d'une autre interface. Le paquet pris en charge à une interface est retransmis via une autre interface représentant le meilleur chemin vers le réseau de destination.

Un protocole routé est un protocole de réseau dont l'adresse de couche réseau fournit suffisamment d'informations pour permettre d'acheminer un paquet d'une machine vers une autre, sur la base du modèle d'adressage. Les protocoles routés définissent le format des champs d'un paquet. Exemples de protocoles routés :

- IP
- IPX

Les protocoles de routage supportent un protocole routé en fournissant les mécanismes de partage des informations de routage. Les routeurs échangent les messages des protocoles de routage. Un protocole de routage permet aux routeurs de communiquer entre eux pour mettre à jour et gérer leurs tables. Exemples de protocoles de routage TCP/IP :

- RIP
- IGRP
- EIGRP
- OSPF
- BGP

Les routeurs peuvent supporter plusieurs protocoles de routage ainsi que plusieurs protocoles routés en même temps. Ces caractéristiques permettent à un routeur de distribuer les paquets de plusieurs protocoles routés sur les mêmes liaisons de données. Il est important de noter qu'il existe alors une table de routage par protocole routé.

Le processus de transmission de l'information se déroule comme suit (durant ce processus, les adresses de couche 3 ne changent pas) :

- L'hôte source détermine si la destination est en local ou distante grâce au couple IP/masque de sous-réseau. Elle calcule ainsi l'IP de sous-réseau de la destination ainsi que la sienne.
- Si les IP de sous-réseau sont les mêmes, alors la source émet la trame avec l'adresse de couche 2 de la destination. L'émission est ainsi directe.
- Par contre, si les IP de sous-réseau sont différentes, alors la source encapsule la trame avec l'adresse de couche 2 de sa passerelle par défaut puis l'envoie.
- La passerelle par défaut, à savoir généralement un routeur, reçoit cette trame. Elle va donc déterminer le chemin à emprunter afin d'atteindre le réseau de destination. Ceci se fait grâce aux informations de couche 3 fournies par le paquet ainsi que par la table de routage.

Il se pose ensuite deux cas :

- Le routeur actuel est le routeur final, c'est-à-dire qu'il est directement connecté au réseau de destination. Dans ce cas précis, on place les adresses de couche 2 de l'interface du routeur comme adresse source, et celle de la destination dans le champ adresse de destination. La trame est alors envoyée sur le réseau de destination.
- Le routeur actuel est un routeur intermédiaire sur le chemin, c'est-à-dire qu'il va falloir passer obligatoirement par un autre routeur afin d'atteindre le réseau de destination. La trame va donc être encapsulée avec l'adresse de couche 2 de l'interface de ce routeur, et celle du prochain saut dans le champ adresse de destination.

20.2. Routage statique et dynamique

Il existe deux types de routage :

- **Statique** : Tout est géré manuellement par un administrateur réseau qui enregistre toutes les informations dans la configuration d'un routeur. Il doit mettre à jour manuellement les entrées de route statique chaque fois qu'une modification de la topologie de l'interréseau le nécessite.
- **Dynamique** : Une fois qu'un administrateur réseau a entré les commandes de configuration pour lancer le routage dynamique, les informations relatives à la route sont mises à jour automatiquement, par un processus de routage, chaque fois que l'interréseau envoie de nouvelles informations.

Le routage statique offre plusieurs applications utiles :

- Le routage dynamique a tendance à révéler toutes les informations connues d'un interréseau, alors que vous souhaiteriez masquer certaines informations pour des raisons de sécurité. Le routage statique vous permet de spécifier les informations que vous souhaitez révéler à propos de réseaux restreints.
- Lorsqu'un réseau n'est accessible que par un seul chemin, une route statique vers ce réseau peut s'avérer suffisante. Ce type de réseau est appelé **réseau d'extrémité**. La configuration d'une route statique vers un réseau d'extrémité permet d'éviter la surcharge liée au routage dynamique.
- Il évite d'avoir une perte en bande passante due aux mises à jour envoyées par les protocoles de routage.

Une **route par défaut** est une entrée de table de routage qui dirige les paquets vers un saut suivant, lorsque ce dernier n'est pas inclus explicitement dans la table de routage. La route par défaut fait partie intégrante du routage statique. Ce type de route est utilisé par exemple pour rediriger les paquets d'un LAN vers Internet.

Le routage dynamique possède comme avantage principal de s'adapter automatiquement aux modifications topologiques.

La mise en œuvre du routage dynamique dépend de deux fonctions de base :

- La gestion d'une table de routage.
- La distribution opportune des informations aux autres routeurs sous la forme de mises à jour du routage.

Le routage dynamique s'appuie sur un protocole de routage pour partager les informations entre les routeurs. Un protocole de routage définit les règles utilisées par un routeur pour communiquer avec les routeurs voisins. Par exemple, un protocole de routage définit les informations suivantes :

- Comment envoyer les mises à jour.
- Les informations contenues dans ces mises à jour.
- Le moment où les informations doivent être envoyées.
- Comment localiser les destinataires des mises à jour.

Lorsqu'un algorithme de routage met à jour une table de routage, son principal objectif est de déterminer les meilleures informations à inclure dans cette table. Chaque algorithme de routage interprète à sa façon les meilleures informations. L'algorithme génère un nombre appelé valeur métrique pour chaque chemin du réseau. En général, plus ce nombre est petit, meilleur est le chemin.

On peut calculer les métriques en fonction d'une seule caractéristique de chemin. On peut aussi calculer des métriques plus complexes en combinant plusieurs caractéristiques. Les plus couramment utilisées par les routeurs sont les suivantes :

- **Bande passante** : Le débit d'une liaison, mesuré en bits par seconde.
- **Délai** : Le temps requis pour acheminer un paquet, pour chaque liaison, de la source à la destination.
- **Charge** : La quantité de trafic sur une ressource réseau telle qu'un routeur ou une liaison.
- **Fiabilité** : Cette notion indique généralement le taux d'erreurs sur chaque liaison du réseau.
- **Nombre de sauts (hop count)** : Le nombre de routeurs par lesquels un paquet doit passer avant d'arriver à destination.
- **Tics** : L'intervalle de temps sur une liaison donnée, en utilisant les signaux (tics) d'horloge d'un ordinateur personnel IBM (environ 55 millisecondes).
- **Coût** : Généralement basée sur une dépense monétaire attribuée à un lien par un administrateur réseau.

La métrique entre le port et chaque réseau directement connecté est égale à 0.

La plupart des algorithmes de routage peuvent être classés selon l'un des algorithmes suivants :

- Vecteur de distance.
- Etat de lien.
- Hybride symétrique.

20.3. Routage à vecteur de distance

La méthode de routage à vecteur de distance détermine la direction (vecteur) et la distance vers un lien de l'interréseau.

Les algorithmes de routage à vecteur de distance transmettent d'un routeur à l'autre des copies périodiques d'une table de routage.

Chaque routeur reçoit une table de routage des routeurs voisins auxquels il est directement connecté. Il rajoute les informations nouvelles ou meilleures dans la table de routage, sachant que les métriques distantes (celles incluses dans les mises à jour) sont cumulées avec la métrique séparant les deux routeurs, afin d'obtenir les nouvelles métriques. La nouvelle table de routage est alors envoyée dans toutes les directions, afin d'informer tous les routeurs des nouveautés.

Les algorithmes de routage à vecteur de distance ne permettent pas à un routeur de connaître la topologie exacte d'un interréseau. En effet, leur vue de la topologie est basée sur celle des voisins. C'est-à-dire que chaque routeur connaît l'existence de toutes les destinations ainsi que les sauts menant à eux, mais ne connaît pas le chemin exact, saut après saut jusqu'à la destination.

Les protocoles de routage suivants utilisent un algorithme à vecteur de distance :

- RIP
- IGRP

20.4. Routage à état des liens

La méthode de routage à état de lien recrée la topologie exacte de l'ensemble de l'inter réseau.

Les algorithmes de routage à état de liens, également appelés algorithmes du plus court chemin d'abord, gèrent une base de données topologique complexe. Ces informations comprennent tous les routeurs distants et leurs interconnexions.

Le routage à état de liens utilise les éléments suivants :

- des mises à jour déclenchées par modifications topologiques.
- une base de données topologiques.
- L'algorithme du plus court chemin d'abord.
- L'arbre du plus court chemin d'abord résultant.
- La table de routage, déduite à partir de l'arbre du plus court chemin d'abord.

Les étapes du processus de création de la table de routage sont :

- Réception d'une mise à jour, contenant des informations destinées à la base de données topologiques.
- L'algorithme du plus court chemin d'abord converti cette base de données en un arbre du plus court chemin d'abord, dont il est la racine.
- Cet arbre contient tous les chemins existants vers chaque destination connue.
- Le contenu de la table de routage est déterminé en parcourant l'arbre, sachant que l'on ne garde dans cette table que la meilleure entrée pour chaque destination.

Chaque fois qu'un paquet de mise à jour entraîne une modification dans la base de données, l'algorithme du plus court chemin d'abord recalcule les meilleurs chemins et met à jour la table de routage.

Le routage à état de liens est lié à deux exigences :

- **Ressource calculatoire** : Un protocole de routage à état des liens requière une puissance processeur importante pour son algorithme du plus court chemin d'abord, afin de transformer sa base de données topologiques en un arbre du plus court chemin d'abord, puis pour traiter cet arbre pour en déduire la table de routage.
- **Ressource mémoire** : Une grande quantité de mémoire RAM est utilisée par un protocole de routage à état des liens car il faut stocker la base de données topologiques ainsi que l'arbre du plus court chemin d'abord, en plus de la classique table de routage.

Le protocole OSPF est le plus connu des protocoles de routage à état des liens.

20.5. Convergence, problème associé et solutions

L'algorithme de routage est essentiel au routage dynamique. Chaque fois que la topologie d'un réseau change, la représentation du réseau doit également être modifiée. La représentation doit refléter une vue précise et cohérente de la nouvelle topologie. Cette vue est appelée convergence.

Lorsque tous les routeurs d'un interréseau utilisent les mêmes informations, l'interréseau est convergent. Une convergence rapide est recommandée pour un réseau, car elle réduit la période au cours de laquelle les routeurs prennent des décisions de routage incorrectes ou inefficaces.

Des boucles de routage peuvent se produire si la convergence lente d'un réseau avec une nouvelle configuration entraîne des entrées de routage incohérentes. Les paquets tournent sans cesse sur une boucle bien que le réseau de destination soit en panne. La métrique est incrémentée chaque fois que le paquet de mise à jour passe par un autre routeur.

Pour tenter de contrer les boucles de routages, il existe :

- La métrique de mesure infinie.
- Le Split Horizon.
- Les compteurs de retenue (Hold Downs).

On ne se préoccupe que de la table de routage avec ces trois solutions, car le problème des paquets en eux-mêmes est réglé automatiquement, et ce grâce au principe de TTL (Time To Live).

Une **métrique de mesure infinie** peut s'avérer nécessaire. Le principe est de définir l'infini en tant que nombre maximum spécifique. Ce nombre se réfère à une métrique de routage. Grâce à cette méthode, le protocole de routage permet à la boucle de routage d'exister jusqu'à ce que la métrique dépasse la valeur maximale autorisée. Le réseau en panne est considéré comme inaccessible lorsque la valeur métrique atteint la valeur maximale.

Le principe du **Split Horizon** est simple : aucune mise à jour ne sera envoyée par le chemin par lequel on a appris la modification de topologie. Ceci permet d'éviter de renvoyer à la source des informations erronées, et donc de limiter la propagation de mises à jour erronées.

On peut aussi utiliser des **compteurs de retenue** qui permettent d'éviter de changer l'état d'une entrée dans la table de routage impunément. Ils ont pour but de laisser le temps à l'information d'atteindre l'intégralité du réseau avant de modifier de nouveau la même entrée.

Ils fonctionnent de la façon suivante :

- Lorsqu'une modification est effectuée sur une entrée de la table de routage, on lance un compteur de retenue pour cette entrée.
- Si une mise à jour contenant une modification pour cette entrée alors que le temps du compteur de retenue est dépassé, alors la modification est appliquée.
- Si une mise à jour contenant une modification pour cette entrée pendant le temps du compteur de retenue, alors le protocole suivra les règles imposées par le principe des compteurs de retenue.

Les règles imposées par le principe des compteurs de retenue sont les suivantes :

- On autorise l'activation ou l'amélioration de qualité (métrique) pour une entrée.
- On refuse la désactivation ou la dégradation de qualité pour l'entrée concernée.

Tous les algorithmes peuvent être configurés afin d'utiliser ces trois méthodes, à l'exception des compteurs de retenue qui sont inutiles pour les protocoles de routage à état des liens, car le temps de convergence est faible.

20.6. Contexte des différents algorithmes de routage

On peut comparer plusieurs aspects fondamentaux du routage à vecteur de distance et du routage à état des liens :

Vecteur de distance	Etat des liens
Vue de la topologie du réseau à partir de la perspective des voisins	Vue commune et complète de l'ensemble de la topologie du réseau
Ajout de vecteurs de distance d'un routeur à l'autre	Calcul du chemin le plus court menant aux autres routeurs grâce à l'arbre du plus court chemin d'abord
Mises à jour périodiques fréquentes : Convergence lente	Mises à jour déclenchées par événements : Convergence plus rapide
Transmission des copies des tables de routage aux routeurs voisins	Transmission des informations de la base de données topologiques

Un troisième type d'algorithme combine les aspects du routage à vecteur de distance et du routage à état de liens. Ce troisième type est appelé routage hybride symétrique.

Les caractéristiques des protocoles de routage hybride symétrique sont :

- Utilisation du principe de base du vecteur de distance (légèreté de la gestion de la table de routage et contenu des mises à jour identique).
- Utilisation des mises à jour déclenchées par modification topologique de l'état des liens (convergence plus rapide que les protocoles de routage à vecteur de distance).

Les protocoles de routage suivants utilisent un algorithme de routage hybride symétrique :

- IS-IS normalisé par l'ISO
- EIGRP de Cisco

20.7. Configuration initiale du routeur

La commande **ip routing** du mode de configuration globale permet d'activer la commutation des paquets entre les interfaces. Cette commande est donc la première à utiliser lors de la configuration du routage, qu'il soit statique ou dynamique.

La commutation des paquets entre les interfaces est désactivée par défaut pour tous les protocoles routés.

La commande **ip route {network} [masque] {adresse ou interface} [distance]** permet de définir une route statique :

- Le paramètre **masque** permet de forcer un masque de sous-réseau pour le paramètre **network**. Il s'agit d'un paramètre optionnel car le masque peut être implicite (déduit grâce au principe de classes d'IP) ou explicite (il doit être précisé dans le cas du subnetting).
- Le paramètre **adresse** correspond à l'adresse IP de l'interface du prochain saut, alors que **interface** fait référence au nom de l'interface locale à utiliser pour atteindre le réseau de destination.
- Enfin, le paramètre **distance** permet de spécifier une mesure de la fiabilité d'une source d'information de routage, exprimée sous forme de valeur numérique comprise entre 1 et 255. Plus la valeur est élevée, plus la fiabilité de la source est faible.

La commande **ip default-network {network}** permet de définir une route par défaut. Un réseau par défaut doit exister dans une table de routage, afin de pouvoir rediriger les paquets destinés à un réseau extérieur au notre (exemple : Internet). Il ne doit pas y avoir plus d'une route par défaut par routeur.

La commande **ip default-network** doit être ajoutée sur tous les routeurs du réseau ou elle doit être utilisée avec la commande **redistribute static** de façon à ce que tous les réseaux connaissent le réseau candidat par défaut.

20.8. Protocoles de routage intérieurs et extérieurs

Un système autonome est un ensemble de routeurs fonctionnant suivant la même administration, par exemple le même protocole de routage.

Le NIC (Network Information Center) attribue aux entreprises un numéro de système autonome unique. Ce numéro est un nombre à 16 bits. Un protocole de routage, tel que le protocole IGRP de Cisco, exige que l'on indique ce numéro unique dans notre configuration.

Les protocoles de routage extérieurs (EGP : Exterior Gateway Protocol) permettent aux systèmes autonomes de communiquer entre eux. Les protocoles de routage intérieurs (IGP : Interior Gateway Protocol) sont utilisés au sein d'un système autonome.

Les différents IGP sont :

- RIP
- IGRP
- EIGRP
- OSPF

Le seul EGP que nous allons citer est :

- BGP

La commande **router {protocole} [option]** lance le processus de routage dynamique en utilisant le **protocole** spécifié (RIP, IGRP, OSPF ou EIGRP). Le paramètre **option** sert à spécifier le système autonome pour les protocoles qui le requièrent. Après avoir entré cette commande, on se retrouve dans le mode de configuration de ce protocole de routage.

Ensuite, dans ce mode de configuration de ce protocole de routage, il faut indiquer les différents réseaux directement connectés qui seront inclus dans les mises à jour de routage. Cela se fait par le biais de la commande **network {numéro de réseau}** dans le mode de configuration du protocole de routage.

21. Protocole RIP

21.1. Théorie

RIP (Routing Information Protocol) est un protocole de routage à vecteur de distance. Il existe en deux versions :

- **RIPv1** (RFC 1058) : Première version du protocole RIP.
- **RIPv2** (RFC 1723) : Evolution permettant le routage Classless (en transmettant les masques de sous-réseaux en plus des préfixes dans les mises à jour) et la transmission des mises à jour en multicast.

RIPv1	RIPv2
Classful	Classless
Broadcast pour les mises à jour	Multicast (224.0.0.9) pour les mises à jour
Préfixes dans les mises à jour	Préfixes et masques de sous-réseau dans les mises à jour
	Support du VLSM
	Authentification des voisins

Les caractéristiques principales de RIP sont :

- Nombre de sauts (hop count) utilisé pour le calcul des métriques.
- Métrique maximale = 15 (métrique de mesure infinie = 16).
- Mises à jour périodiques toutes les 30 secondes.

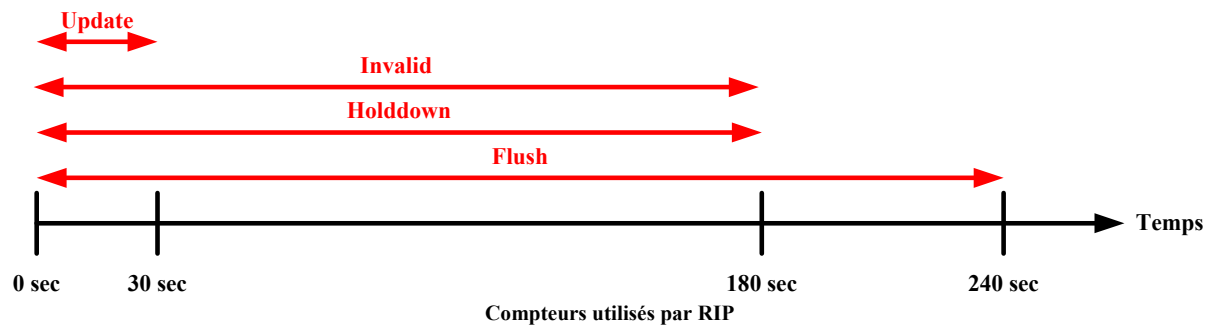
Avantages	Inconvénients
Processus léger	Temps de convergence lent
Implémenté sur tous les systèmes d'exploitation	Nombre de sauts pour calculer les métriques
	Nombre de sauts limité à 15

RIP n'a pas de notion de système autonome. Ceci signifie qu'il ne connaît rien d'autre que lui-même. Le seul moyen de pouvoir sortir du système autonome RIP est par conséquent une route statique par défaut.

L'implémentation Cisco de RIP supporte les mises à jour déclenchées. De plus, les caractéristiques de ce protocole font de RIP le protocole de prédilection pour les réseaux LAN homogènes de petite taille.

En tant que protocole de routage à vecteur de distance, RIP utilise quatre compteurs :

- **Update** : Intervalle de temps entre les mises à jour périodiques (30 secondes par défaut).
- **Invalid** : Intervalle de temps après réception de la dernière mise à jour pour chaque entrée dans la table de routage avant de la considérer comme périmée. Après ce temps, l'entrée concernée ne sera plus analysée lors du parcours de la table de routage (180 secondes par défaut).
- **Holddown** : Intervalle de temps après réception de la dernière mise à jour avant d'autoriser le remplacement de cette route par une autre moins bonne (180 secondes par défaut).
- **Flush** : Intervalle de temps après réception de la dernière mise à jour pour chaque entrée dans la table de routage avant de la supprimer de la table de routage (240 secondes par défaut).



21.2. Configuration

21.2.1. Commandes

Les commandes liées à la configuration du protocole RIP sont :

- **router rip**
 - o Mode de configuration globale
 - o Active le protocole RIP
 - o Passe dans le mode de configuration du routeur
- **network {préfixe}**
 - o Mode de configuration du routeur
 - o Spécifie le réseau qui sera inclut dans les mises à jour de routage
 - o Détermine les interfaces appartenant à ce réseau qui participent au processus de routage
 - o Le **préfixe** doit être un réseau directement connecté au routeur
- **neighbor {IP}**
 - o Mode de configuration du routeur
 - o Définit l'adresse IP d'un voisin avec lequel RIP échangera des mises à jour de routage
 - o Par défaut, aucun voisin n'est défini
- **passive-interface {type} {numéro}**
 - o Mode de configuration du routeur
 - o Empêche l'interface indiquée d'envoyer des mises à jour
- **[no] ip split-horizon**
 - o Mode de configuration d'interface
 - o Active/désactive Split Horizon sur l'interface courante
- **timers basic {update} {invalid} {holddown} {flush}**
 - o Mode de configuration du routeur
 - o Définit les intervalles de temps, en secondes, utilisés par RIP

- **default-information originate**
 - o Mode de configuration du routeur
 - o Propage le réseau candidat par défaut aux autres routeurs RIP du système autonome
- **maximum-paths {nombre}**
 - o Mode de configuration du routeur
 - o Spécifie le nombre maximum de liens ayant la même métrique pouvant être utilisés pour la répartition de charge
 - o Par défaut à 4 et maximum à 6 ou 16 (IOS >= 12.3(2)T)
- **redistribute static**
 - o Mode de configuration du routeur
 - o Injecte les routes statiques locales et les propagent dans les mises à jour RIP

22. Protocole RIPv2

22.1. Spécifications de RIPv2

RIPv2 est une version améliorée de son prédécesseur et partage donc certaines caractéristiques :

- Tous deux sont des IGP (Interior Gateway Protocol).
- RIPv1 et RIPv2 sont des protocoles de routage à vecteur de distance.
- Ils utilisent une métrique basée sur le nombre de saut.
- Ils emploient un nombre maximum de saut, des compteurs de retenue d'on la valeur est fixé à 180s par défaut, ainsi que le split horizon et le route poisoning pour limiter les effets de boucles de routage.

RIPv2 apporte également des fonctionnalités supplémentaires tels que :

- Le support du routage classless et du VLSM
- La diffusion du masque réseau dans les mises à jour de routage.
- La diffusion des mises à jour de routage par multicast avec l'adresse de classe D 224.0.0.9.
- L'authentification de la source de la mise à jour de routage par un texte en clair (**actif par défaut**), ou un texte crypté suivant l'algorithme MD5 (Message-Digest 5).
- L'utilisation d'indicateurs de route externe (**route tag**) afin de pouvoir différencier les routes apprises d'autre protocole de routage et redistribué dans RIP.

22.2. Configuration

22.2.1. Commandes générales

- **router rip**
 - o Mode de configuration globale
 - o Active le protocole RIP.
- **version 2**
 - o Mode de configuration du protocole de routage
 - o Permet d'utiliser RIPv2 à la place de RIPv1
- **ip rip {send | receive} version {1 | 2 | 1 2}**
 - o Mode de configuration d'interface
 - o Spécifie précisément le type (RIPv1 et/ou RIPv2) de mises à jour envoyées ou reçues
- **network {adresse réseau}**
 - o Mode de configuration du protocole de routage
 - o Permet d'indiquer les réseaux directement connectés au routeur.
- **ip default-network {adresse réseau}**
 - o Mode de configuration globale
 - o Permet de spécifier une route par défaut.
- **default-information originate**
 - o Mode de configuration du protocole de routage

- o Permet de propager la route par défaut dans les mises à jour de routage.
- **no auto-summary**
 - o Mode de configuration du protocole de routage
 - o Désactive l'auto-agrégation.

22.2.2. Authentification

- **key-chain {nom}**
 - o Mode de configuration globale
 - o Permet d'identifier un groupe de clef d'authentification.
- **key {id}**
 - o Mode de configuration de clé
 - o Permet de créer une clef dans un groupe de clef. L'identifiant de clef peut prendre une valeur de 0 à 2147483647. L'identifiant de clef peut ne pas être consécutif.
- **key-string {mot de passe}**
 - o Mode de configuration de clé
 - o Permet de définir un mot de passe pour une clef.
- **ip rip authentication key-chain {nom}**
 - o Mode de configuration d'interface
 - o Active l'authentification RIP sur une interface
- **ip rip authentication mode {text | md5}**
 - o Mode de configuration d'interface
 - o Permet de spécifier le type d'authentification en clair ou crypté.

23. Protocole IGRP

23.1. Théorie

IGRP (Interior Gateway Routing Protocol) est un protocole de routage à vecteur de distance propriétaire Cisco. Il a été conçu au milieu des années 1980 pour remplacer RIP. En effet, des incohérences de routage peuvent survenir avec RIP sur des réseaux hétérogènes.

IGRP est donc capable de fonctionner sur des réseaux hétérogènes de très grande taille, tout en proposant un calcul des métriques basé sur les critères suivants :

- Bande passante
- Délai
- Fiabilité
- Charge

Les métriques IGRP sont des nombres sur 24 bits (de 0 à 16 777 215) calculés à l'aide de cette formule :

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

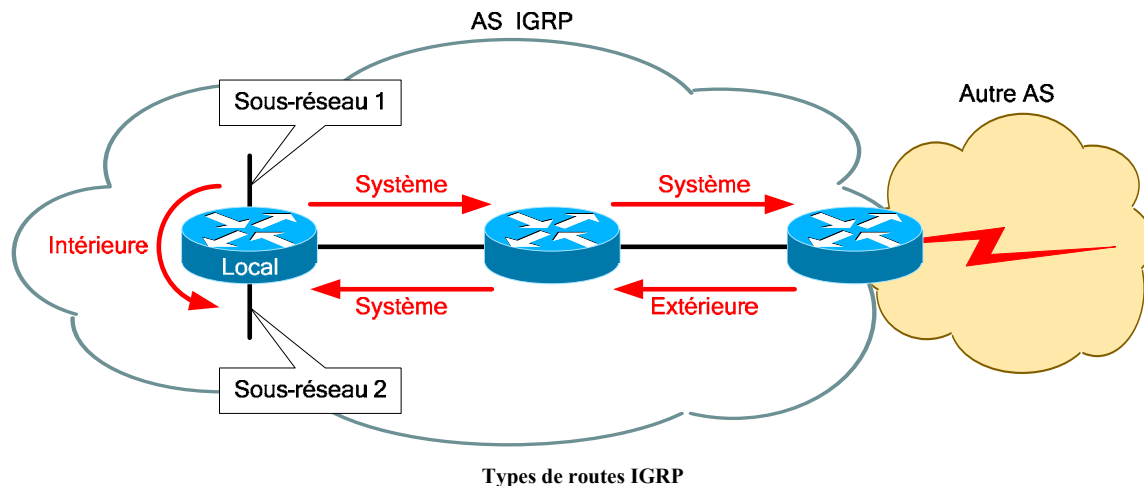
Les différents paramètres de cette formule sont les suivants :

- **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 $10^7 \div \text{BP}$, 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 $\Sigma_{\text{délais}}$.
- **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 :

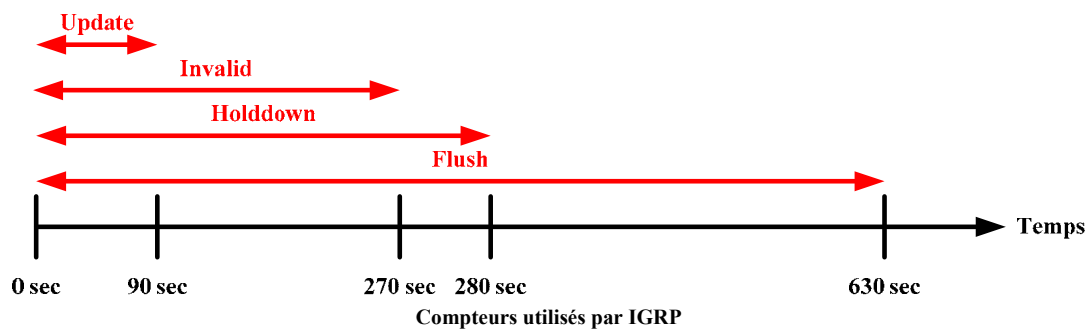
$$\text{Métrique} = \text{Bandwidth} + \text{Delay}$$

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



Il peut y avoir jusqu'à 4 routes pour une même destination dans la table de routage. Ces routes peuvent être de 3 types :

- **Intérieure** : Route entre des sous-réseaux directement connectés au routeur local.
- **Système** : Route interne au système autonome propagée par un routeur.
- **Extérieure** : Route externe à l'AS qui a été redistribuée dans l'AS IGRP (inclus aussi les routes statiques redistribuées).



En tant que protocole de routage à vecteur de distance, IGRP utilise quatre compteurs :

- **Update** : Intervalle de temps entre les mises à jour périodiques (90 secondes par défaut).
- **Invalid** : Intervalle de temps après réception de la dernière mise à jour pour chaque entrée dans la table de routage avant de la considérer comme périmée. Après ce temps, l'entrée concernée ne sera plus analysée lors du parcours de la table de routage (270 secondes par défaut, ou 3 fois l'Update).
- **Holddown** : Intervalle de temps après réception de la dernière mise à jour avant d'autoriser le remplacement de cette route par une autre moins bonne (280 secondes par défaut).
- **Flush** : Intervalle de temps après réception de la dernière mise à jour pour chaque entrée dans la table de routage avant de la supprimer de la table de routage (630 secondes par défaut, ou 7 fois l'Update).

IGRP utilise aussi les mises à jour Poison Reverse. Ceci permet de placer des routes directement à l'état Holddown. Toute route dont la métrique augmentant d'un facteur de 1,1 fera l'objet d'une mise à jour Poison Reverse.

23.2. Configuration

23.2.1. Commandes

Les commandes pouvant être utilisées pour la configuration du protocole IGRP sont les suivantes :

- **router igrp {AS}**
 - o Mode de configuration globale
 - o Active le protocole de routage IGRP sur le routeur pour le système autonome indiqué en paramètre
 - o Permet de passer dans le mode de configuration du routeur
- **network {préfixe}**
 - o Mode de configuration du routeur
 - o Spécifie le réseau qui sera inclut dans les mises à jour de routage
 - o Détermine les interfaces appartenant à ce réseau qui participent au processus de routage
 - o Le **préfixe** doit être un réseau directement connecté au routeur.
- **neighbor {IP}**
 - o Mode de configuration du routeur
 - o Définit l'adresse IP d'un voisin avec lequel IGRP échangera des mises à jour de routage
 - o Par défaut, aucun voisin n'est définit
- **passive-interface {type} {numéro}**
 - o Mode de configuration du routeur
 - o Empêche l'interface indiquée d'envoyer des mises à jour
- **[no] ip split-horizon**
 - o Mode de configuration d'interface
 - o Active/désactive Split Horizon sur l'interface courante
- **maximum-paths {nombre}**
 - o Mode de configuration du routeur
 - o Spécifie le nombre maximum de liens ayant la même métrique pouvant être utilisés pour la répartition de charge
 - o Par défaut à 4 et maximum à 6 ou 16 (IOS >= 12.3(2)T)
- **variance {valeur}**
 - o Mode de configuration du routeur
 - o Permet la répartition de charge entre des liens n'ayant pas la même métrique
 - o **valeur** est un entier pouvant aller de 1 à 128 (défaut = 1)
 - o La variance est un coefficient multiplicateur permettant de sélectionner les routes ayant des métriques identiques à la variance près pour faire de la répartition de charge pondérée (Weighted Round Robin)
- **metric weights {TOS} {K1} {K2} {K3} {K4} {K5}**
 - o Mode de configuration du routeur
 - o Spécifie les valeurs pour les coefficients utilisés pour le calcul des métriques.
 - o **TOS** doit toujours être à 0

- **timers basic {update} {invalid} {holddown} {flush}**
 - Mode de configuration du routeur
 - Définit les intervalles de temps, en secondes, utilisés par IGRP
- **metric maximum-hops {valeur}**
 - Mode de configuration du routeur
 - Indique le nombre maximum de sauts (diamètre du système autonome)
 - **valeur** peut aller de 1 à 255 (défaut = 100)
- **ip default-network {préfixe}**
 - Mode de configuration globale
 - Définit un réseau candidat par défaut à propager dans le système autonome
 - Le réseau indiqué doit être connu des routeurs IGRP et doit être directement connecté
 - La route propagée sera vue par les autres routeurs IGRP comme une route externe
- **redistribute static**
 - Mode de configuration du routeur
 - Injecte les routes statiques locales et les propagent dans les mises à jour IGRP
- **bandwidth {BP}**
 - Mode de configuration d'interface
 - Définit la bande passante de la liaison
 - Cette valeur est utilisée par IGRP et EIGRP pour le calcul de leurs métriques.
 - Le paramètre BP est exprimé en Kbps

23.3. Vérification

Comme pour RIP, IOS fournit des commandes de visualisation d'état et de débogage pour IGRP :

- **show ip protocols** : Affiche les différentes instances d'IGRP, avec leur numéro d'AS, les compteurs, les coefficients utilisés pour le calcul des métriques, les réseaux avertis ainsi que les interfaces participant au processus de routage.
- **debug ip igrp events** : Affiche en temps réel les événements d'IGRP.
- **debug ip igrp transactions** : Affiche en temps réel les échanges d'IGRP.

24. Protocole OSPF

24.1. Caractéristiques

Le protocole **OSPF** (Open Shortest Path First) est un protocole de routage à état de lien créé en 1988 par l'IETF (RFC 2328). C'est à l'heure actuelle l'**IGP** (Interior Gateway Protocol) le plus répandu. OSPF est un protocole libre.

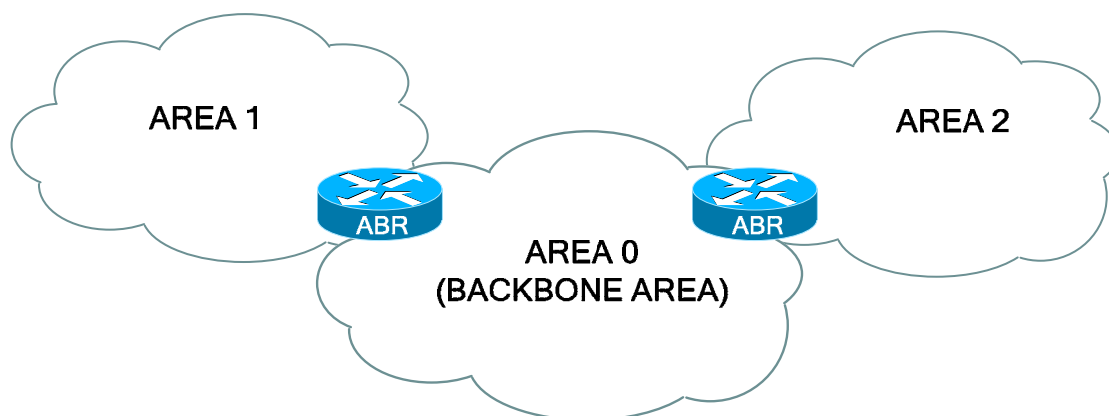
Principales caractéristiques d'OSPF :

- Emission des mises à jour déclenchées par modification(s) topologique(s).
- Connaissance exacte et complète de la topologie du réseau.
- Chaque nœud connaît l'existence de ses voisins adjacents.
- Utilisation d'un arbre du plus court chemin d'abord (SPF Tree) et d'un algorithme du plus court chemin d'abord (Algorithme SPF appelé aussi l'algorithme de Dijkstra) pour générer la table de routage.
- Envoi des mises à jour topologiques via une adresse multicast et non broadcast.
- Utilisation moindre de la bande passante
- Protocole de routage classless supportant le VLSM.
- Requiert des routeurs plus puissants.
- Domaines de routage exempts de boucles de routage
- Métrique utilisée : le coût (chaque liaison a un coût).
- Détermination et utilisation d'un ou plusieurs domaines de routage appelés Areas (ou aires) au sein d'un même système autonome (AS).

Les interfaces OSPF distinguent quatre types de réseaux :

- Les réseaux multi-accès broadcast comme Ethernet.
- Les réseaux point-à-point.
- Les réseaux multi-accès non broadcast ou encore Nonbroadcast multi-access (NBMA), tel que Frame Relay.
- Les réseaux point-à-multipoint configuré manuellement par un administrateur

L'établissement de la base de données topologique, ainsi que le calcul du plus court chemin d'abord impose une grande charge de traitements pour chaque routeur. Pour diminuer la taille de la base donnée topologique, les routeurs peuvent être regroupés en plusieurs aires (**area**) au sein d'un même système autonome (**SA**). On parle alors de **multiple area OSPF** (voir schéma ci-dessous), mais le cursus CCNA 3 ne s'attarde que sur l'emploi de **single area OSPF**.



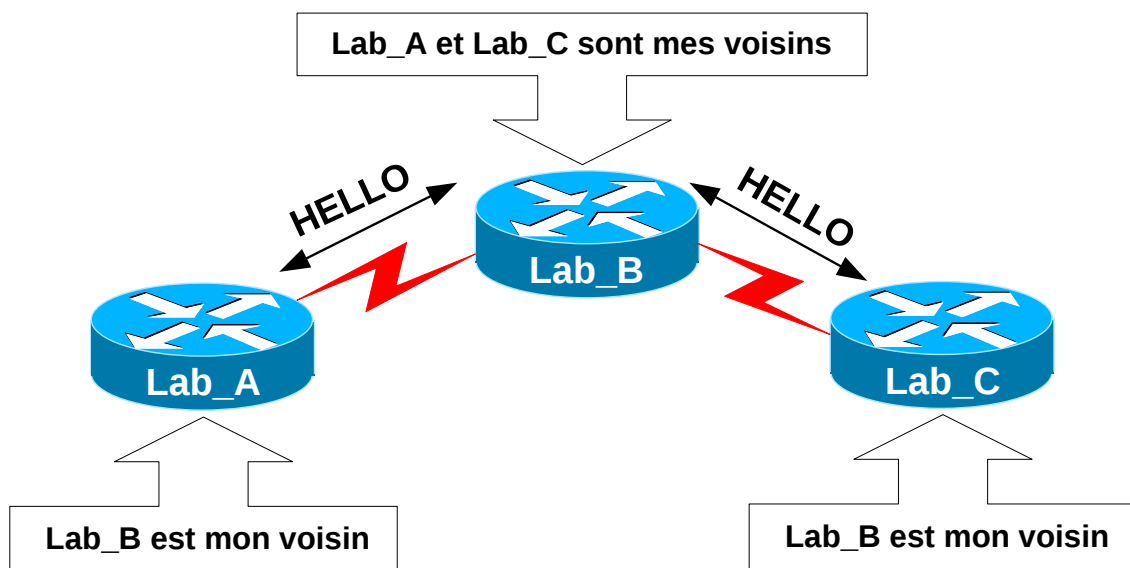
24.2. Définitions

- **Neighbor**
 - Routeur voisin sur le même réseau.
- **HELLO**
 - Protocole permettant la découverte et le maintien de liens entre les voisins. Les paquets HELLO sont transmis toutes les 10s pour un réseau de type broadcast multi-access et toutes les 30s pour un réseau de type NBMA.
- **LSU**
 - Paquet de mise à jour de données topologique. Permet d'avoir des informations sur l'évolution topologique du réseau.
- **LSA**
 - Contenu dans les LSUs ils permettent d'avertir qu'une modification topologique a lieu.
- **SPF tree**
 - L'arbre du plus court chemin d'abord résultant de l'application de l'algorithme de Dijkstra.
- **Algorithme de Dijkstra**
 - L'algorithme de Dijkstra (ou algorithme SPF), publié par le scientifique allemand du même nom en 1959 est utilisé pour le calcul de l'arbre du plus court chemin d'abord.
- **Adjacencies database**
 - Base de données contenant les informations relatives aux voisins.
- **Topological database**
 - Base de données qui contient toutes les informations sur la topologie du réseau.
- **Routing table**
 - Table de routage avec les meilleures routes à destination de tous les sous-réseaux de la topologie.
- **Flooding**
 - Processus qui consiste à envoyer par tous les ports.
- **DR (Designated Router)**
 - Routeur élu pour centraliser toutes les informations topologiques.
- **BDR (Backup Designated Router)**
 - Routeur élu pour prendre le relais du DR en cas de panne.
- **NBMA (Non Broadcast Multi-access)**
 - Réseau multi-accès Non broadcast tel que Frame Relay.
- **ABR (Area Border Router)**
 - Routeur situé à la bordure d'une ou plusieurs aires.

24.3. Fonctionnement dans un réseau ne comportant qu'une aire

24.3.1. Découverte des routeurs voisins

Avant tout échange d'informations de données topologiques, le routeur implémentant OSPF doit s'assurer qu'il existe d'autres routeurs adjacents à celui-ci qui utilisent eux aussi OSPF. Ces routeurs adjacents sont appelés des « voisins » et chacun d'entre eux peut être voisin d'un ou de plusieurs routeurs.



Pour découvrir leurs voisins, chaque routeur utilisant OSPF comme protocole de routage va devoir recourir au protocole **HELLO** qui permet d'établir et de maintenir un échange avec les routeurs voisins.

Celui-ci va permettre à chaque routeur d'envoyer des paquets HELLO à intervalles réguliers sur chacune de leurs interfaces en utilisant l'adresse multicast **224.0.0.5**. Les voisins découverts seront ensuite enregistrés dans une base de données de voisinage appelée **Neighbor Database**.

24.3.2. Etablissement des bases de données topologiques

24.3.2.1. Dans un réseau point-à-point

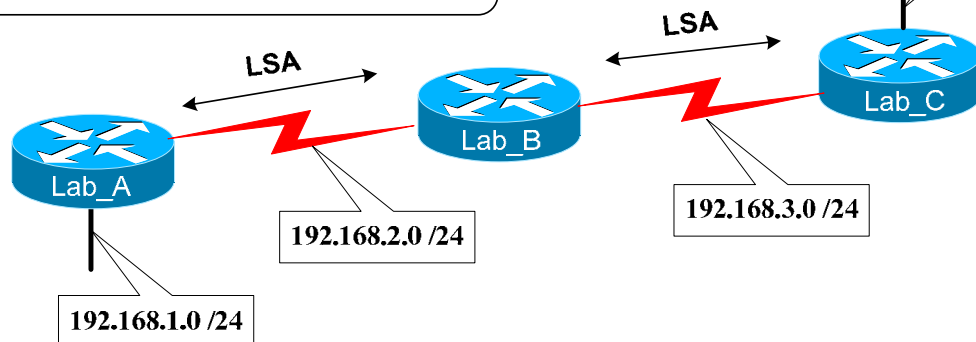
Une fois que chaque routeur a appris l'existence de ses voisins, il va leur envoyer les informations concernant tous les réseaux directement connectés à celui-ci.

Ces informations envoyées vont permettre à chaque nœud de mettre rapidement à jour leur base de données topologique (**Topological Database**) et d'obtenir ainsi une connaissance complète de la topologie réseau.

Ces mises à jour topologiques, déclenchées à l'initialisation du protocole OSPF sur les routeurs et par la suite lors de chaque modification topologique, se font grâce à l'envoi de paquets LSU (Link State Update) contenant des LSA (Link State Advertisement) comme le montre le schéma ci-dessous.

Informations envoyées par **Lab_A** à **Lab_B** :

« J'ai accès aux réseaux **192.168.1.0** et **192.168.2.0** »



24.3.2.2. Dans un réseau multiaccès

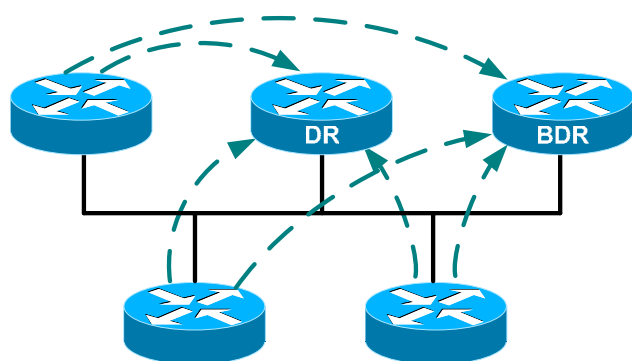
Les réseaux multiaccès fonctionnent suivant le même principe que les réseaux point-à-point à la différence que dans les réseaux multiaccès tous les routeurs sont voisins.

Cela pose cependant un problème puisque chaque routeur maintient un lien avec tous ses voisins pour l'échange d'informations topologiques. Par conséquent plus il y a de routeurs sur le réseau, plus ces derniers sont sollicités à envoyer des paquets de mises à jour topologiques.

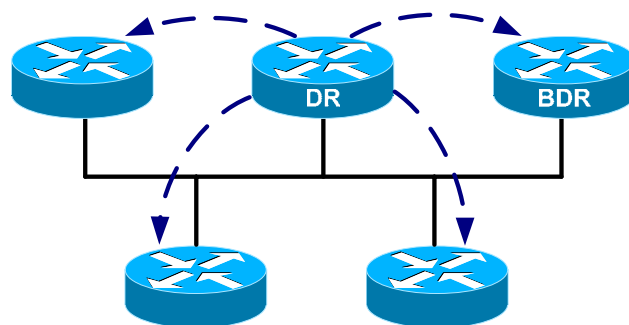
Pour palier à ce problème, le protocole HELLO va élire un **DR** (Designated Router) qui sera chargé de centraliser toutes les informations de modifications topologiques et de les retransmettre par la suite à tous les autres routeurs.

Il y aura ensuite l'élection d'un **BDR** (Backup Designated Router) servant de secours au cas où le DR ne pourrait plus assurer son rôle.

Tous les routeurs transmettront donc leurs informations topologiques au DR (ainsi qu'au BDR) en utilisant l'adresse multicast 224.0.0.6, tandis que le DR redistribuera ces informations avec l'adresse multicast 224.0.0.5 à tous les autres routeurs comme indiqué ci-dessous.



Envoi de paquets LSA au DR et au BDR en utilisant l'adresse multicast 224.0.0.6



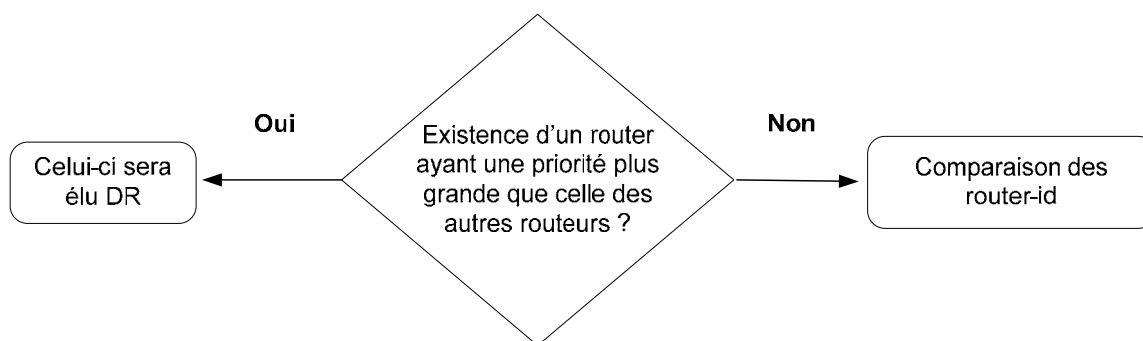
Envoi de paquets LSA à tous les routeurs en utilisant l'adresse multicast 224.0.0.5

24.4. Opérations OSPF

24.4.1. Election du DR / BDR

Un routeur doit répondre à plusieurs critères pour être désigné DR dans le réseau multi-accès. L'élection se fait grâce aux paquets HELLO qui contiennent l'ID du routeur et une priorité.

Lors du processus d'élection, le routeur ayant la plus grande priorité sur le réseau multi-accès sera élu DR. Dans le cas d'une égalité des priorités, les routeurs devront comparer leur router-id. Le routeur qui aura dans ce cas le plus grand router-id sera élu DR.



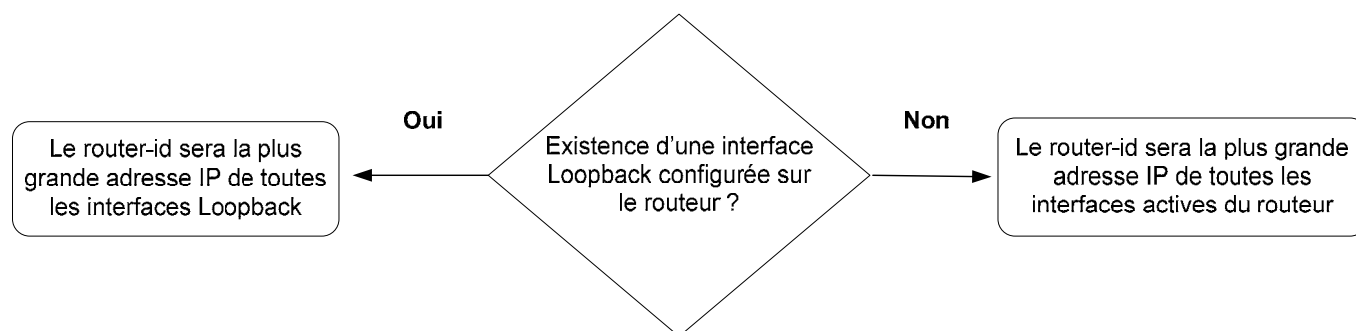
Une fois le DR désigné, le processus d'élection devra ensuite déterminer le BDR, correspondant au routeur ayant la deuxième plus haute priorité ou le deuxième plus grand router-id sur le réseau multi-accès.

24.4.2. Détermination du Router-ID

Lorsqu'une instance OSPF est initialisée, un identifiant de routeur appelé router-id est déterminé. Ce router-id n'est autre qu'une adresse IP qui servira d'identifiant à un routeur sur les réseaux auxquels il est raccordé.

Le router-id est déterminé selon les critères suivant :

- S'il y a présence d'une ou plusieurs interfaces Loopback sur le routeur, son router-id correspondra à la plus grande adresse IP de toutes les interfaces Loopback configurées sur celui-ci.
- Si aucune interface Loopback n'est présente sur le routeur alors son router-id sera la plus grande adresse IP de toutes les interfaces actives configurées sur celui-ci.



Pour fonctionner, un processus OSPF nécessite qu'il y ait au moins une interface active configurée sur le routeur. Il est donc conseillé, pour éviter toute interruption du processus OSPF, de faire usage des interfaces Loopback lorsque l'on configure ce protocole de routage sur un équipement.

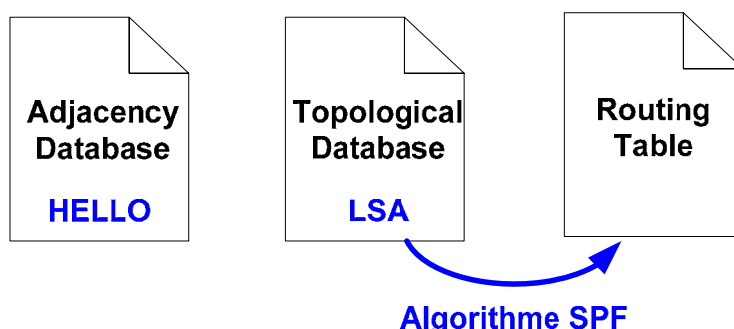
24.5. Construction de la table de routage

Une fois que tous les routeurs ont convergé, c'est-à-dire qu'ils ont tous la même vue complète du réseau, chacun d'entre eux va construire, à partir de sa base de données topologique, un arbre du plus court chemin d'abord (**SPF Tree**).

Cette construction va être réalisée grâce à l'algorithme **SPF** (Shortest Path First), aussi appelé l'algorithme de Dijkstra, qui va parcourir la base de données topologique et considérer chaque routeur comme étant des sommets reliés par des liens point-à-point. Le routeur qui l'implémente sera placé à la racine de l'arbre du plus court chemin d'abord.

La métrique utilisée par OSPF étant le coût, calculée par les composants Cisco à l'aide de la formule suivante : **coût=10⁸/bande passante** (s'exprime en bps), chaque lien va donc avoir un coût. La métrique d'une route est par conséquent calculée en faisant la somme de la bande passante de chaque lien de la route.

L'algorithme de Dijkstra va parcourir ensuite cet arbre du plus court chemin afin de déterminer les meilleures routes pour atteindre chaque réseau de destination (routes dont le coût est le plus bas). Ces routes seront ensuite ajoutées à la table de routage.



Au niveau de la table de routage, chaque route apprise par le protocole de routage OSPF se manifestera par la lettre « O » devant celle-ci et aura une distance administrative de 110.

24.6. Commandes

24.6.1. Commandes générales

- **router ospf {id de processus}**
 - o Mode de configuration globale
 - o Active le protocole OSPF.
 - o Plusieurs processus peuvent être lancés sur un routeur.
- **network {préfixe}**
 - o Mode de configuration du routeur
 - o Permet de spécifier les réseaux devant participer au processus de routage.
 - o Le préfixe doit être un réseau directement connecté au routeur
- **interface loopback {number}**
 - o Mode de configuration globale
 - o Permet de créer une interface logique.
- **bandwidth**
 - o Mode de configuration d'interface
 - o Permet de spécifier la bande passante sur l'interface.
- **ip ospf priority {number}**
 - o Mode de configuration d'interface
 - o Permet de modifier la priorité d'une interface pour l'élection du DR.
 - o La valeur peut aller de 0 à 255. Attention, une priorité de 0 empêche le routeur d'être élu DR.
- **ip ospf cost {number}**
 - o Mode de configuration d'interface
 - o Permet de spécifier la valeur du coût.

24.6.2. Authentification

- **area {numéro de l'aire} authentication**
 - o Mode de configuration du routeur
 - o Active l'authentification OSPF pour le mot de passe en clair.
- **area {numéro de l'aire} authentication message-digest**
 - o Mode de configuration du routeur
 - o Active l'authentification pour le mot de passe encrypté.
- **ip ospf message-digest-key {key-id} md5 {type d'encryption}**
 - o Mode de configuration d'interface
 - o Permet l'encryption du mot de passe.
- **ip ospf authentication-key {mot de passe}**
 - o Mode de configuration d'interface

- o Spécifie le mot de passe utilisé pour générer les données d'authentification de l'entête de paquets OSPF.

24.6.3. Timers

- **ip ospf hello-interval {intervalle}**
 - o Mode de configuration d'interface
 - o Définit la fréquence d'émission des paquets HELLO.
- **ip ospf dead-interval {intervalle}**
 - o Mode de configuration d'interface
 - o Définit la durée pendant laquelle un lien sera considéré comme actif, après que le routeur est reçu un paquet HELLO d'un routeur voisin.

24.6.4. Commandes show associées

- **show ip ospf interface**
 - o Mode privilégié
 - o Permet d'afficher la priorité de l'interface.
- **show ip protocols**
 - o Mode privilégié
 - o Affiche les informations sur les protocoles de routage configurés sur le routeur.
- **show ip route**
 - o Mode privilégié
 - o Affiche la table de routage du routeur.
- **show ip ospf**
 - o Mode privilégié
 - o Affiche la durée pendant laquelle le protocole est activé, ainsi que la durée durant laquelle il n'y a pas eu de modification topologique.
- **show ip ospf neighbor detail**
 - o Mode privilégié
 - o Affiche une liste détaillée des voisins, leur priorité et leur statut.
- **show ip ospf database**
 - o Mode privilégié
 - o Affiche le contenu de la base de données topologique (router-Id, process-Id).

25. Protocole EIGRP

25.1. 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
- Distance administrative de 90

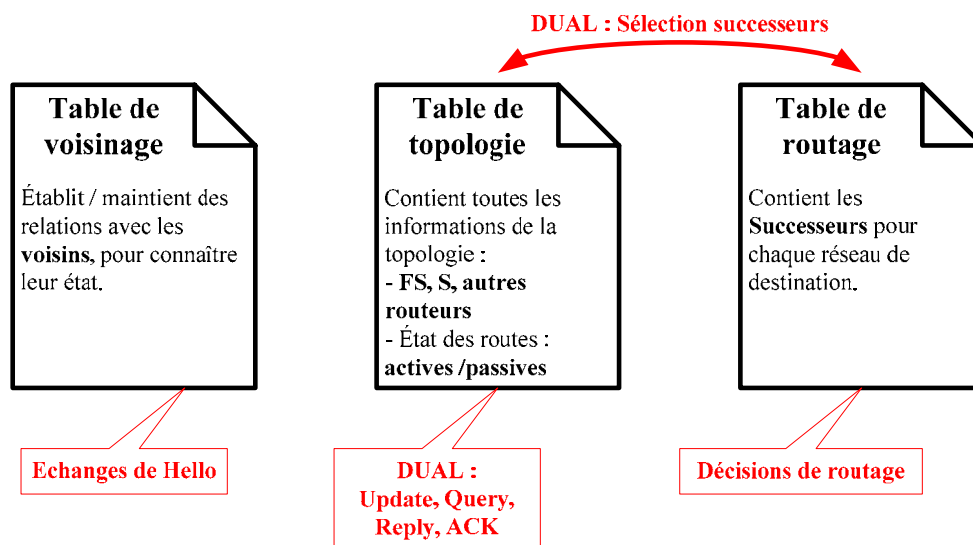
Pour chaque protocole routé utilisé, EIGRP maintient 3 tables distinctes :

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

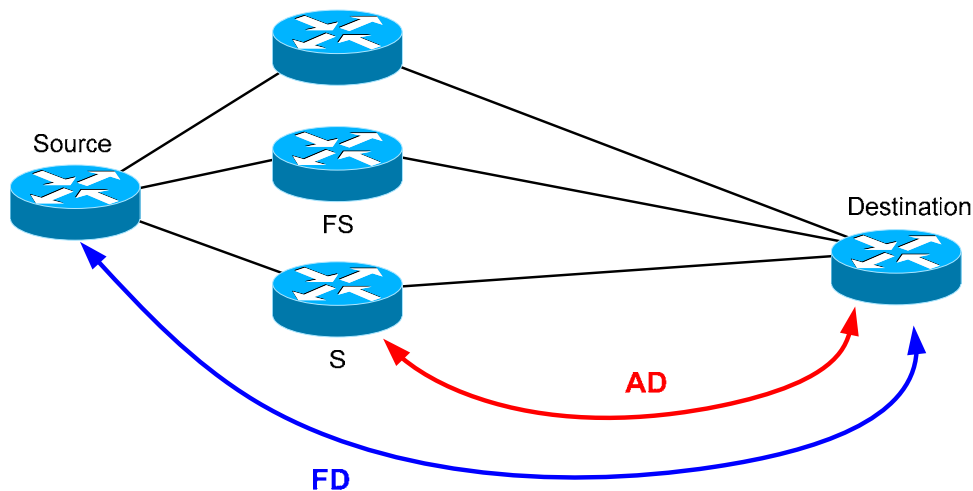
25.2. Termes et définition

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. Ca valeur est 3 fois celle de l'intervalle de transmission des messages Hello. Passé ce délai, le voisin sera considéré comme mort.
- **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)**
 - o 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)**
 - o Autre nom pour l'AD.
- **Feasible Distance (FD)**
 - o Plus petite métrique pour une destination donnée. C'est la métrique pour la route actuellement dans la table de routage.
- **Feasible Condition (FC)**
 - o Condition vérifiée quand un voisin informe une AD plus petite que la FD du routeur local pour une même destination.
- **Feasible Successor (FS)**
 - o Voisin vérifiant la FC. Il est potentiellement éligible en tant que successeur.
- **Successor**
 - o Voisin utilisé comme prochain saut pour une destination donnée. C'est le FS ayant la plus petite métrique.
- **Stuck In Active (SIA) (aussi appelé Query Scoping)**
 - o Etat d'un routeur lorsqu'une route reste active après dépassement d'un certain temps.



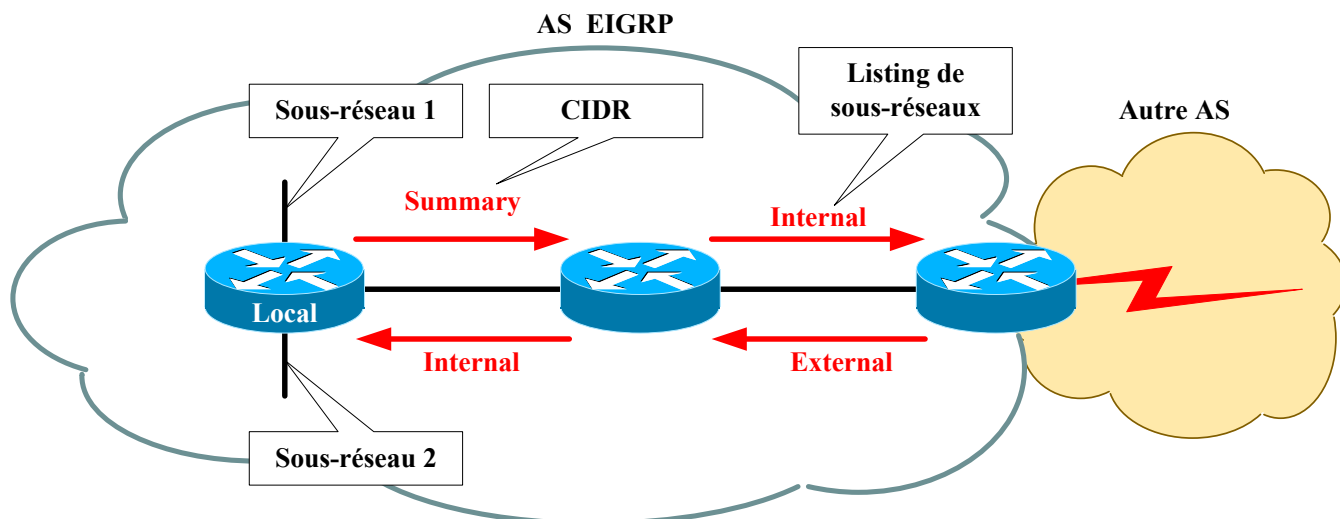
Représentation schématique de quelques définitions

25.3. 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)



Ces routes sont représentées ainsi dans la table de routage :

- **D** : Routes internes et agrégées
- **D EX** : Routes externes

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 :

$$\begin{aligned} \text{Métrique} &= \text{Bandwidth} + \text{Delay} \\ \text{Métrique} &= (10^7 \div \text{BP} + \Sigma_{\text{délais}}) \times 256 \end{aligned}$$

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.

25.4. 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**
 - o Emis périodiquement
 - o Non orienté connexion
 - o Toutes les 5 secondes sur les liaisons LAN
 - o Toutes les 60 secondes sur les liaisons WAN
- **Update**
 - o 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.
 - o Orienté connexion avec RTP
 - o 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**
 - o 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.
 - o Envoyé lorsqu'une ou plusieurs destinations passent à l'état Active
 - o Orienté connexion avec RTP
 - o Ce type de paquet est toujours envoyé en multicast.
 - o 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**
 - o Identique à un paquet Update sauf que celui-ci est envoyé uniquement en réponse à un paquet Query.
 - o Orienté connexion avec RTP
 - o Ce paquet est un unicast vers le voisin ayant émis le paquet Query.
- **ACK**
 - o Accusé de réception pour les paquets envoyés orientés connexion
 - o Envoyé sous la forme d'unicast
 - o C'est un paquet Hello sans données qui contient un numéro d'accusé de réception différent de 0.
 - o 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.
 - o 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.

25.4.1. 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 identiques** 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.

25.4.2. 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 porteur.

25.5. DUAL

Cet algorithme a pour but 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.

25.6. 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.

- **passive-interface {type} {numéro}**
 - o Mode de configuration du protocole de routage
 - o 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}**
 - o Mode de configuration du protocole de routage
 - o Modifie des coefficients entrants en jeu dans le calcul des métriques d'EIGRP.
 - o La valeur de **TOS** doit toujours être de 0.

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

- **show ip route [eigrp [n° AS]]**
 - o Visualise uniquement les routes EIGRP de la table de routage.
- **show ip eigrp neighbors [{type} {numéro} [n° AS]] [detail]**
 - o Fournit toutes les informations sur les voisins, l'état de la relation de voisinage ainsi que les interfaces et adresses par lesquelles ils communiquent.
- **show ip eigrp topology [all | n° AS | [IP] masque]**
 - o 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]**
 - o Donne les informations regroupées sur le trafic total envoyé depuis et vers le processus EIGRP.
- **show ip eigrp interfaces [n° AS] [detail]**
 - o 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.

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

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

26. ACL

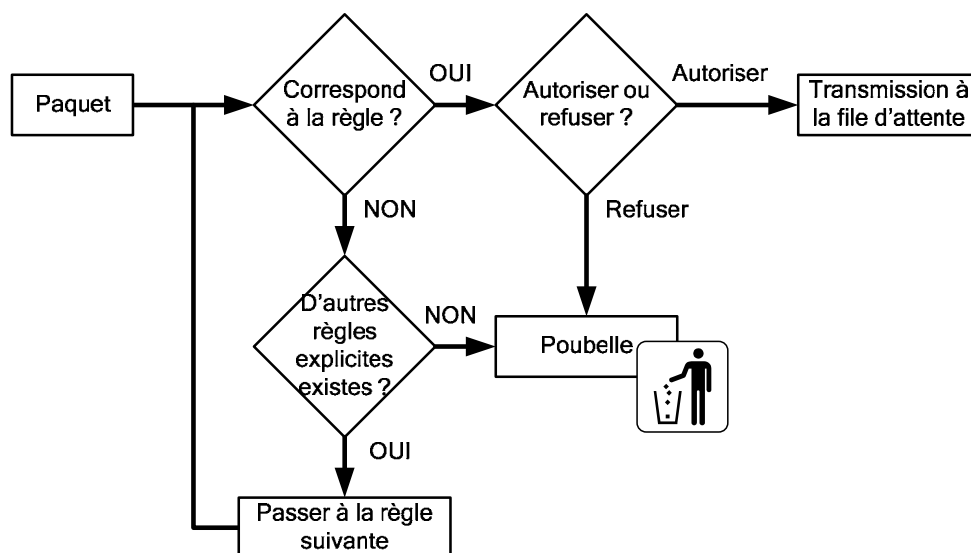
26.1. Théorie

26.1.1. Principe fondamental

Une ACL (Access Control List) est une liste séquentielle de critères utilisée pour du filtrage des paquets. Les ACLs sont capables d'autoriser ou d'interdire des paquets, que ce soit en entrée ou en sortie.

Cette liste est parcourue de la première à la dernière instruction jusqu'à trouver une correspondance. Si le paquet répond aux critères d'une instruction, le reste des instructions est ignoré et le paquet est autorisé ou refusé. Si aucune correspondance n'est trouvée dans les critères explicités par l'administrateur, le paquet est implicitement supprimé.

Il ne peut y avoir qu'une seule ACL par protocole, par interface et par direction (entrée/sortie).



Parcours des instructions d'une ACL

Les ACLs permettent ainsi d'autoriser ou d'interdire des trafics en fonctions de critères tels que les adresses sources et destinations, les protocoles utilisés et les numéros de ports.

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

- ACL standard (numérotée)
- ACL étendue (numérotée)
- ACL nommée (peut être de type standard ou étendue)

Plage de numéros	Type d'ACL associé
1 à 99 et 1300 à 1999	Standard pour IP
100 à 199 et 2000 à 2699	Etendue pour IP
600 à 699	AppleTalk
800 à 899	Standard pour IPX
900 à 999	Etendue pour IPX
1000 à 1099	IPX/SAP

L'avantage principal des ACLs est donc de fournir une base de sécurité réseau en filtrant les trafics traversant un routeur.

Le principal inconvénient est malheureusement un traitement supplémentaire à effectuer pour chaque paquet entrant et/ou sortant du routeur, rallongeant ainsi à la latence réseau et à la surcharge CPU.

La configuration des ACLs se fait en deux parties distinctes, à savoir :

- Création de l'ACL
- Application de l'ACL sur une interface réseau

Quelques précautions sont à prendre en compte lors de la configuration ou de l'utilisation des ACLs :

- Les instructions sont toujours parcourues de la première à la dernière, jusqu'à correspondance des critères.
- Si aucune instruction ne correspond au paquet, la dernière instruction implicite indique alors de supprimer ce paquet.
- Une ACL appliquée sur une interface mais dont les instructions ne sont pas configurées n'a pour seule instruction que la dernière qui bloque tout. Tout trafic serait alors interdit.
- Lors de la création des instructions, il faut toujours procéder du plus précis (exceptions) jusqu'au plus générique.
- Une ACL IP qui interdit un paquet enverra automatiquement un message ICMP Host Unreachable.
- Une ACL pour un trafic sortant n'affecte pas le trafic originaire du routeur local.

26.1.2. Masque générique

Les instructions utilisées dans les ACLs utilisent les masques génériques (Wildcard Mask) conjointement à des préfixes réseaux pour identifier des plages d'adresses.

Un masque générique est une valeur 32 bits noté sous la forme décimale pointée (comme les IP et les masques de sous-réseaux), sachant que :

- **"0" binaire** : Doit correspondre
- **"1" binaire** : Peut varier

On observe donc qu'un masque générique est l'inverse binaire d'un masque de sous-réseaux, ou, du point de vue décimal pointé, est le complément à 255 du masque de sous-réseau correspondant :

Masque de sous-réseau	1111 1111.1111 1111.1110 000.0000 0000
Masque générique	0000 0000.0000 0000.0001 111.1111 1111

$$\begin{array}{rcl} 255 . 255 . 224 . 0 & \text{(Masque de sous-réseau)} \\ + 0 . 0 . 31 . 255 & \text{(Masque générique)} \\ \hline = 255 . 255 . 255 . 255 \end{array}$$

Par conséquent, un masque générique ne peut prendre que ces valeurs (pour chaque octet) :

0	1	3	7	15	31	63	127	255
---	---	---	---	----	----	----	-----	-----

Au niveau syntaxique, deux masques génériques précis (les deux extrêmes, à savoir tout ou rien) peuvent s'écrire normalement, sous la forme préfixe/masque générique, ou sous une forme plus conviviale. Ces deux exceptions d'écriture sont les suivantes :

- **{IP} {0.0.0.0} = host {IP}**
- **{IP} {255.255.255.255} = any**

26.2. ACL standard

Une ACL standard permet d'autoriser ou d'interdire des adresses spécifiques ou bien un ensemble d'adresses ou de protocoles, sachant que, dans les instructions d'une ACL standard, on ne peut indiquer que des adresses sources.

Ce sont les ACLs les plus simples et, par conséquent, les moins gourmandes en ressources CPU. Elles sont par exemple utilisées pour autoriser ou interdire toute une plage d'adresses réseaux ou encore pour le filtrage des informations contenues dans des mises à jour de routage.

Pour configurer une instruction pour une ACL standard pour IP, il faut utiliser la commande suivante :

- **access-list {numéro} {permit | deny} {préfixe} [masque générique] [log]**
- **access-list {numéro} {remark} {commentaire}**
 - o Mode de configuration globale
 - o Si le masque générique n'est pas précisé, le masque générique par défaut 0.0.0.0 est utilisé.
 - o **log** permet de garder en mémoire le nombre de paquets correspondant à l'instruction en cours.
 - o Le mot clé **remark** suivi d'un commentaire permet d'indiquer l'utilité de l'instruction.

L'ordre de parcours des instructions dépend de l'ordre dans lequel on a configuré les instructions. Une nouvelle instruction est donc obligatoirement ajoutée à la fin de la liste, et il est impossible de supprimer une instruction particulière.

Pour toute modification, il est donc conseillé d'utiliser un éditeur de texte, de copier la liste des instructions de l'ACL devant être modifiée, de supprimer cette ACL sur le routeur, d'éditer les instructions pour faire les modifications voulues puis de les insérer dans le routeur.

26.3. ACL étendue

Une ACL étendue permet de faire un filtrage plus précis qu'une ACL standard. En effet, une ACL étendue permet de filtrer en fonction de :

- Protocole utilisé (couche 3 et 4)
- Adresse source
- Adresse de destination
- Numéro de port

La commande permettant de configurer une ACL étendue pour IP est :

- **access-list {numéro} {permit | deny} {protocole} {préfixe source} {masque source} [{opérateur} {opérande}] {préfixe destination} {masque destination} [{opérateur} {opérande}] [icmp-type] [log] [established]**

- **access-list {numéro} {remark} {commentaire}**
 - o Mode de configuration globale
 - o **protocole** peut être soit le nom (IP, TCP, UDP, ICMP, IGRP, etc.) soit le numéro du protocole (de 0 à 255).
 - o Le couple **opérateur/opérande** est pour les numéros de ports TCP ou UDP uniquement, et peut être spécifié pour la source et/ou pour la destination :

Opérateur	Signification
eq	Egal à
neq	Différent de
lt	Inférieur à
gt	Supérieur à
range	Entre (nécessite 2 numéros de port)

- o Le paramètre **icmp-type** ne peut être utilisé que pour le protocole ICMP, et correspond au nom ou au numéro du type de message ICMP devant être vérifié.
- o Le paramètre **established** ne peut être utilisé que pour le protocole TCP et permet de faire correspondre uniquement les sessions TCP déjà établies (drapeaux ACK, FIN, PSH, RST, SYN ou URG).

Pour l'ordre de parcours ou la modification, les règles sont les mêmes qu'avec une ACL standard.

26.4. ACL nommée

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.

Une ACL nommée peut être de type standard ou étendue.

Deux nouveaux modes de configuration sont donc étudiés :

Mode de configuration	Invite de commande associée
ACL nommée standard	(config-std-nacl)#
ACL nommée étendue	(config-ext-nacl)#

Les ACLs nommées permettent :

- D'identifier intuitivement les listes de contrôle d'accès à l'aide d'un code alphanumérique.
- De supprimer une instruction particulière sans avoir à tout supprimer et réécrire.

Les commandes suivantes permettent de configurer une ACL nommée :

- **ip access-list {standard | extended} {nom}**
 - o Mode de configuration globale
 - o Permet de créer une ACL nommée standard ou étendue
 - o Permet de passer dans le mode de configuration de l'ACL nommée
- **{permit | deny} {préfixe} [masque] [log]**
 - o Mode de configuration d'ACL nommé standard
 - o Les paramètres sont identiques que pour une ACL standard numérotée.

- **{permit | deny} {protocole} {préfixe source} {masque source} [{opérateur} {opérande}] {préfixe destination} {masque destination} [{opérateur} {opérande}] [icmp-type] [log] [established]**
 - o Mode de configuration d'ACL nommée étendue
 - o Les paramètres sont identiques que pour une ACL étendue numérotée
- **remark {commentaire}**
 - o Mode de configuration d'ACL nommée (standard ou étendue)
 - o Fournit un commentaire pour indiquer l'utilité de l'ACL

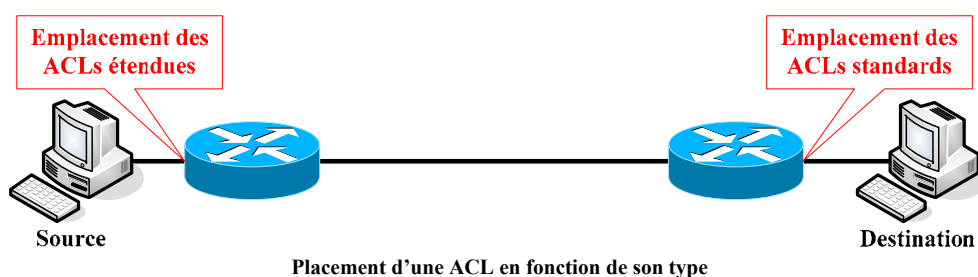
26.5. Mise en place et vérification des ACLs

La création des ACLs étant faite, il faut maintenant les appliquer en utilisant les commandes suivantes :

- **ip access-group {numéro | nom} {in | out}**
 - o Mode de configuration d'interface
 - o Applique une ACL (standard, étendue ou nommée) sur l'interface pour filtrer le trafic entrant ou sortant
- **access-class {numéro | nom} {in | out}**
 - o Mode de configuration de ligne
 - o Applique une ACL sur la ligne pour filtrer les accès à cette dernière
- **no access-list {numéro}**
 - o Mode de configuration globale
 - o Supprime complètement une ACL numérotée

Les commandes suivantes servent à vérifier le placement des ACLs, ainsi que leurs instructions :

- **show access-lists [numéro | nom]** : Affiche la liste des ACLs créées sur le routeur, leurs instructions ainsi que le nombre de correspondance pour chaque instruction
- **show ip interface [{type} {numéro}]** : Permet entre autres de voir quelles sont les ACLs appliquées sur les interfaces et pour quelle direction



Parce que les ACLs standards ne permettent que de filtrer en fonction d'adresses sources, il faut les placer au plus près de la destination, et inversement pour les ACLs étendues qui doivent toujours être placées au plus près de la source.

De plus, les ACLs standards, interdisant intégralement un trafic pour une source donnée, bloquent implicitement le trafic dans le sens opposé (explicitement bloqué de la source vers la destination et implicitement bloqué de la destination à la source).

27. NAT et PAT

27.1. Adressage privé et public

La très forte croissance et popularité d'Internet dans le début des années 90 ont menée très rapidement à la saturation des adresses pouvant être fournies par le protocole IP version 4. C'est entre autres pourquoi le système d'adressage privé a été élaboré, de manière à ralentir l'inévitable, à savoir l'épuisement de toutes les adresses IPv4.

Les plages d'adresses privées définies par la RFC 1918 sont les suivantes :

Classe d'adresses	Plage d'adresses privées	CIDR correspondant
A	De 10.0.0.0 à 10.255.255.255	10.0.0.0/8
B	De 172.16.0.0 à 172.31.255.255	172.16.0.0/12
C	De 192.168.0.0 à 192.168.255.255	192.168.0.0/16

Ces plages d'adresses privées utilisées conjointement à la translation d'adresses, permettent à plusieurs réseaux d'utiliser les mêmes adresses. La translation d'adresse prend alors tout son intérêt en traduisant, ou remplaçant, les adresses privées en une ou plusieurs adresses publiques afin de transiter sur Internet.

Ceci crée donc plusieurs « cellules » d'adresses privées pouvant être identiques pour différents réseaux, sachant que chaque cellule ne serait accessible depuis Internet que par la ou les adresses publiques attribuées à chaque entreprise.

Les adresses privées étant réservée à un usage interne, ces adresses ne peuvent pas être utilisées directement sur Internet. C'est pourquoi les routeurs de bordure des FAI sont configurés pour empêcher le routage de ces adresses.

27.2. Translation d'adresses

La translation d'adresse est un processus générique permettant la substitution d'une adresse par une autre, et permet ainsi de masquer les adresses privées des réseaux locaux derrière une adresse publique.

Ce processus existe sous deux variantes :

- **NAT** (Network Address Translation)
 - o Statique
 - o Dynamique
- **PAT** (Port Address Translation)

27.2.1. Principe du NAT

Le NAT a été fait pour économiser des adresses IP en permettant la translation d'adresses IP privées (RFC1918), internes à une entité (une entreprise, une école etc.) en une ou plusieurs adresses IP publiques routable sur Internet.

Remarque : l'adresse IP utilisée pour la translation n'est pas forcément une adresse IP public et peut être à nouveau une adresse IP privée qui, à son tour, pourra être traduite.

Laboratoire Supinfo des Technologies Cisco

Site : www.labo-cisco.com - Mail : labo-cisco@supinfo.com

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

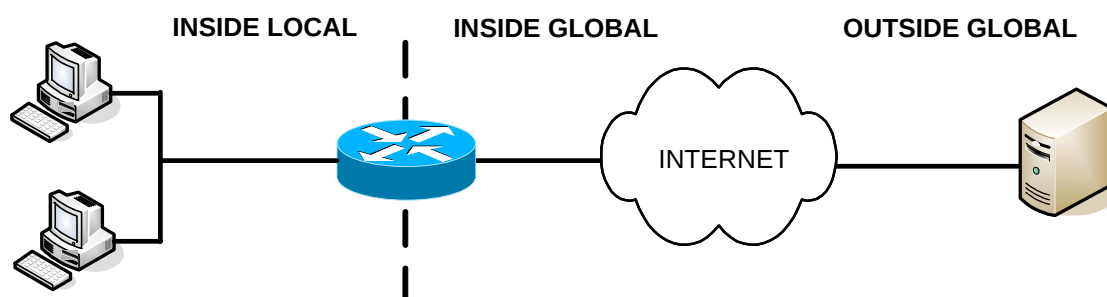
Cette translation d'adresse est effectuée principalement sur les routeurs de bordure d'une entreprise connectée à Internet. Le réseau utilisant les adresses IP privées est ainsi appelé le réseau interne (**inside**), tandis que la partie du réseau utilisant des adresses IP publiques (Internet) est appelé le réseau externe (**outside**).

Quand un utilisateur du réseau interne (inside) souhaite communiquer avec un hôte du réseau externe (outside), le routeur reçoit le paquet avec l'adresse IP privée et réécrit le paquet en changeant l'adresse IP source avec l'adresse IP public du routeur (c'est l'opération de translation).

Le routeur consulte ensuite sa table de routage pour acheminer le paquet jusqu'à la bonne destination. Le destinataire recevra le paquet avec comme source l'adresse IP public du routeur et non l'adresse IP privée de l'hôte qui envoie le paquet dans le réseau interne.

Au-delà des appellations « inside » et « outside », Cisco définit 4 types d'adresses pour le NAT :

- **Inside local address**
 - o Adresse IP attribuée à un hôte dans le LAN.
- **Inside global address**
 - o Adresse(s) IP attribuée(s) par le FAI reconnue(s) par l'Internet pour représenter le LAN.
- **Outside local address**
 - o Adresse IP d'un hôte du réseau externe telle qu'elle est connue par les utilisateurs du réseau interne. La plupart du temps, celle-ci est identique à l'« outside global address ».
- **Outside global address**
 - o Adresse IP attribuée à un hôte dans le réseau externe.



Le NAT peut être utilisé dans plusieurs cas, cependant il peut être configuré de deux manières différentes statiquement ou dynamiquement.

- **Le NAT statique** traduit une adresse IP privée avec toujours la même adresse IP publique. S'il y a 4 utilisateurs nécessitant une traduction d'adresse, il faudra donc utiliser 4 adresses IP publiques.
- **Le NAT dynamique** traduit une adresse privée avec une adresse IP publique appartenant à un pool d'adresses. L'adresse IP publique utilisée pour la traduction n'est donc pas toujours la même. S'il n'y a pas assez d'adresses IP publiques disponibles les utilisateurs devront attendre qu'une adresse se libère pour pouvoir être traduite.

L'avantage du NAT, en plus de la grande économie d'adresses IP, est de ne pas avoir à refaire tout l'adressage IP lorsque l'on change de fournisseur d'accès Internet.

Cette technologie apporte également de la sécurité au sein du réseau interne puisque les machines qui s'y trouvent ne sont pas accessibles depuis l'extérieur.

27.2.2. Principe du PAT

Le PAT (Port Address Translation) ou Overloading permet d'attribuer une seule adresse IP publique pour la translation de plusieurs adresses IP privées. Chaque utilisateur est différencié grâce à un numéro de port unique qui lui est attribué lorsqu'il souhaite communiquer.

Etant donné qu'il existe 65536 ports différents, un routeur pourrait traduire jusqu'à 65536 adresses IP privées différentes. Cependant en réalité, un équipement ne peut gérer en moyenne que la translation d'environ 4000 ports par adresse IP publique.

27.3. Configuration

27.3.1. Commandes

- **ip nat inside**
 - o Mode de configuration d'interface
 - o Spécifie l'interface inside.
 - o Complémentaire des autres commandes NAT
- **ip nat outside**
 - o Mode de configuration d'interface
 - o Spécifie l'interface outside
 - o Complémentaire des autres commandes NAT
- **ip nat inside source static {local-ip} {global -ip}**
 - o Mode de configuration globale
 - o Etablit une translation statique entre une 'Inside local address' et une 'Inside global address'
- **access-list {numéro} permit {prefix} {wildcard_mask}**
 - o Mode de configuration globale
 - o Spécifie le ou les réseaux autorisés à être traduits
- **ip nat inside source list {numéro} pool {nom_du_pool}**
 - o Mode de configuration globale
 - o Définit le pool qui va être traduit
- **ip nat pool {nom_du_pool} {première-ip} {dernière-ip} netmask {masque_de_sous-reseau}**
 - o Mode de configuration globale
 - o Spécifie le pool d'adresses IP : toutes les adresses IP entre première-ip et dernière-ip
- **ip nat inside source list {numéro} interface type {numéro} overload**
 - o Mode de configuration globale
 - o Configuration du PAT sur l'interface outside
- **clear ip nat translation**
 - o Mode privilégié
 - o Configuration du PAT sur l'interface outside

27.3.2. Procédure de configuration

- Spécifier les interfaces outside et inside (ip nat outside / inside)
 - o NAT statique :
 - Spécifier chaque adresse une par une (ip nat inside source static ip1 ip2)
 - o NAT dynamique :
 - Spécifier le bloc privé
 - Spécifier le pool public
 - Activer le NAT avec le bloc privé et le pool public en argument.
 - o PAT :
 - Spécifier le bloc privé
 - Activer le NAT sur l'interface outside avec le bloc privé en argument.

27.3.3. Vérification

- show ip nat translations
 - o Mode privilégié
 - o Affiche des informations sur chaque translation en cours en particulier le temps depuis lequel elle est active.
- show ip nat statistics
 - o Mode privilégié
 - o Configuration du PAT sur l'interface outside
- show running-config
 - o Mode privilégié
 - o Affiche la configuration du routeur.
- debug ip nat
 - o Mode privilégié
 - o Affiche en temps réel toute les paquets traduits.

28. Protocole DHCP

28.1. Introduction

DHCP (Dynamic Host Configuration Protocol) est un protocole fonctionnant en mode Client – Serveur. Il fournit aux clients une configuration de couche 3 : principalement une adresse (IP), mais aussi des adresses de passerelle ou de serveur DNS, NETBIOS, noms de domaines, ...

Ce protocole permet une gestion dynamique de l'adressage de niveau 3. Il allège ainsi grandement les tâches de l'administrateur réseau.

Les **clients DHCP** sont fournis aux utilisateurs sur la plupart des systèmes d'exploitation. Grâce à l'envoi d'une requête au serveur, ceux-ci peuvent se voir attribuer une adresse de couche 3. Seuls les équipements utilisateurs doivent bénéficier de ce service, les serveurs et équipements réseaux devant être adressés de façon statique.

Le DHCP fonctionne sur un principe de location ou bail. Le serveur attribue une adresse à un client pour une durée prédéterminée (en jours, minutes, secondes). Le client doit donc effectuer à nouveau une demande pour voir son bail reconduit.

Il existe trois types d'allocation d'adresse :

- **Automatique** : une adresse IP permanente est attribuée automatiquement au client. Un mappage statique (mac – IP) permet de retrouver la même adresse lors d'une déconnexion / reconnexion.
- **Manuelle** : l'attribution est faite manuellement par l'administrateur réseau (mappage statique). Le protocole DHCP se charge d'envoyer ces informations au client lors d'une demande.
- **Dynamique** : l'attribution se fait à la volée. Une IP libre est attribuée à un client en faisant la demande, pour une durée déterminée.

Les **serveurs DHCP** sont généralement gérés par des serveurs d'entreprise (service généralement assuré par l'OS), mais ils peuvent également être configurés sur les routeurs.

28.1.1. Comparatif entre BOOTP et DHCP

BOOTP (Bootstrap Protocol) est l'ancêtre du protocole DHCP. Son but était d'attribuer une configuration de couche 3 aux stations de travail fonctionnant sans disque dur. DHCP reprend plusieurs de ses caractéristiques :

- Fonctionne en mode client - serveur
- Utilise les ports UDP 67 (serveur) et 68 (client), appelés ports BOOTP
- Attribue une adresse IP
- Attribue un masque de sous-réseau
- Attribue une adresse de passerelle
- Attribue une adresse de serveur DNS

Le protocole BOOTP alloue les adresses de façon statique : le serveur BOOTP doit posséder au préalable une table de correspondance mac – IP pour attribuer une IP. BOOTP n'a pas de notion de bail et fait donc une liaison permanente entre un hôte et l'adresse IP qu'il lui donnera.

Enfin, le protocole DHCP peut fournir jusqu'à 30 options de configuration, contre 4 seulement pour BOOTP (IP, masque, adresse de passerelle, adresse du DNS).

28.1.2. Opération DHCP

La configuration d'un client avec le protocole DHCP se fait en 4 étapes :

1) **DHCP DISCOVER :**

- Lorsqu'une configuration DHCP cliente est présente sur un poste utilisateur, celui-ci envoie une requête en broadcast aux serveurs DHCP, appelée DHCP DISCOVER.

2) **DHCP OFFER :**

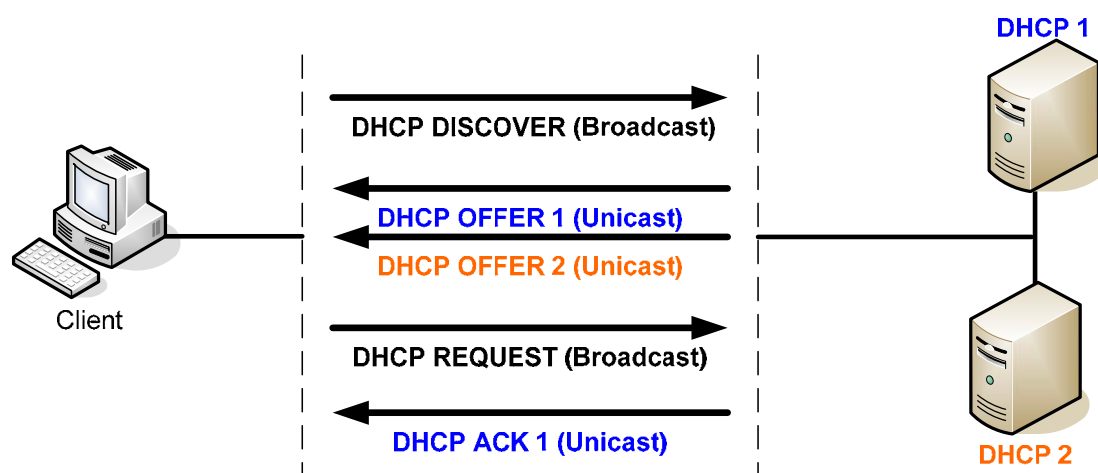
- Les serveurs DHCP recevant le broadcast et pouvant répondre à la demande, envoient une requête en unicast au client. Ce DHCP OFFER contient toutes les informations nécessaires au client (IP, adresse de passerelle, durée du bail, serveur DNS, WINS, etc.).

3) **DHCP REQUEST :**

- Le client émet ensuite une requête en broadcast afin de confirmer l'offre qu'il a sélectionnée (celle qui lui est arrivée en premier).
- S'il y avait plusieurs serveurs DHCP, tous sont alors au courant et peuvent libérer leur offre en conséquence.
- S'il s'agit d'un renouvellement de bail, le client propose au serveur l'IP qu'il veut se voir réattribuer.

4) **DHCP ACK :**

- Cette confirmation est envoyée en unicast par le serveur DHCP au client. Une fois le DHCP ACK reçu, le client peut alors utiliser l'adresse IP ainsi que le reste de la configuration attribuée.



Il existe trois autres requêtes DHCP :

- **DHCP DECLINE :** Si le client détecte l'IP qu'on lui a proposée sur le même segment réseau, il envoie cette requête au serveur. Le processus redémarre alors.
- **DHCP NACK :** Lorsqu'un serveur détecte que l'IP pour laquelle il doit renvoyer un ACK est déjà présente sur le réseau, il envoie un DHCP NACK. Le processus doit alors redémarrer pour le client concerné.

- **DHCP RELEASE** : Lorsqu'un client veut annuler le bail (arrêt du système, commande ipconfig /release sous Windows), cette requête est envoyée au serveur afin qu'il libère la réservation d'adresse.

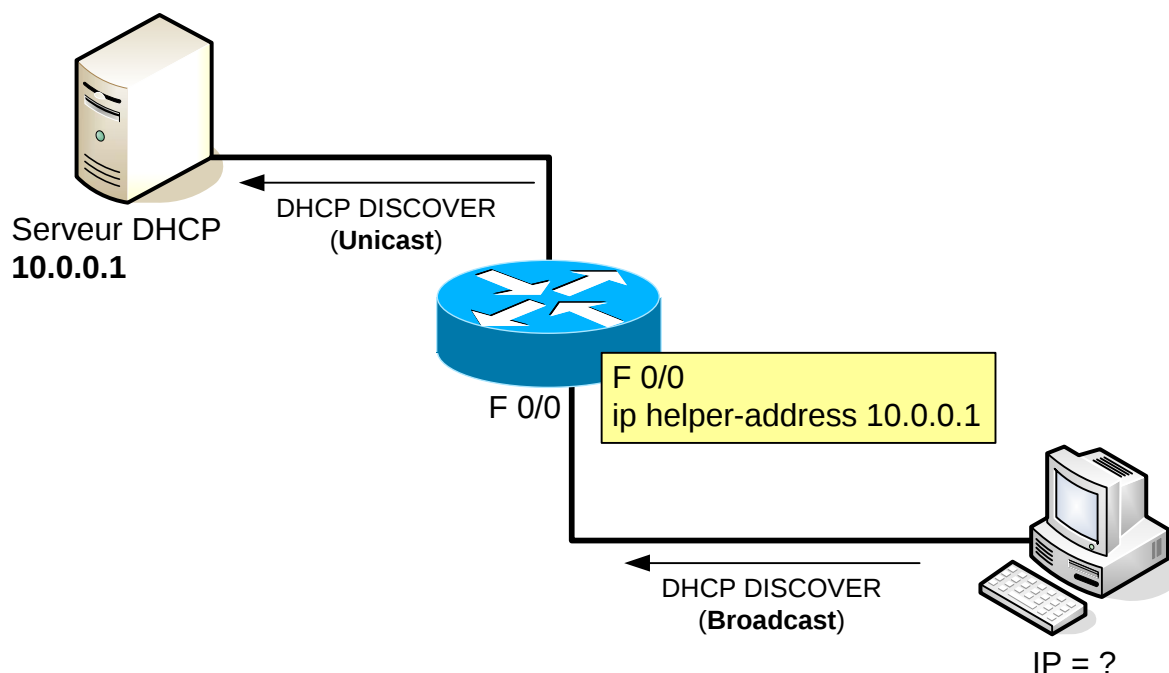
28.1.3. Relais DHCP

Les serveurs DHCP font partie des serveurs d'entreprise. Il est très courant que ces serveurs soient placés sur un sous-réseau différent de celui des utilisateurs.

Un problème se pose donc : les requêtes clientes étant envoyées au serveur DHCP en broadcast, un routeur segmentant le réseau arrêtera également ces broadcast. Il en va de même pour les services DNS, TFTP, TACACS (service d'authentification), etc.

Il est possible d'éviter ce problème en appliquant la commande ip helper-address sur l'interface d'un routeur. Celle-ci permet de relayer les broadcast UDP vers une adresse unicast définie. Ce relais se fait au niveau des services UDP suivants :

- Protocole Time
- TACACS
- Le protocole DNS
- Le service BOOTP/DHCP
- TFTP
- Le service NetBIOS



28.2. Configuration

Comme pour le NAT, la configuration DHCP nécessite la définition de groupe(s) de plages d'adresses attribuables.

28.2.1. Commandes

- **ip dhcp pool {nom_groupe}**
 - o Mode de configuration globale
 - o Passe en mode de configuration DHCP
 - o Spécifie et nomme un groupe d'adresses
- **ip dhcp excluded-address {prefix} [prefix2]**
 - o Mode de configuration globale
 - o Spécifie l'adresse ou la plage d'adresses à exclure du DHCP
- **[no] service dhcp**
 - o Mode de configuration globale
 - o Active/désactive le service DHCP
 - o Actif par défaut
- **network {prefix} {masque}**
 - o Mode de configuration DHCP
 - o Spécifie la plage d'adresses attribuables
- **default-router {prefix}**
 - o Mode de configuration DHCP
 - o Spécifie la passerelle par défaut
- **dns-server {prefix} [prefix2, prefix3, ...]**
 - o Mode de configuration DHCP
 - o Spécifie le(s) serveur(s) DNS
- **netbios-name-server {prefix}**
 - o Mode de configuration DHCP
 - o Spécifie l'adresse du serveur NETBIOS WINS
- **domain-name {nom}**
 - o Mode de configuration DHCP
 - o Spécifie le nom du domaine
- **lease {infinite | jours [heures] [minutes]}**
 - o Mode de configuration DHCP
 - o Spécifie la durée du bail
 - o Valeur par défaut : 1 jour
- **ip helper-address {prefix}**
 - o Mode de configuration d'interface
 - o Relay les broadcast UDP (reçus sur l'interface) vers l'adresse unicast spécifiée.

28.2.2. Procédure de configuration

Voici la procédure permettant de configurer le service DHCP sur un routeur Cisco :

- Définir le nom du groupe d'adresses (commande ip dhcp pool)
- Définir les plages d'adresses attribuables (commande network)
- Spécifier la passerelle par défaut (commande default-router)
- Exclure les adresses IP statiques (commande ip dhcp excluded-address)

Commandes optionnelles :

- Spécifier l'adresse du serveur DNS (commande dns-server)
- Spécifier la durée du bail (commande lease)
- Spécifier l'adresse du serveur NETBIOS (commande netbios-name-server)
- Spécifier le nom de domaine (commande domain-name)
- Relayer les broadcast vers le serveur concerné (commande ip helper-address)

28.2.3. Vérification

Deux commandes show permettent de vérifier le bon fonctionnement du protocole DHCP :

- show ip dhcp binding
 - o Mode privilégié
 - o Affiche les liaisons créées par DHCP (mac – IP)
 - o Affiche la date de fin du bail
 - o Affiche le type d'allocation d'adresse (Automatique, Manuel, Dynamique)
- show ip dhcp server statistics
 - o Mode privilégié
 - o Affiche les requêtes DHCP envoyées et reçues

29. Réseau WAN

29.1. Qu'est-ce qu'un réseau WAN

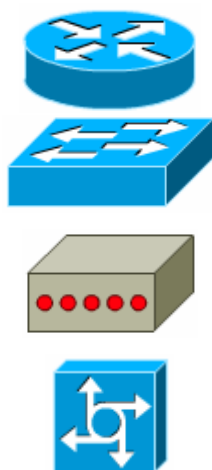
Par définition, un réseau WAN est :

- Un réseau longue distance.
- Un réseau qui interconnecte des réseaux LAN qui sont habituellement séparés par de vastes étendues géographiques.

Les principales caractéristiques des réseaux WAN sont les suivantes :

- Ils fonctionnent au niveau des couches physique et liaison de données du modèle de référence OSI.
- Ils fonctionnent au-delà de la portée géographique des réseaux LAN.
- Ils utilisent les services d'opérateurs Télécoms.
- Ils utilisent diverses connexions série pour communiquer.

29.2. Les différents dispositifs WAN



- **Routeurs** : offrent de nombreuses fonctionnalités dont l'interconnexion de réseaux et des ports d'interface WAN. Ce sont des dispositifs de couche 3 car ils basent leur décision d'acheminement sur les adresses de couche 3.
- **Commutateurs** : se connectent au réseau WAN pour la transmission de la voix, des données et des images.
- **Modems** : Services d'interface de voix ; unités CSU/DSU servant d'interface pour les services T1-E1 ; adaptateurs de terminal/terminaison de réseau 1 (NT1) servant pour les services RNIS. Ce sont des unités de couche 1 car elles n'agissent qu'au niveau de la forme du signal électrique.
- **Serveurs de communication** : Concentrent les communications utilisateur entrantes et sortantes.

29.3. Normes WAN

Les protocoles WAN de couche physique décrivent comment fournir des connexions électriques, mécaniques, opérationnelles et fonctionnelles pour les services WAN. Ces services sont le plus souvent assurés par des fournisseurs d'accès WAN, comme les sociétés Télécoms.

Les protocoles WAN de liaison de données décrivent la façon dont les trames sont transportées entre des systèmes par une liaison unique. Ils incluent les protocoles conçus pour fonctionner avec des services point à point, multipoints et commutés multi-accès, tels que les services Frame Relay.

La couche physique WAN décrit l'interface entre l'équipement ETTD (Equipement Terminal de Traitement de Données) et l'équipement ETCD (Equipement de Terminaison de Circuit de Données). Ces deux parties sont les extrémités d'une liaison WAN. Il y a donc uniquement un ETTD et un ETCD sur une même liaison.

D'un point de vue définition :

- L'ETTD est la partie client d'une liaison WAN. C'est lui qui gère les données.
- L'ETCD est la partie fournisseur de services de la liaison WAN. Il a pour but d'acheminer les données fournies par l'ETTD.

Les encapsulations courantes de liaison de données associées aux lignes série synchrones sont délimitées entre les dispositifs de couche 2 minimum (ponts, commutateurs, routeurs) et sont les suivantes :

- **HDLC** : Norme de l'IEEE ; peut être incompatible avec différents fournisseurs en raison de la façon dont chacun d'eux a choisi de l'implanter. La norme HDLC supporte les configurations point à point et multipoints avec un minimum de surcharge système.
- **Frame Relay** : Utilise des installations numériques de haute qualité ; utilise un verrouillage de trame simplifié, sans mécanisme de correction des erreurs, ce qui signifie qu'il peut envoyer les informations de couche 2 beaucoup plus rapidement que d'autres protocoles WAN.
- **PPP** : Décrit par la RFC 1661 ; deux normes développées par l'IETF ; contient un champ de protocole pour identifier le protocole de couche réseau et permet de configurer un système d'authentification.
- **SDLC** : Protocole de liaison de données WAN conçu par IBM pour les environnements à architecture SNA ; en grande partie remplacé par le protocole HDLC, plus polyvalent.
- **SLIP** : Protocole de liaison de données WAN très répandu pour le transport des paquets IP ; est remplacé dans de nombreuses applications par le protocole PPP, plus polyvalent.
- **Procédure d'accès en mode équilibré (LAPB)** : Protocole de liaison de données utilisé par X.25 ; offre des fonctions étendues de vérification des erreurs.
- **LAPD** : Protocole de liaison de données WAN utilisé sur les canaux D RNIS.
- **Procédure de liaison LAPF** : Concerne les services en mode trame ; protocole de liaison de données WAN, semblable au protocole de liaison LAPD, utilisé avec les technologies Frame Relay.

29.4. Technologies WAN

Il existe trois grands types de services à commutation :

- **Services à commutation de circuits** : Une liaison WAN est généralement une liaison point à point, c'est-à-dire qu'il n'y a que deux extrémités (ETTD et ETCD). On va donc commuter physiquement des aiguillages (comme des centraux téléphoniques) grâce à un code d'identification du destinataire (numéro de téléphone) pour établir une liaison physique directe entre la source et la destination.
- **Services à commutation de paquets** : Pour ces services, les liaisons distinctes existent déjà entre le fournisseur de services WAN et le client. Il ne reste plus qu'à rediriger l'information correctement dans cet espace commuté. Le traitement s'effectue donc au niveau de la couche liaison de données du modèle OSI. Ce traitement est effectué logiciellement.
- **Services à commutation de cellules** : Le principe est le même que pour les services à commutation de paquets, sauf que le traitement est effectué au niveau matériel, grâce à l'emploi de cellules de taille fixe et restreinte.

Services à commutation de circuits :

- **Réseau téléphonique analogique (POTS)** : Il ne s'agit pas d'un service de données informatiques, mais il est présenté ici pour deux raisons : bon nombre de ses technologies font partie de l'infrastructure télécoms en pleine expansion, qui transporte les données et il constitue un modèle de réseau de communication longue distance incroyablement fiable et facile à utiliser. Le média type est le fil de cuivre à paires torsadées.

- **RNIS** : Premier service commuté entièrement numérique. Son usage varie grandement d'un pays à l'autre. Coût modéré. Bande passante maximale de 128 Kbps pour l'interface de base RNIS. Le média type est le fil de cuivre à paires torsadées.

Services à commutation de paquets :

- **X.25** : Technologie plus ancienne, mais encore largement utilisée. Offre des fonctions étendues de vérification des erreurs héritées du passé où les liaisons WAN étaient plus sujettes aux erreurs, ce qui la rend fiable mais limite sa bande passante. Bande passante pouvant atteindre jusqu'à 2 Mbps. Usage assez répandu et coût modéré. Le média type est le fil de cuivre à paires torsadées.
- **Frame Relay** : Plus efficace que X.25, mais avec des services similaires. Bande passante maximale de 1,544 Mbps. Technologie très répandue. Coût : de modéré à faible. Les médias types comprennent le fil de cuivre à paires torsadées et la fibre optique.

Services à commutation de cellules :

- **ATM** : Utilise des petites cellules de longueur fixe (53 octets) pour transporter les données. Bande passante maximale actuelle de 622 Mbps. Les médias types sont le fil de cuivre à paires torsadées et la fibre optique. Usage répandu et croissant. Coût élevé.
- **Service de commutation de données haut débit** : Étroitement lié à ATM et généralement utilisé dans les réseaux métropolitains. Bande passante maximale de 44,736 Mbps. Les médias types sont le fil de cuivre à paires torsadées et la fibre optique. Usage assez peu répandu. Coût très élevé.

Services numériques dédiés :

- **T1, T3, E1 et E3** : Les services T offerts aux États-Unis et les services E en Europe sont des technologies WAN très importantes. Elles utilisent le multiplexage temporel pour "découper" et assigner des tranches de temps pour la transmission des données. Les médias utilisés sont le fil de cuivre à paires torsadées et la fibre optique. Leur usage est largement répandu et leur coût est modéré.
- **xDSL** : Nouvelle technologie WAN en développement pour usage domestique. Offre une bande passante qui diminue en fonction de la distance par rapport à l'équipement de l'opérateur. Des vitesses maximales de 51,84 Mbps sont possibles près d'un central téléphonique, mais des débits largement inférieurs sont plus courants (de quelques centaines de Kbps à plusieurs Mbps). Usage peu répandu, qui augmente rapidement. Coût modéré en baisse. Le caractère x indique l'ensemble de la famille de technologies DSL, dont HDSL, SDSL, ADSL, VDSL et RADSL.
- **SDH** : Une famille de technologies propre à la couche physique, offrant de très hauts débits et conçue pour la fibre optique, elle peut aussi être utilisée avec des câbles de fil de cuivre. Elle offre une série de débits de données disponibles avec désignations spéciales. Elle est mise en œuvre à différents niveaux d'opérateur optique, de 51,84 Mbps (OC-1) à 9,952 Mbps (OC-192). Ces débits exceptionnels peuvent être atteints grâce au multiplexage de longueur d'onde, permettant aux lasers d'être réglés sur des couleurs (longueurs d'onde) légèrement différentes, afin d'envoyer d'énormes quantités de données sur un câble optique. D'un usage répandu sur le backbone Internet, cette technologie reste d'un coût élevé.

Autres services WAN :

- **Modem commuté** (analogique) : Limité au niveau du débit mais très polyvalent. Fonctionne avec le réseau téléphonique actuel. Bande passante maximale d'environ 56 Kbps. Coût faible. Usage encore très répandu. Le média type est la ligne téléphonique à paires torsadées.
- **Modem câble** (analogique partagé) : Envoyant des signaux de données sur le même câble que les signaux de télévision, cette technologie augmente en popularité dans les régions où le câble coaxial de la télévision câblée est très répandu (90 % des foyers aux États-Unis). La bande passante maximale peut atteindre 10 Mbps, mais elle diminue avec le nombre d'utilisateurs qui se relie à un segment donné du réseau (de la même manière qu'un réseau LAN non commuté). D'un coût relativement faible, elle est d'un usage peu répandu, bien qu'en constante augmentation. Le média utilisé est le câble coaxial.
- **Sans fil** : Aucun média n'est nécessaire car les signaux sont des ondes électromagnétiques. Il existe une variété de liaisons WAN sans fil, dont les liaisons terrestres et par satellite.

30. Tests de base et résolution de problèmes

30.1. Commandes de vérification

Les problèmes d'adressage sont les problèmes les plus fréquemment rencontrés sur les réseaux IP. Il est important de vérifier la configuration de vos adresses avant de passer aux autres étapes de configuration.

Trois commandes vous permettent de vérifier la configuration des adresses dans votre inter réseau :

- **telnet {IP ou nom d'hôte} [tcp-port-number]** : Elle vérifie le logiciel de la couche application entre les stations d'origine et de destination. Elle constitue le mécanisme de test le plus complet.
- **ping {IP ou nom d'hôte}** : Elle utilise le protocole ICMP (Internet Control Message Protocol) pour vérifier la connexion matérielle et l'adresse de couche réseau du modèle OSI. Elle constitue un mécanisme de test de base.
- **trace {IP ou nom d'hôte}** : Elle utilise des valeurs TTL pour générer des messages à partir de chaque routeur utilisé tout au long du chemin. Cette commande est un outil très efficace, car cela permet de détecter les pannes entre les stations d'origine et de destination. C'est un mécanisme de test de couche réseau.

La commande **telnet** permet donc, en plus d'offrir un accès à un hôte pour pouvoir l'administrer, de vérifier l'état fonctionnel d'un service. Il nous est possible par conséquent d'explicitier le service, par le biais du port TCP qui lui est rattaché, afin d'en vérifier le bon fonctionnement.

La commande **ping** nous renvoie des informations de la forme suivante :

Caractère	Description
!	Réception réussie d'une réponse d'écho
.	Délai d'attente dépassé pour la réponse du datagramme
U	Erreur due à une destination inaccessible
C	Paquet ayant rencontré une congestion de trafic
I	Vérification ping interrompue (par exemple avec la combinaison Ctrl-Maj-6)
?	Type de paquet inconnu
&	Durée de vie du paquet dépassée

On a à notre disposition une commande ping étendue, nous permettant de spécifier les options d'en-tête Internet prises en charge. Cette version de la commande ping n'est disponible qu'en mode privilégié. On y accède en tapant juste **ping** à l'invite de commande.

Lorsqu'on utilise la commande **trace**, les noms de machine s'affichent si les adresses sont converties de façon dynamique ou si elles sont converties par le biais d'entrées de la table d'hôtes statique. Les délais représentent le temps de retour requis pour chacun des trois analyseurs. Si il y a un problème quelconque, les résultats ne seront pas ces temps, mais seront parmi les suivants :

!H	La sonde d'analyse a été reçue par le routeur, mais elle n'a pas été transmise, probablement en raison d'une liste d'accès
P	Le protocole était inaccessible
N	Le réseau était inaccessible
U	Le port était inaccessible
*	Le délai d'attente a été dépassé

30.2. Erreurs courantes au niveau des trois premières couches du modèle OSI

Les erreurs courantes au niveau de la couche 1 sont les suivantes :

- Des câbles rompus.
- Des câbles déconnectés.
- Des câbles raccordés à des ports inappropriés.
- Des connexions instables.
- Des câbles inappropriés pour la tâche à accomplir (les câbles console, les câbles d'interconnexion et les câbles droits doivent être employés à bon escient).
- Des problèmes d'émetteur-récepteur.
- Des problèmes de câblage ETCD.
- Des problèmes de câblage ETTD.
- Des unités hors tension.

Les erreurs courantes au niveau de la couche 2 sont les suivantes :

- Des interfaces série configurées de façon incorrecte.
- Des interfaces Ethernet configurées de façon incorrecte.
- Un ensemble d'encapsulation inapproprié (HDLC est utilisé par défaut pour les interfaces série).
- Une fréquence d'horloge inappropriée pour les interfaces WAN.

Les erreurs courantes au niveau de la couche 3 sont les suivantes :

- Un protocole de routage non activé.
- Un protocole de routage activé mais incorrect.
- Des adresses IP incorrectes.
- Des masques de sous-réseau incorrects.
- Des liens DNS-IP incorrects.

30.3. Debugging

IOS met à notre disposition toute une panoplie de commandes nous permettant de vérifier en temps réel certains aspects. Cela nous permet de vérifier le bon fonctionnement du routeur et, le cas échéant, d'avoir des informations sur les problèmes rencontrés.

Il faut utiliser les commandes **debug** avec parcimonie car elles exigent un temps processeur important. Elles sont disponibles depuis le mode privilégié.

Les commandes suivantes sont celles que nous sommes amenés à rencontrer au cours de la résolution de problème sur le module 2 du programme CNA :

- **no debug all** ou **undebug all** : Permet de stopper tous les déboguages en cours.
- **debug ip rip** : Affiche les mises à jour du routage RIP lors de leur envoi et de leur réception.
- **debug ip igrp events** : Idem mais pour IGRP.
- **debug ip http** : Informations concernant les connections au serveur HTTP.
- **debug cdp events** : Tous les événements CDP.
- **debug cdp adjacency** : Affiche les informations sur les voisins CDP.
- **debug interface {type} {numéro}** : Tout ce qui concerne l'état de l'interface choisie.
- **debug all** : Affiche l'intégralité des informations de déboguage disponibles.

30.4. Procédure de récupération des mots de passe d'un routeur

Pour pouvoir accéder à un routeur, sachant que l'on ne dispose pas du ou des mots de passe appropriés, nous avons à notre disposition une procédure de récupération des mots de passe :

- On doit éteindre, attendre quelques secondes puis rallumer le routeur.
- Avant les 60 secondes qui suivent la remise sous tension du routeur, il faut appuyer simultanément sur les touches **Ctrl** et **Pause**. A ce moment, un caractère d'interruption est envoyé au routeur pour interrompre la séquence d'amorçage. On se retrouve alors dans le mode RXBoot, dont l'invite de commande est un chevron.
- Il faut maintenant modifier la valeur du registre de configuration. Pour un routeur Cisco de la gamme 2500, il faut utiliser la commande **o/r 0x2142**. Cette modification demande au routeur d'ignorer le fichier de configuration stocké dans la mémoire NVRAM.
- On tape le caractère **i** pour continuer le processus de démarrage le routeur.
- Ensuite, à la question concernant la configuration initiale (mode setup), il faut répondre par **n**, car on dispose déjà d'une configuration en NVRAM que l'on va pouvoir exploiter.
- On peut maintenant passer en mode privilégié sans mot de passe. A ce moment, il faut copier le fichier de configuration de la mémoire NVRAM vers la RAM, grâce à la commande **copy startup-config running-config**.
- Il faut ensuite modifier les différents mots de passe, principalement ceux du mode privilégié et de la ligne console, pour avoir les accès nécessaires.
- Il faut maintenant remettre le registre de configuration dans son état d'origine, pour que le routeur puisse accéder à fichier de configuration lors des prochains redémarrages. Pour cela, il faut entrer la commande **config-register 0x2102** en mode de configuration globale.
- La dernière opération est de sauvegarder le fichier de configuration modifié avec les nouveaux mots de passe dans la mémoire NVRAM, et ce à l'aide de la commande **copy running-config startup-config**.
- A ce moment, on peut relancer le routeur avec la commande **reload**.

Il est important de redémarrer le routeur à la fin de cette procédure, car les mots de passe des lignes console et auxiliaire ne sont pris en compte qu'après ce redémarrage.

Pour cette procédure, il faut avoir impérativement un accès physique au routeur, par le biais du port console.

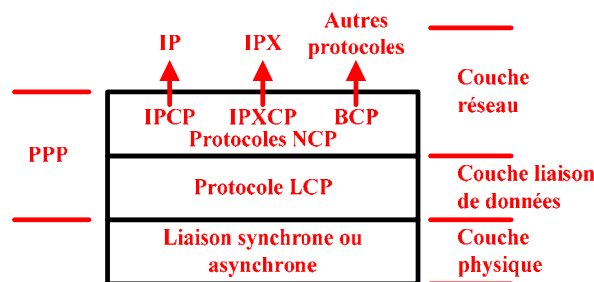
31. Protocole PPP

31.1. Étude du protocole

C'est le protocole de réseau WAN le plus répandu, successeur du protocole SLIP, permettant :

- Connexion entre routeurs ou entre un hôte et un routeur.
- Gestion des circuits synchrones et asynchrones.
- Contrôle de la configuration des liaisons.
- Possibilité d'attribution dynamique des adresses de couche 3.
- Multiplexage des protocoles réseau (Possibilité de faire passer plusieurs paquets de protocoles différents sur la même connexion).
- Configuration des liaisons et vérification de leur qualité.
- Détection des erreurs.
- Négociation d'options (Adresses de couche 3, Compression, etc.).

Le protocole PPP est composé de trois parties distinctes indispensables :



- **Un mode d'encapsulation** : La trame PPP est une trame générique HDLC modifiée.
- **Le protocole LCP (Link Control Protocol)** : Etablissement et contrôle d'une session.
 - o Trame LCP d'établissement de liaison.
 - o Trame LCP de fermeture de liaison.
 - o Trame LCP de maintenance de liaison.
- **Une famille de protocoles NCP (Network Control Protocol)** : Gestion des protocoles de couche 3.
 - o IPCP (Internet Protocol Control Protocol).
 - o IPXCP (Internetwork Packet eXchange Control Protocol).
 - o BCP (Bridge Control Protocol).

Une trame PPP est de la forme :

Drapeau (1 octet)	Adresse (1 octet)	Contrôle (1 octet)	Protocole (2 octets)	Données (Taille variable)	FCS (2 ou 4 octets)
------------------------------------	------------------------------------	-------------------------------------	---------------------------------------	--	--------------------------------------

- **Drapeau** : Indicateur de début ou fin de trame (Valeur = 01111110).
- **Adresse** : Adresse de broadcast standard (Valeur = 11111111), car PPP n'attribue pas d'adresse d'hôte (Couche 2).
- **Contrôle** : Fourniture d'un service non orienté connexion (semblable au LLC) (Valeur = 00000011).
- **Protocole** : Identification du protocole encapsulé (IP, IPX, etc.).
- **Données** : Contient soit la valeur zéro, soit des données (1500 octets maximum).
- **FCS** : Séquence de contrôle de trame pour une vérification des erreurs.

31.2. Établissement d'une session

Les quatre phases d'une session PPP, pour l'établissement des communications sur une liaison point-à-point, sont :

- **Établissement de la liaison.**
- **Détermination de la qualité de la liaison.**
- **Configuration du ou des protocoles de couche réseau.**
- **Fermeture de la liaison.**

Ce sont les trames LCP qui se chargent du bon déroulement de ces quatre phases.

Phase 1 - Etablissement de la liaison :

- Le nœud d'origine envoie des trames LCP pour configurer et établir la liaison.
- Négociation des paramètres de configuration grâce au champ d'option des trames LCP (MTU, compression, authentification, etc.). Ces options peuvent donc être explicite (indiquées dans les trames LCP) ou implicites (Utilisation des valeurs par défaut).
- Fin de cette phase par l'émission et la réception d'une trame LCP d'accusé de réception de la configuration.

Phase 2 - Détermination de la qualité de la liaison :

- Cette phase est facultative.
- Vérification de la qualité suffisante pour activer les protocoles de couche 3.
- Une fois la liaison établie, le processus d'authentification est lancé, si nécessaire.

Phase 3 - Configuration du ou des protocoles de couche réseau :

- Émission de paquets NCP pour configurer les protocoles de couche 3 choisis.
- Configuration individuelle des protocoles de couche 3 grâce au protocole NCP approprié.
- Activation et fermeture à tout moment des protocoles de couche 3.
- Les paquets des protocoles de couche 3 sont émis une fois configuré par son NCP correspondant.

Phase 4 - Fermeture de la liaison :

- Fermeture par le biais de trames LCP ou de paquets NCP spécifiques (Si LCP ferme la liaison, il informe les protocoles de couche 3 par l'intermédiaire du NCP correspondant).
- Fermeture à cause d'un événement extérieur (délai d'attente, perte de signaux, etc.).
- Fermeture en cas de demande d'un utilisateur.

On peut vérifier l'état des protocoles LCP et NCP grâce à la commande **show interfaces**.

31.3. Authentification/configuration

Le protocole PPP peut prendre en charge plusieurs modes d'authentification :

- Aucune authentification.
- Utilisation du protocole PAP.
- Utilisation du protocole CHAP.

Les caractéristiques du protocole PAP sont :

- **Échange en deux étapes** (après la demande d'authentification) :
 - Envoie des informations d'authentification.
 - Acceptation ou refus.
- **Méthode simple d'authentification** : Emission de la combinaison utilisateur/password de façon répétée jusqu'à :
 - Confirmation de l'authentification.
 - Interruption de la connexion.
- **PAP n'est pas très efficace** :
 - Mots de passe envoyés en clair.
 - Aucune protection (Lecture répétée des informations, attaques répétées par essais et erreurs).
- Le nœud s'authentifiant contrôle la fréquence et la durée des tentatives d'authentification.

Pour le protocole PAP, on a le choix entre une authentification :

- **Unidirectionnelle** : Seul le client est authentifié sur le serveur de compte.
- **Bidirectionnelle** : Chaque hôte authentifie l'autre.

Celles du protocole CHAP sont :

- **Échange en trois étapes** (après la demande d'authentification) :
 - Confirmation.
 - Réponse.
 - Acceptation ou refus.
- **Méthode d'authentification plus évoluée** :
 - Vérification régulière de l'identité du nœud distant (A l'établissement puis à tout moment).
 - Authentification dans les deux sens.
 - Impossibilité de tenter une authentification sans avoir reçu une demande de confirmation.
 - Authentification cryptée via l'algorithme MD5 lors du transit sur la liaison.
- **Efficacité contre le piratage** :
 - Utilisation d'une valeur de confirmation variable, unique et imprévisible.
 - Répétition des demandes de confirmation visant à limiter la durée d'exposition aux attaques.
- Chaque côté contrôle la fréquence et la durée des tentatives d'authentification.

Les commandes permettant de configurer tous les différents aspects du protocole PPP sont les suivantes :

- **username {nom} password {mot_de_passe} :**
 - o Mode de configuration globale.
 - o Paramètre **nom** : Nom d'hôte qu'on souhaite accepter.
 - o Paramètre **mot_de_passe** : Mot de passe à utiliser pour l'authentification. Celui-ci doit correspondre au mot de passe du mode privilégié crypté du routeur distant si on utilise CHAP. Ce mot de passe doit être le même sur les deux routeurs.
 - o Définir un compte d'utilisateur localement, afin de permettre l'authentification d'un hôte distant.
- **encapsulation PPP :**
 - o Mode de configuration d'interface.
 - o Spécifier le mode d'encapsulation pour l'interface courante.
- **ppp authentication {chap | chap pap | pap chap | pap} [callin] :**
 - o Mode de configuration d'interface.
 - o Définir la méthode d'authentification voulue. On a la possibilité de définir deux méthodes différentes. Dans ce cas, la première est utilisée, et en cas de refus ou de suggestion de la deuxième, la deuxième méthode sera utilisée.
 - o Le paramètre **callin** est utilisé pour différencier l'authentification unidirectionnelle de la bidirectionnelle.
- **ppp pap sent-username {nom} password {mot_de_passe} :**
 - o Mode de configuration d'interface.
 - o Indique les informations qui seront envoyées lors d'une demande d'authentification PAP. Les informations doivent correspondre au compte utilisateur défini sur le routeur distant.
- **ppp chap hostname {nom} :**
 - o Mode de configuration d'interface.
 - o Permettre l'authentification sur plusieurs routeurs en donnant toujours le même nom d'hôte.
- **ppp chap password {mot_de_passe} :**
 - o Mode de configuration d'interface.
 - o Idem que pour le hostname, mais pour le mot de passe. Ceci permet de limiter le nombre d'entrées utilisateur/password.
- **ppp quality {pourcentage} :**
 - o Mode de configuration d'interface.
 - o Permet de configurer le LQM (Link Quality Monitor) sur la liaison PPP courante. Si la qualité de la liaison tombe en dessous du pourcentage spécifié, le routeur coupera la liaison.

Pour tout problème concernant l'authentification et la négociation de liaison par rapport au protocole PPP, nous avons à notre disposition les commandes suivantes :

- **debug ppp authentication**
- **debug ppp negotiation**

32. Technologie RNIS

32.1. Technologie

Il existe deux types de services RNIS :

- **BRI** : Accès de base.
 - o Aussi appelé canal 2B+D.
 - o 2 canaux B à 64 Kbits/s (8 bits).
 - o 1 canal D à 16 Kbits/s (2 bits).
 - o Débit binaire de 192 Kbits/s (8000 trames de 24 bits).
 - o Débit réel de 144 Kbits/s (2 canaux B + 1 canal D).
- **PRI** : Accès primaire (fonctionnant sur des lignes dédiées).
 - o **T1** (Débit de 1.544 Mbits/s) :
 - 23 canaux B à 64 Kbits/s (8 bits).
 - 1 canal D à 64 Kbits/s (8 bits).
 - 1 bit de verrouillage de trame.
 - 8000 trames par seconde.
 - o **E1** (Débit de 2.048 Mbits/s) :
 - 30 canaux B à 64 Kbits/s (8 bits).
 - 1 canal D à 64 Kbits/s (8 bits).
 - 1 canal à 8 bits pour le verrouillage de trame.

La vitesse de transmission est toujours de 8000 trames par seconde et par canal.

Ces deux services utilisent plusieurs canaux, qui sont répartis en deux types :

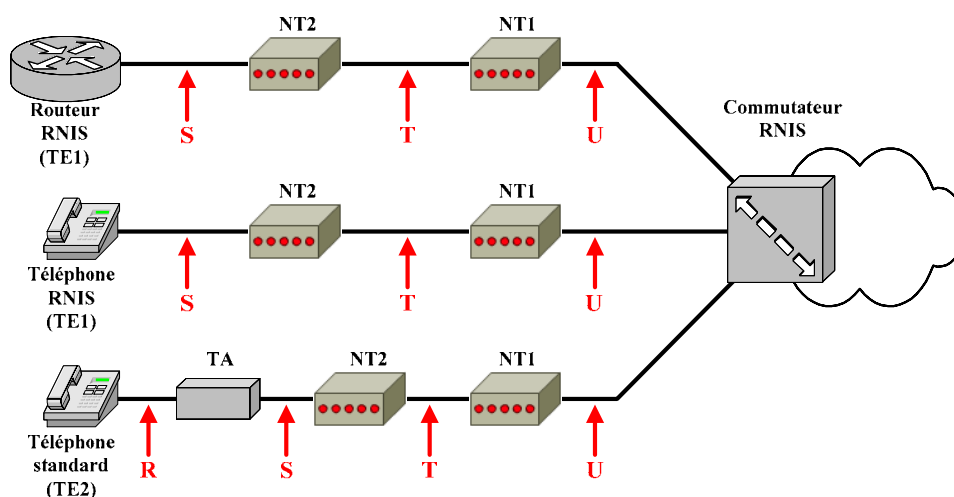
- **Canal B (Bearer)** :
 - o Acheminement du trafic de voix et de données.
 - o Le RNIS offre une grande souplesse d'utilisation, car il est possible d'utiliser chaque canal B séparément, pour transmettre à la fois la voix (Téléphone) et les données (Informatique).
 - o Le protocole PPP multilaision s'occupe du regroupement de la bande passante lorsque plusieurs canaux B sont utilisés pour le trafic de données.
 - o Utilisation éventuelle d'un SPID par canal B. Cet identificateur permet de déterminer la configuration de ligne, et ressemble à un numéro de téléphone. Le commutateur peut ainsi relier les services demandés à la connexion.
- **Canal D (Delta)** :
 - o Canal de signalisation des instructions de traitement des données des canaux B.
 - o Le protocole de signalisation de ce canal s'exécute au niveau des couches 1 à 3 du modèle OSI.

Le protocole LAPD (Couche 2) est utilisé sur le canal D et permet une circulation et une réception adéquate des flux d'information de contrôle et de signalisation. Ce protocole est similaire à HDLC et à LAPB (X.25).

Il est possible de connecter plusieurs unités utilisateur sur un même circuit RNIS. Dans ce cas, des collisions peuvent apparaître. Le canal D prend en charge des fonctions permettant de déterminer des conflits sur la liaison. Il a été mis en place un principe simple afin de permettre à chaque terminal de transmettre :

- Un terminal ne peut transmettre sur le canal D que lorsqu'il détecte un nombre précis de 1 (indiquant l'absence de signal), ce qui correspond à un niveau de priorité prédéterminé.
- Si le terminal détecte un bit E (Voir normes RNIS) qui est différent de ses bits du canal D, il doit cesser immédiatement la transmission.
- Dès que le message du canal D a été transmis, le niveau de priorité du terminal est réduit.
- Un terminal ne peut passer à un niveau de priorité supérieur que si tous les autres terminaux sur la même ligne n'ont pas eu la possibilité d'émettre un message de canal D.
- La connexion téléphonique est prioritaire aux autres services (Données, etc.).
- L'information de signalisation est prioritaire aux autres types d'informations.

32.2. Termes & équipements



Les différents équipements que l'on peut trouver sur un réseau RNIS sont :

- **Commutateur RNIS :** Dispositif de couche 2 permettant la commutation entre les différentes liaisons RNIS.
- **NT1 (Terminaison réseau 1) :**
 - o Unité reliant le câblage à quatre fils de l'utilisateur à la boucle locale à deux fils classique.
- **NT2 (Terminaison réseau 2) :**
 - o Unité dirigeant le trafic des différentes unités terminales (TE1 et TE2) vers le NT1.
 - o Assure les fonctions de commutation et de concentration (Permet de connecter plusieurs TE sur un NT1).
 - o Généralement présent dans les autocommutateurs numériques (PABX).
- **TA (Adaptateur de terminal) :**
 - o Unité convertissant des signaux standard (provenant d'un TE2) au format RNIS.
 - o Raccordée en amont sur une unité NT 1 ou 2.
- **TE1 (Equipement terminal 1) :**
 - o Unité compatible RNIS.
 - o Raccordée sur une unité NT 1 ou 2.
 - o Reliée au réseau au moyen d'une liaison numérique à paires torsadées de quatre fils.
- **TE2 (Equipement terminal 2) :**
 - o Unité non compatible RNIS.
 - o Raccordée sur une unité TA.

Les points de référence RNIS sont regroupés sous quatre désignations :

- **R** : Interface entre une unité TE2 et un TA.
- **S** : Interface entre un NT2 et un TE1 ou TA. C'est la partie qui active les appels entre les différentes parties du CPE.
- **T** : Idem électriquement que S mais correspond à la connexion entre un NT2 et un NT1 ou le réseau RNIS.
- **S/T** : Interface entre un TE1 ou un TA et directement un NT1 (car le NT2 est optionnel).
- **U** : Interface entre un NT1 et le réseau RNIS (uniquement aux USA, car NT1 n'est pas pris en charge par l'opérateur).

32.3. Normes

La technologie RNIS a été mise au point en vue d'uniformiser les services proposés par les opérateurs aux abonnés. Cette uniformisation comprend l'**interface UNI** (Correspond aux informations génériques de base ainsi qu'à des fonctions réseau). En plus de cette interface UNI, une pile complète de protocoles (Couches 1 à 3) a été définie.

Les différents protocoles définis pour le RNIS sont classés dans trois catégories :

- **E** : Normes de réseau téléphonique RNIS.
 - **E.164** : Adressage international RNIS.
- **I** : Concepts, terminologie et méthodes générales.
 - **Série I.100** : Concepts généraux.
 - **Série I.200** : Aspects des services RNIS.
 - **Série I.300** : Aspects réseau.
 - **Série I.400** : Comment est fournie l'interface UNI.
- **Q** : Fonctionnement de la commutation et de la signalisation.
 - **Q.921** : Décrit les processus du protocole LAPD (Canal D).
 - **Q.931** : Précise les fonctions de couche 3 (entre le point d'extrémité et le commutateur RNIS).

La norme Q.931 n'impose pas de recommandation de bout en bout. Cette norme a donc pu être mise en œuvre de diverses façons en fonction du fournisseur et du type de commutateur. Ce point est à préciser lors de la configuration.

Les différentes normes que nous étudierons en fonction des couches du modèle OSI sont :

- **Couche physique** :
 - **I.430** : Spécification de couche physique du BRI.
 - **I.431** : Spécification de couche physique du PRI.
- **Couche liaison de données** :
 - **Q.920 à Q.923** : Spécification fondée sur LAPD.
- **Couche réseau** :
 - **Q.930 (I.450)** et **Q.931 (I.451)** : Définition des connexions entre utilisateurs, à commutation de circuits ou de paquets. La signalisation d'établissement, maintien et fermeture des connexions réseau RNIS est le principal objectif de ces deux normes. Elles s'occupent aussi de fournir une variété de messages (Configuration, connexion, libération, information sur les utilisateurs, annulation, état et déconnexion).

Il existe deux formats de trames pour le RNIS :

- **Trame TE** : Trame sortante (Terminal au réseau).
- **Trame NT** : Trame entrante (Réseau au terminal).

Elles ont une taille de 48 bits, dont 36 de données. Il s'agit en réalité de deux trames successives de 24 bits (deux canaux B à 8 bits + un canal D à 2 bits + 6 bits de verrouillage de trame) :

Trame NT

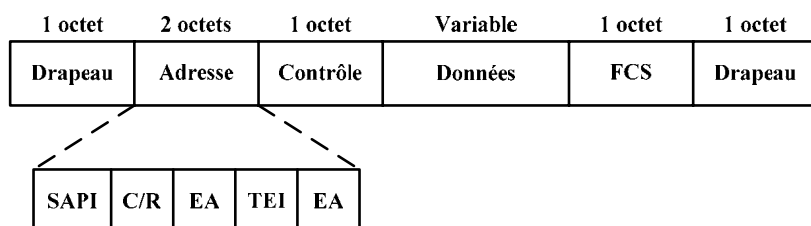
F	L	B1	L	D	L	F	L	B2	L	D	L	B1	L	D	L	B2	...
1	1	8	1	1	1	1	1	8	1	1	1	8	1	1	1	8	3

Trame TE

F	L	B1	E	D	A	F	F	B2	E	D	S	B1	E	D	S	B2	...
1	1	8	1	1	1	1	1	8	1	1	1	8	1	1	1	8	3

- **A** : Bit d'activation (Activation d'unités).
- **B1** : Bits de canal B1.
- **B2** : Bits de canal B2.
- **D** : Bit de canal D.
- **E** : Echo du bit D précédent (Résolution de conflits pouvant survenir lorsque plusieurs terminaux sur un bus passif rivalisent pour un canal).
- **F** : Bit de verrouillage de trame (Synchronisation).
- **L** : Bit d'équilibrage de charge (Ajustement de la valeur moyenne de bit).
- **S** : Bit de réserve (Non affecté).

Ces deux types de trame sont sous la forme d'une trame LAPD générique :

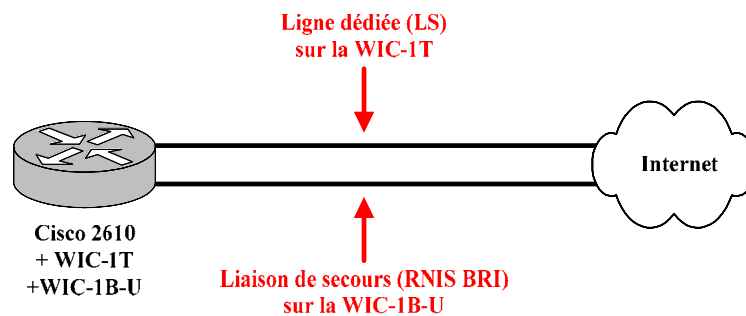


- **Drapeau** : Similaire au champ HDLC.
- **Adresse** : Peut comporter 1 ou 2 octets (Dépend de la valeur des bits EA).
 - o **SAPI** : Bits d'identification du point d'accès (6 bits). Indique le portail où les services LAPD sont fournis à la couche 3.
 - o **C/R** : Bit de commande/réponse.
 - o **EA** : Bit d'adressage étendu. Si le premier EA est défini, alors l'adresse comporte 1 octet, sinon elle en comporte 2.
 - o **TEI** : Identificateur de point d'extrémité de terminal. Ce champ précise le nombre de terminaux, ou s'il s'agit d'un broadcast.
- **Contrôle** : Similaire au champ HDLC.
- **Données** : Données fournies par l'intermédiaire des canaux B.
- **FCS** : Séquence de contrôle de trame (Contrôle d'erreurs).

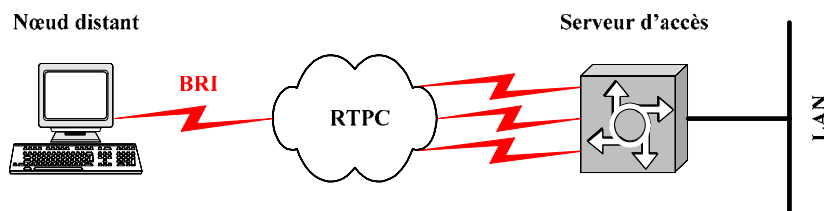
32.4. Utilisation/implémentation

La technologie RNIS a de nombreuses applications :

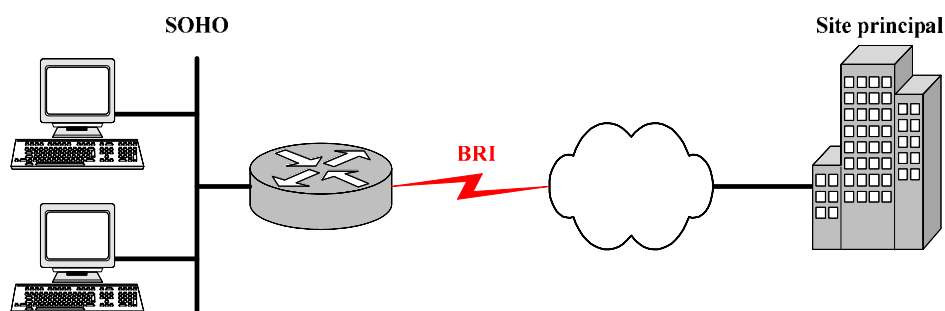
- Solution alternative aux lignes dédiées.
- Accès à distance :
 - o Nœuds distants.
 - o Connectivité des petits bureaux et bureaux à domicile (SOHO – Small Office / Home Office).



L'utilisation du RNIS en tant qu'alternative aux lignes dédiées permet d'avoir une continuité de service en cas de défaillance de la liaison principale. L'utilisation de la liaison de secours se fait automatiquement, car la route ayant une meilleure métrique passant par la liaison principale sera désactivée, laissant ainsi comme seul choix le passage par la liaison de secours.



L'accès à distance pour un nœud isolé (Employés itinérants, etc.) permet une connectivité éphémère. L'environnement présenté à l'utilisateur est identique à celui qu'il verrait s'il était en local (Utilisation du VPN). La seule différence pour le nœud distant est que la liaison est relativement lente comparée à celle d'un LAN, et passe par l'intermédiaire d'un serveur d'accès, qui fournit les services LAN.



L'accès à distance pour une SOHO (Succursale de l'entreprise, etc.) permet à un petit groupe d'utilisateurs d'avoir un accès aux ressources du site principal. C'est le routeur de la SOHO qui s'occupe de la translation d'adresse, afin de fournir des services à plusieurs travailleurs en utilisant une seule connexion WAN (Une seule IP).

32.5. Routage à établissement de la connexion à la demande (DDR)

Le principe du DDR est d'ouvrir ou de fermer dynamiquement une session de communication, et ce sur une liaison WAN de type commutation de circuits (Exemples : POTS, RNIS).

La notion de trafic intéressant pour le DDR est un trafic, ou ensemble de paquets, que le routeur doit acheminer par le biais de la liaison WAN. Ceci peut être basé :

- Sur les adresses de couche 3.
- Sur les services réseaux spécifiques, en se basant sur les numéros de port des protocoles de couche 4.

Principe de fonctionnement du DDR :

- Lorsque le routeur reçoit un trafic intéressant, il va ouvrir une session, afin de transmettre ce trafic.
- Cette session sera fermée après expiration du délai du compteur d'inactivité.
- Ce compteur d'inactivité est réinitialisé uniquement si un trafic intéressant est reçu.

Les avantages du DDR sont nombreux :

- Plus économique que des liaisons spécialisées ou multipoints, lorsque le trafic devant être émis ne nécessite pas un circuit continu.
- Partage de charges, lorsque l'on a par exemple plusieurs liaisons séries, ce qui permet d'utiliser le nombre de liaison nécessaire uniquement. Dans ce cas, il faudrait configurer le DDR afin d'ouvrir la session uniquement lorsque la liaison précédente est surchargée.
- Liaison de secours pour une liaison spécialisée. Le DDR permet d'offrir un moyen de communication de secours en cas de défaillance de la liaison principale (liaison spécialisée).

Le trafic empruntant une liaison utilisant le DDR est moins important et plus intermittent que le trafic passant au travers d'un réseau LAN ou par une liaison spécialisée.

Les étapes de la configuration du DDR sur un routeur sont les suivantes :

- **Utilisation des ACL** : Permet de préciser les adresses de couche 3 (source et destination), ainsi que les protocoles de couche 4 et numéro de port associés. Cela définit ce que nous voulons considérer comme trafic intéressant.
- **Définition des interfaces utilisant le DDR** : Indique le groupe de numérotations qui associe l'interface WAN voulue avec les ACL pour le DDR.

32.6. Commandes

Les commandes qu'il est nécessaire de connaître en vue de pouvoir configurer un routeur branché sur une liaison RNIS sont :

- **interface bri {numéro} :**
 - o Mode de configuration globale.
 - o Permet de passer dans le mode de configuration d'une interface BRI.
- **interface dialer {numéro} :**
 - o Mode de configuration globale.
 - o Permet de passer dans le mode de configuration d'une interface de connexion à la demande.
- **isdn switch-type {isdn_swith_type} :**
 - o Mode de configuration globale.
 - o Permet de spécifier le type de commutateur RNIS sur lequel on est raccordé.
 - o Le paramètre **isdn_switch_type** peut prendre les valeurs **basic-1tr6** (Allemagne), **basic-5ess** (USA), **basic-dms100** (Angleterre), **basic-net3** (Angleterre et Europe), **basic-ni**, **basic-qsig**, **basic-ts013** (Australie), **ntt** (Japon), **vn3** (France).
- **isdn spid1 {valeur_spid_1} :**
 - o Mode de configuration d'interface BRI.
 - o Configure le SPID pour le canal B1.
- **isdn spid2 {valeur_spid_2} :**
 - o Mode de configuration d'interface BRI.
 - o Configure le SPID pour le canal B2.
- **dialer-list {numéro_groupe} protocol {proto} {permit | deny | list {numéro_acl}} :**
 - o Mode de configuration globale.
 - o Cette commande permet de définir le trafic intéressant pour le DDR.
 - o Le paramètre **numéro_groupe** indique le groupe pour lequel on attribut le trafic intéressant.
 - o **proto** permet de spécifier le protocole de couche 3 dont fera partie le trafic intéressant.
 - o Le dernier paramètre permet de rentre intéressant tout le protocole spécifié (**permit**), tout sauf le protocole spécifié (**deny**), ou bien de limiter le trafic intéressant à tout ce qui correspond à l'ACL indiquée (**list**).
- **dialer-group {numéro_groupe} :**
 - o Mode de configuration d'interface BRI ou Dialer.
 - o Permet d'affecter un trafic intéressant spécifique (**dialer-list** correspondant) sur l'interface actuelle.
- **dialer pool {numéro} :**
 - o Mode de configuration d'interface Dialer.
 - o Permet le regroupement d'interfaces Dialer sur une interface BRI spécifique (**dialer pool-member**).
- **dialer pool-member {numéro} :**
 - o Mode de configuration d'interface BRI.
 - o Permet de spécifier l'interface BRI qui sera la source des interfaces Dialer (**dialer pool**).
- **dialer string {numéro} :**
 - o Mode de configuration d'interface Dialer.
 - o Permet de configurer le numéro de téléphone de la destination à appeler.
- **dialer wait-for-carrier-time {temps} :**
 - o Mode de configuration d'interface BRI ou Dialer.
 - o Configuration du temps pendant lequel le routeur attendra le signal de porteuse.
- **dialer idle-timeout {temps} :**
 - o Mode de configuration d'interface BRI ou Dialer.
 - o Configuration du temps de déconnexion après inactivité.

- **dialer remote-name {nom_distant} :**
 - o Mode de configuration d'interface Dialer.
 - o Permet de spécifier le nom d'hôte du nœud distant.
- **dialer in-band :**
 - o Mode de configuration d'interface BRI ou Dialer.
 - o Indique que l'on va faire passer le flux de signalisation dans le canal de données
- **dialer map {protocole} {adresse} name {nom} {numéro} :**
 - o Mode de configuration d'interface BRI ou Dialer.
 - o Précise le numéro de téléphone à appeler pour atteindre l'adresse de destination indiquée.
 - o Ne pas utiliser cette commande avec la commande **dialer string** en même temps.
- **dialer load-threshold {charge} [inbound | outbound | either] :**
 - o Mode de configuration d'interface.
 - o Spécifie à quel pourcentage de charge de la liaison un nouveau canal B sera utilisé (Uniquement avec PPP), que ce soit en entrée (**inbound**), sortie (**outbound**) ou les deux (**either**).
 - o Charge doit être un nombre entre 1 et 255 (255 = 100 %).
- **PPP multilink :**
 - o Mode de configuration d'interface.
 - o Indique que le protocole PPP sur l'interface courante pourra prendre en charge la gestion de liaisons multiples.

Afin de permettre une résolution des problèmes éventuels ainsi qu'une surveillance de l'état des protocoles et des connexions, IOS fournit différentes commandes :

- **show interfaces bri {numéro}:{bearer} :** Permet de visualiser l'état d'un canal B particulier de l'interface BRI voulue.
- **show isdn status :** Etat de la liaison RNIS. Cette commande indique le type de commutateur RNIS configuré, les statuts au niveau des couches 1 et 2, ainsi que le nombre de connexions actives sur la liaison.
- **show isdn active :** Affichage des connexions actives.
- **show dialer :** Affichage des paramètres et des statistiques concernant l'interface DDR (Dialer).
- **debug isdn events :** Permet d'obtenir des informations sur les événements RNIS.
- **debug isdn q921 :** Permet la vérification d'une connexion au commutateur RNIS (Problèmes liés aux SPID).
- **debug isdn q931 :** Permet d'identifier les problèmes entre le routeur et le commutateur (Problème lié à une mauvaise configuration du type de commutateur RNIS).
- **debug dialer [events | packets] :** Permet une visualisation sur l'état du DDR.

32.7. Configuration

On peut choisir entre plusieurs types d'encapsulation lors de la configuration d'une liaison RNIS :

- HDLC (Par défaut).
- PPP (Généralement utilisé).

Les tâches à accomplir sont :

- Détermination du type de commutateur RNIS sur lequel on est relié.
- Choix de l'encapsulation pour notre liaison (HDLC, ou PPP avec ou sans authentification).
- Définir les SPID pour les canaux B (Si nécessaire).
- Configurer une ou plusieurs interfaces Dialer, en fonction des besoins :
 - o Indiquer le numéro à appeler.
 - o Indiquer le rattachement de l'interface Dialer courante à une interface BRI.
 - o Préciser le type de trafic qui devra être transmis (DDR).
 - o Créer une route statique pour diriger le trafic sur la bonne interface.

33. Technologie Frame Relay

33.1. Technologie

La technologie Frame Relay dispose des caractéristiques suivantes :

- Destinée pour des équipements numériques haut de gamme et à haut débit.
- Fonctionne au niveau des couches 1 et 2 du modèle OSI.
- Utilise des circuits virtuels dans un environnement commuté.
- Technologie à commutation de paquets, et à accès multiples.
- L'ETTD et l'ETCD sont respectivement généralement le routeur client et le commutateur de l'opérateur.
- Remplace des réseaux point-à-point, trop coûteux.
- Se base sur l'encapsulation HDLC.
- Utilise le multiplexage pour partager la bande passante totale du nuage Frame Relay.

Cette technologie comporte quelques inconvénients, dont :

- Capacité de vérification des erreurs et fiabilité minimale (laissées aux protocoles de couches supérieures).
- Affecte le fonctionnement de certains aspects (Split Horizon, broadcasts, etc.).
- Ne diffuse pas les broadcasts. Pour en effectuer, il faut envoyer un paquet à chaque destination du réseau.

Un réseau Frame Relay peut être conçu suivant deux topologies :

- **Maillage global** : Chaque extrémité est reliée par l'intermédiaire d'un PVC distinct vers chaque autre destination.
- **Maillage partiel** : Egalement appelé topologie en étoile ou "hub-and-spokes". Chaque extrémité n'est pas reliée à toutes les autres.

Définitions :

- **Tarif d'accès** : Vitesse d'horloge de la connexion.
- **DLCI (Identificateur de connexion de liaison de données)** : C'est un numéro désignant un point d'extrémité. Le commutateur Frame Relay mappe deux DLCI (Source et destination) afin de créer un PVC. Il a une portée locale.
- **PVC (Circuit virtuel permanent)** : Circuit virtuel agissant comme une liaison point-à-point dédiée pour relier deux extrémités dans un environnement commuté.
- **LMI (Interface de supervision locale)** : Norme de signalisation entre le point d'extrémité et le commutateur Frame Relay chargé de la gestion et maintenance de l'état entre les unités.
- **CIR (Débit de données garanti)** : Débit de données que le fournisseur s'engage à fournir.
- **Bc (Débit garanti en rafale)** : Nombre maximum de bits que le commutateur accepte de transférer sur une période donnée.
- **Be (Débit garanti en excès)** : Nombre maximum de bits non garantis que le commutateur tentera de transférer au-delà du CIR. Il est généralement limité par la vitesse du port de la boucle locale. Les trames émises en excès ont leur bit d'éligibilité à la suppression mis à 1.
- **FECN (Notification explicite de congestion au destinataire)** : Bit défini dans une trame qui signale à l'unité réceptrice de lancer des procédures de prévention de congestion.
- **BECN (Notification explicite de congestion à la source)** : Idem mais pour l'unité source. Un routeur recevant cette notification réduira le débit de transmission de 25%.
- **Bit d'éligibilité à la suppression** : Bit qui indique que la trame peut être supprimée en priorité en cas de congestion.

Le format des trames Frame Relay est le suivant :

1 octet	2 octets	Variable	2 octets	1 octet
Drapeau	Adresse	Données	FCS	Drapeau

- **Drapeau** : Indique le début et la fin de la trame.
- **Adresse** : Contient l'adresse d'extrémité (10 premiers bits), ainsi que les mécanismes de notification de congestion (3 derniers bits).
 - o **DLCI**.
 - o **FECN**.
 - o **BECN**.
 - o **Bit d'éligibilité à la suppression**.
- **Données** : Informations encapsulées de couche supérieure.
- **FCS** : Séquence de contrôle de trame.

33.2. Interface LMI & DLCI

La mise en œuvre et le fonctionnement de la technologie Frame Relay repose essentiellement sur les interfaces LMI, dont les fonctions de base sont :

- Déterminer la fonctionnalité des PVC connus du routeur.
- Transmettre des messages de veille, pour éviter que le PVC ne se ferme pour cause d'inactivité.
- Indiquer au routeur les PVC disponibles.

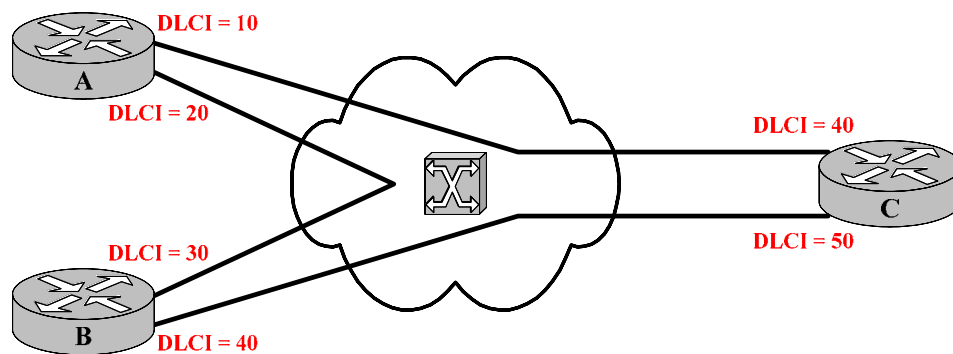
Il existe des extensions LMI, qui sont optionnelles :

- **Messages d'état des circuits virtuels** (Extension universelle) : Signalisation périodique sur les PVC (Nouveaux, supprimés, leur intégrité, etc.).
- **Diffusion multicast** (Extension facultative) : Permet la diffusion des messages de protocole de routage et ARP, qui doivent être normalement transmis à plusieurs destinataires. Cela utilise les DLCI 1019 à 1022.
- **Adressage global** (Extension facultative) : Portée globale des DLCI au lieu d'être locale. Permet d'avoir un DLCI unique sur le réseau Frame Relay.
- **Contrôle de flux simple** (Extension facultative) : Contrôle de flux de type XON/XOFF, destiné aux unités dont les couches supérieures ne peuvent pas utiliser les bits de notification de congestion, mais nécessitant un niveau de contrôle de flux.

1 octet	2 octets	1 octet	1 octet	1 octet	1 octet	Variable	2 octets	1 octet
Drapeau	DLCI LMI	Indicateur d'informations non numéroté	Indicateur de protocole	Référence d'appel	Type de message	Éléments d'information	FCS	Drapeau

Le schéma ci-dessus représente une trame Frame Relay spécifique aux messages LMI.

- **DLCI LMI** : DLCI pour les messages LMI. Il est fixé à 1023.
- **Indicateur de protocole** : Défini sur une valeur précisant l'interface LMI.
- **Type de message** : Deux types ont été définis, qui permettent de vérifier l'intégrité des liaisons logiques et physiques.
 - o **Message d'état** : Emis en réponse à un message de demande d'état. Message de veille ou message d'état sur chaque DLCI défini pour la liaison.
 - o **Message de demande d'état**.
- **Éléments d'information (IE)** : Contient un ou plusieurs éléments d'information d'1 octet chacun, et un ou plusieurs octets de données.



Les identificateurs DLCI sont reconnus localement, ce qui implique qu'ils ne sont pas forcément uniques dans le nuage Frame Relay (Exception faite si on utilise l'extension LMI d'adressage global). Deux unités ETTD peuvent utiliser une valeur DLCI identique ou différente pour désigner le PVC les reliant.

L'espace d'adressage DLCI est limité à 10 bits. Une partie de la plage d'adresse (0 à 1023) est utilisable pour les adresses d'extrémité (Transport des données utilisateur), et le reste est réservé à des fins d'implémentation par le constructeur (Messages LMI, adresses de multicast, etc.).

La portion exploitable de la plage d'adresse DLCI est définie par le type LMI utilisé :

- **ansi** : La plage de DLCI hôte va de 16 à 992.
- **cisco** : Les DLCI hôte vont de 16 à 1007.
- **q933a** : Même plage DLCI que la version **ansi**.

33.3. Fonctionnement, table de commutation & processus de transmission

La norme Frame Relay de base ne supporte que des PVC reconnus localement. Il n'y a pas d'adresses pour désigner les nœuds distants. Il est donc impossible d'utiliser un processus classique de résolution d'adresses. Pour palier à ce problème, il y a deux possibilités :

- Créer manuellement des cartes statiques avec la commande **frame-relay map**.
- Opter pour l'extension LMI sur l'adressage global. Ainsi, chaque nœud aura un DLCI unique.

La carte Frame Relay comporte trois champs :

- DLCI local par lequel passer pour atteindre la destination.
- L'adresse de couche 3 du nœud distant correspondant.
- L'état de la connexion :
 - **Active state** : Connexion active. Les routeurs peuvent échanger des données.
 - **Inactive state** : La connexion locale au commutateur est en service, mais la connexion du routeur distant au commutateur ne l'est pas.
 - **Deleted state** : Soit aucun LMI n'est reçu du commutateur, soit aucun service n'est assuré entre le routeur local et le commutateur.

Il existe un mécanisme de résolution d'adresse inverse (Inverse-ARP), qui permet à un routeur d'élaborer automatiquement la carte Frame Relay :

- Le routeur prend connaissance des DLCI au moment de l'échange LMI initial avec le commutateur.
- Il envoie alors une requête Inverse-ARP à chaque DLCI pour chaque protocole de couche 3 configurés localement.
- Les informations renvoyées sont utilisées pour remplir la carte Frame Relay.

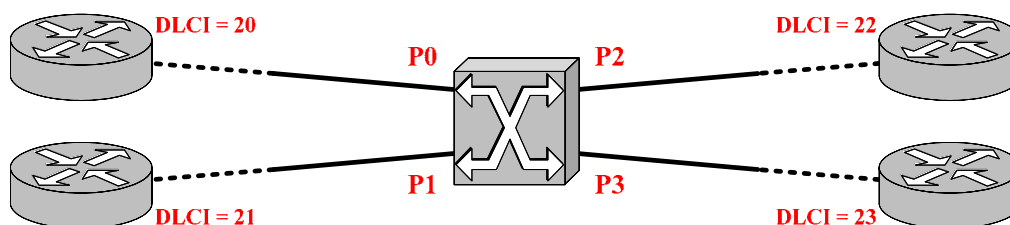


Table de commutation du port P0

IN_Port	IN_DLCI	OUT_Port	OUT_DLCI
P0	20	P1	21
		P2	22
		P3	23

La table de commutation Frame Relay dispose de quatre colonnes :

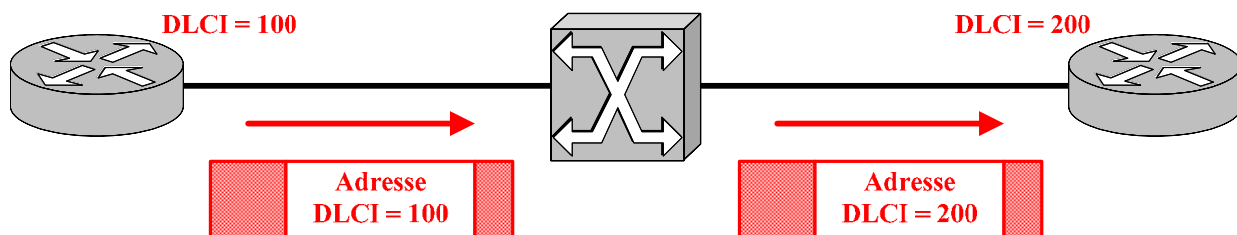
- Port d'entrée.
- DLCI d'entrée.
- Port de sortie.
- DLCI de sortie.

Cette table de commutation est basée sur un port du commutateur, il y a donc autant de tables qu'il y a de ports fonctionnels. De plus, elle est administrée, ce qui signifie que c'est l'opérateur qui décide du contenu de chaque table. Elle sert :

- Au moment du premier échange LMI, afin d'informer le routeur des DLCI des nœuds distants qui lui sont accessibles.
- Durant la transmission des données, où elle fonctionne comme une table de commutateur LAN.

Le processus de découverte est le suivant :

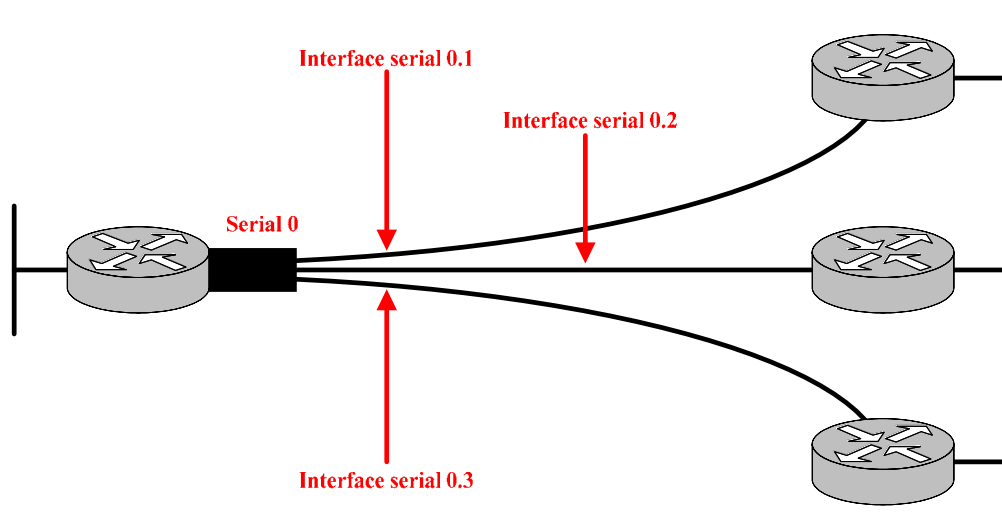
- Émission d'un message de demande d'état au commutateur Frame Relay (donne l'état du routeur local et demande celui des connexions des routeurs distants).
- Le commutateur répond avec un message d'état, contenant les DLCI des routeurs distants qui sont accessibles au routeur local.
- Pour chaque DLCI actif, le routeur envoie un paquet Inverse-ARP afin de se présenter et de demander aux routeurs distants de s'identifier (Adresse de couche 3).
- Le routeur mappe dans sa carte chaque adresse de nœud distant qu'il reçoit par le biais d'un message de résolution d'adresse inverse.
- Les messages de résolution d'adresse inverse sont ensuite échangés toutes les 60 secondes.
- Les messages de vieille sont envoyés toutes les 10 secondes au commutateur.



Le processus de transmission de données au travers d'un réseau Frame Relay est :

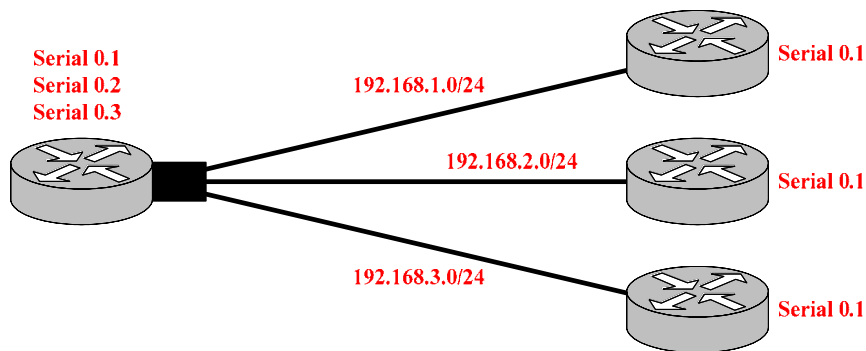
- Le routeur source encapsule les données à transmettre dans une trame Frame Relay, dont la valeur du champ Adresse correspond au DLCI de l'émetteur, puis l'envoie.
- Le commutateur reçoit cette trame, et utilise la table de commutation du port d'entrée afin de déterminer le port de sortie, et donc le DLCI de sortie.
- Le commutateur modifie la trame en plaçant le DLCI de la source, afin que la destination puisse savoir quelle est cette source.
- Le routeur de destination reçoit la trame émise par le commutateur. Il répondra, si besoin est, en émettant une trame vers le DLCI indiqué dans la trame reçue.

33.4. Sous-interfaces Frame Relay



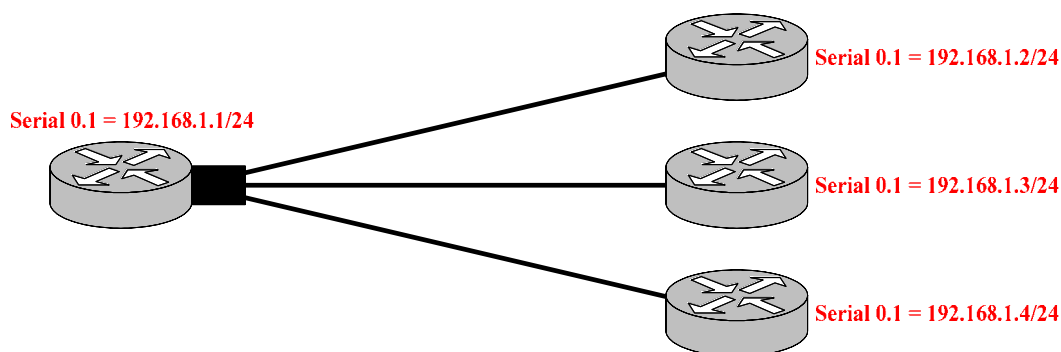
Les sous-interfaces sont des subdivisions logiques d'une interface physique et peuvent être de deux types :

- **Point-à-point.**
- **Multipoint.**



Les caractéristiques des sous-interfaces point-à-point sont :

- Une sous-interface par PVC.
- Une attribution statique de DLCI par sous-interface.
- Chaque connexion point-à-point est son propre sous-réseau.
- Chaque interface possède un seul DLCI.
- Split horizon ne fonctionne pas comme on voudrait qu'il fonctionne dans le principe, car il ne connaît pas le principe de sous-interface, ce qui veut dire que les mises à jour de routage ne seront pas propagées vers les autres sous-interfaces.



Les caractéristiques des sous-interfaces multipoints sont :

- Une seule sous-interface pour établir plusieurs PVC.
- Autant d'attributions statiques de DLCI qu'il y a de PVC (Destinataires).
- Toutes les interfaces font partie du même sous-réseau.
- Chaque interface possède son DLCI local.
- Split horizon fonctionne avec ce type de sous-interface.

33.5. Commandes

Les commandes concernant Frame Relay sont les suivantes :

- **interface serial {numéro} :**
 - Mode de configuration globale.
 - Permet de passer dans le mode de configuration de l'interface souhaitée.
- **interface serial {numéro.sous-numéro} {multipoint | point-to-point} :**
 - Mode de configuration globale.
 - Permet de passer dans le mode de configuration de la sous-interface souhaitée.
 - Le paramètre **multipoint** ou **point-to-point** définit le type de sous-interface utilisée.
 - Il faut utiliser **multipoint** si on veut que le routeur envoie les broadcast et les mises à jour de routage qu'il reçoit.
- **encapsulation frame-relay [ietf] :**
 - Mode de configuration d'interface.
 - Précise l'encapsulation des trames pour l'interface courante.
 - Le paramètre **cisco** est la valeur par défaut, et est à utiliser si on est raccordée à un autre équipement Cisco.
 - Le paramètre **ietf** est utile pour se connecter à un dispositif non Cisco.
- **frame-relay interface-dlci {dlci} :**
 - Mode de configuration de sous-interface.
 - Affecte un DLCI pour la sous-interface courante.
- **frame-relay local-dlci {dlci} :**
 - Mode de configuration d'interface.
 - Permet d'affecter manuellement le DLCI pour l'interface courante (normalement attribué automatiquement par le LMI).
 - Il faut utiliser cette commande dans les environnements ne supportant pas les interfaces LMI.
- **frame-relay lmi-type {ansi | cisco | q933a} :**
 - Mode de configuration d'interface.
 - La valeur **cisco** est par défaut.
 - Cette commande est à utiliser uniquement pour une version d'IOS ancienne car, avec les versions 11.2 et ultérieure, le type de LMI est détecté automatiquement.
- **bandwidth {bp} :**
 - Mode de configuration d'interface.
 - Permet de spécifier la bande passante de la liaison sur un ETDD, à titre d'information (Pour un protocole de routage).
- **frame-relay inverse-arp {protocole} {dlci} :**
 - Mode de configuration d'interface.
 - Active la résolution d'adresse inverse pour le protocole de couche 3 indiqué en paramètre.
 - Cette résolution est active par défaut.
- **frame-relay map {protocole} {adresse} {dlci} [broadcast] :**
 - Mode de configuration d'interface.
 - Permet de mapper localement une adresse de couche 3 distante avec le DLCI local par lequel passer pour atteindre cette destination.

- **frame-relay intf-type {dte | dce | nni}** :
 - o Mode de configuration d'interface.
 - o Permet d'expliciter le type d'interface Frame Relay locale.
 - o La valeur par défaut est **dte**.
 - o **dce** est à utiliser pour l'interface du commutateur Frame Relay reliée au DTE (ETTD), et **nni** est pour les interfaces reliant les commutateurs Frame Relay.
- **frame-relay switching** :
 - o Mode de configuration globale.
 - o Permet d'activer la commutation de PVC sur une unité ETCD (Commutateur Frame Relay).
 - o Active l'interface LMI.
- **frame-relay route {dlci_src} interface {type} {numéro} {dlci_dest}** :
 - o Mode de configuration d'interface.
 - o Permet de créer une entrée dans la table de commutation Frame Relay.
 - o Il faut indiquer le DLCI source, l'interface locale de sortie et celui de la destination.
 - o Cette commande est à utiliser sur un commutateur Frame Relay uniquement.

IOS met à notre disposition des commandes de visualisation d'état et de débogage afin de pouvoir vérifier le bon fonctionnement des points spécifiques à Frame Relay, ainsi que d'identifier les problèmes éventuels :

- **show interfaces serial {numéro}** : Affichage des informations sur les DLCI utilisés et sur l'indicateur de connexion de liaison de données LMI utilisé.
- **show frame-relay pvc** : Affichage de l'état de chaque connexion configurée ainsi que les statistiques sur le trafic. Cette commande permet aussi de savoir le nombre de paquets BECN et FECN reçus par le routeur.
- **show frame-relay map** : Affichage de l'adresse de couche 3 ainsi que le DLCI associé à chaque destination distante connectée au routeur local.
- **show frame-relay lmi** : Affichage des statistiques sur le trafic LMI.
- **show frame-relay route** : Affichage des routes Frame Relay configurées avec leur statut.
- **show frame-relay traffic** : Affichage des statistiques Frame Relay globales (Requêtes ARP, etc.).
- **debug frame-relay events** : Affichage des réponses aux requêtes ARP.
- **debug frame-relay lmi** : Affichage des échanges de paquets LMI entre le routeur et le commutateur.
- **debug frame-relay packet** : Analyse des paquets Frame Relay envoyés.

33.6. Configuration

La procédure de configuration d'une interface (DTE) en Frame Relay passe par les étapes suivantes :

- Passer dans le mode de configuration de l'interface voulue (Commande **interface serial {numéro}**).
- Définir une adresse de couche 3 (Commande **ip address {IP} {SM}**).
- Définir le type d'encapsulation (Commande **encapsulation frame-relay**).
- Définir le DLCI local en cas de non support de l'interface LMI (Commande **frame-relay local-dlci {dlci}**).
- Définir optionnellement la bande passante de la liaison (Commande **bandwidth {bp}**).
- Activer l'interface (Commande **no shutdown**).

Cette même procédure change un peu lorsqu'il s'agit de sous-interfaces :

- Passer dans le mode de configuration de l'interface voulue.
- Enlever toute adresse de couche 3 (Commande **no ip address**).
- Définir le type d'encapsulation.
- Passer dans le mode de configuration de la sous-interface voulue (Commande **interface serial {if.subif} {point-to-point | multipoint}**).
- Définir une adresse de couche 3.
- Définir le ou les DLCI locaux, car le LMI ne supporte pas les sous-interfaces (Commande **frame-relay interface-dlci {dlci}**).
- Définir optionnellement la bande passante de la liaison.
- Activer la sous-interface.

Il est possible de simuler un commutateur Frame Relay à l'aide d'un routeur. Les interfaces utilisées sont alors obligatoirement de type DCE. Pour ce faire, il faut utiliser une configuration distincte, et ce pour chaque interface :

- Activer la commutation Frame Relay sur le routeur (Commande **frame-relay switching**).
- Passer dans le mode de configuration de chaque interface utilisée.
- Enlever toute adresse de couche 3.
- Définir le type d'encapsulation.
- Définir la vitesse de fonctionnement de la liaison (Commande **clock rate {valeur}**).
- Définir le type d'interface Frame Relay.
- Définir une route pour chaque destinations accessibles depuis la source raccordée sur l'interface courante (Commande **frame-relay route {dlci_src} interface serial {numéro} {dlci_dest}**).
- Activer l'interface.