

Sommaire

1. INTRODUCTION AUX RESEAUX.....	5
1.1. CONNEXION A UN RESEAU.....	5
1.1.1. <i>Matériel.....</i>	5
1.2. SYSTEMES DE NUMERATION.....	6
1.2.1. <i>Représentation des données informatiques.....</i>	6
1.2.2. <i>Systèmes de numération.....</i>	7
1.2.3. <i>Conversions.....</i>	8
1.3. TERMINOLOGIE DE BASE DES RESEAUX.....	9
1.4. UNITES DE MESURE.....	10
2. MODELES OSI ET TCP/IP.....	11
2.1. MODELE OSI.....	11
2.2. MODELE TCP/IP.....	13
2.3. COMPARAISON ENTRE LE MODELE TCP/IP ET LE MODELE OSI.....	14
3. COUCHE 1 : MEDIAS ET EQUIPEMENTS RESEAU.....	15
3.1. LES NOTIONS DE BASE SUR LES SIGNAUX ET LE BRUIT DANS LES SYSTEMES DE COMMUNICATION.....	15
3.1.1. <i>Comparaison des signaux analogique et numériques.....</i>	15
3.1.2. <i>La représentation d'un bit dans un média physique.....</i>	16
3.1.3. <i>Les facteurs pouvant affecter un bit.....</i>	16
3.2. MEDIAS DE CUIVRES.....	18
3.2.1. <i>Le câble à paires torsadées non blindées.....</i>	18
3.2.2. <i>Le câble à paires torsadées blindées.....</i>	19
3.2.3. <i>Le câble coaxial.....</i>	20
3.2.4. <i>Les connecteurs RJ-45.....</i>	21
3.3. MEDIAS OPTIQUES.....	22
3.3.1. <i>Phénomènes physiques :.....</i>	22
3.3.2. <i>Composants optiques.....</i>	24
3.4. MEDIAS SANS FIL.....	27
3.4.1. <i>Fonctionnement d'un réseau sans fil.....</i>	27
3.4.2. <i>Authentification et sécurité.....</i>	28
3.4.3. <i>Modes d'implémentations.....</i>	29
3.5. EQUIPEMENTS DE COUCHE 1.....	30
3.5.1. <i>Répéteur.....</i>	30
3.5.2. <i>Concentrateur.....</i>	30
3.5.3. <i>Emetteur/récepteur.....</i>	30
3.6. LES TOPOLOGIES DE BASE UTILISEES DANS LES RESEAUX.....	31
3.6.1. <i>La topologie en bus.....</i>	31
3.6.2. <i>La topologie en anneau.....</i>	31
3.6.3. <i>La topologie en étoile.....</i>	32
3.6.4. <i>La topologie en étoile étendue.....</i>	32
3.6.5. <i>La topologie hiérarchique.....</i>	33
3.6.6. <i>La topologie complète (maillée).....</i>	33
4. COUCHE 2 : TECHNOLOGIES ETHERNET.....	34
4.1. INTRODUCTION AUX TECHNOLOGIES LAN.....	34
4.2. INTRODUCTION A ETHERNET.....	34
4.2.1. <i>Ethernet et le modèle OSI.....</i>	34
4.2.2. <i>Spécifications et normes.....</i>	34
4.2.3. <i>Trames Ethernet et IEEE 802.3.....</i>	35
4.3. FONCTIONNEMENT D'ETHERNET.....	36
4.3.1. <i>MAC.....</i>	36
4.3.2. <i>Erreurs possibles.....</i>	36

5. COUCHE 2 : COMMUTATION ETHERNET	38
5.1. DOMAINE DE COLLISION.....	38
5.2. SEGMENTATION	38
5.2.1. Segmentation par ponts.....	38
5.2.2. Segmentation par commutateurs.....	38
5.2.3. Spanning Tree	39
6. COUCHE 3 : PROTOCOLE IP	40
6.1. PROTOCOLES ROUTABLES	40
6.1.1. Protocoles orientés connexion et non orientés connexion	40
6.1.2. Protocoles routés	40
6.2. PROTOCOLE IP	41
6.2.1. Paquet IP.....	41
6.2.2. Adressage IP	41
6.2.3. Classes d'adresses IP.....	42
6.2.4. IPv4 et IPv6 (IPng / IP next generation).....	43
6.3. GESTION DES ADRESSES IP.....	43
6.3.1. Méthodes d'obtention.....	43
6.3.2. Résolution d'adresses.....	44
6.3.3. Le protocole ICMP.....	45
7. COUCHE 3 : SUBNETTING	46
7.1. INTERET DU SUBNETTING.....	46
7.2. METHODES DE CALCUL	46
7.2.1. Méthode classique.....	46
7.2.2. Méthode du nombre magique.....	48
8. COUCHE 3 : INTRODUCTION AU ROUTAGE.....	50
8.1. PRINCIPES FONDAMENTAUX.....	50
8.2. DOMAINE DE BROADCAST	50
8.3. LES EQUIPEMENTS DE COUCHE 3 : LES ROUTEURS	50
8.4. DETERMINATION DU CHEMIN	51
8.5. SYSTEMES AUTONOMES, IGP ET EGP	52
8.6. ROUTAGE STATIQUE ET DYNAMIQUE.....	52
9. COUCHE 4 : COUCHE TRANSPORT	54
9.1. INTRODUCTION	54
9.2. TCP ET UDP	54
9.2.1. Numéros de ports	55
9.2.2. Structures d'un segment TCP.....	55
9.2.3. Structure d'un datagramme UDP	56
9.3. METHODE DE CONNEXION TCP.....	56
9.3.1. Connexion ouverte/échange en 3 étapes	56
9.3.2. Positive Acknowledgement Retransmission	57
9.3.3. Fenêtrage	57
10. COUCHE 5 : COUCHE SESSION	58
10.1. CONTROLE DU DIALOGUE.....	58
10.2. SYNCHRONISATION DU DIALOGUE	59
10.3. DIVISION DU DIALOGUE	59
11. COUCHE 6 : COUCHE PRÉSENTATION.....	61
11.1. FONCTIONS ET NORMES.....	61
11.2. LE CRYPTAGE DES DONNEES	62
11.3. LA COMPRESSION DES DONNEES.....	62

12. COUCHE 7 : COUCHE APPLICATION	63
12.1. INTRODUCTION:	63
12.2. DNS.....	63
12.2.1. <i>Présentation du protocole DNS</i>	63
12.2.2. <i>Les noms d'hôtes et le « domain name system »</i>	64
12.2.3. <i>Codes des domaines internet</i>	64
12.3. FTP ET TFTP	65
12.3.1. <i>FTP</i>	65
12.3.2. <i>TFTP</i>	65
12.4. HTTP.....	65
12.5. SMTP	66
12.6. SNMP	66
12.7. TELNET	67
12.7.1. <i>Présentation du protocole Telnet</i>	67
12.7.2. <i>La notion de terminal virtuel</i>	67

1. Introduction aux réseaux

A l'origine, un réseau était un rassemblement de personnes ou d'objets. De nos jours on entend par réseau, les réseaux d'entreprises, qui connectent différentes machines afin de pouvoir les faire communiquer entre elles. Que ce soit pour le partage de fichiers ou l'envoi de messages, la plupart des entreprises sont aujourd'hui dotées d'un réseau afin d'être plus efficaces (il est quand même plus simple de transférer un fichier par Internet que de l'envoyer sur CD par la poste).

Au cours de cet essentiel nous allons étudier comment les informations (fichier, données, etc.) circulant sur des réseaux de petite taille (PAN, LAN) ou plus grande taille (MAN, WAN), ainsi que la connectique utilisée.

1.1. Connexion à un réseau

1.1.1. Matériel

Un ordinateur est composé de divers éléments. Avant de connecter votre ordinateur sur un réseau, il est nécessaire que vous connaissiez ce qui le compose, afin qu'en cas de panne vous sachiez identifier si cela provient du réseau ou non. De plus, cela vous permettra d'être plus familier avec une machine et pourra sûrement vous aider en cas de panne d'un ordinateur.

Voici la liste des différents composants de votre pc, ainsi que leurs descriptions :

Liste des composants	Description
Carte mère	La carte électronique principale dans un ordinateur. La carte mère contient les bus, le microprocesseur, et des circuits intégrés utilisés pour commander tous les périphériques extérieurs tels que le clavier, l'affichage graphique, les ports série et les ports parallèles, ou encore les ports USB ou Firewire.
Processeur	Puce de silicium effectuant tous les calculs arithmétiques et logiques dans un ordinateur. Il gère aussi les flux d'informations dans un ordinateur.
RAM (Random Access Memory)	Mémoire vive permettant de stocker les instructions en attente de traitement, autant que les données temporaires. Une fois l'ordinateur éteint cette mémoire se vide, contrairement au disque dur.
Disque Dur	Aussi appelé HDD (Hard Disk Drive en Anglais). Disque de stockage de données. C'est sur le disque dur que vous enregistrez vos données. Contrairement à la RAM, le disque dur conserve vos données même si l'ordinateur est éteint.
Bus	Canal de communication interne à un ordinateur par lequel transitent les données entre les différents composants.
Alimentation	Composant fournissant l'alimentation nécessaire à votre ordinateur.
ROM (Read Only Memory)	Mémoire accessible uniquement en lecture une fois la mémoire écrite. Ce genre de composant sert à stocker des informations qui ne doivent pas être effacées.
Lecteur de CD-ROM	Dispositif permettant de lire des CD-ROM

Il existe aussi des composants de fond de panier (backplane en Anglais) qui permettent d'ajouter des extensions à votre carte mère.

Liste des composants	Descriptions
Carte Vidéo	Carte d'extension permettant d'afficher un visuel sur un moniteur
Carte Son	Carte d'extension permettant de manipuler et de produire des sons via des hauts parleurs ou tout autre périphérique de sortie sonore (casque, etc.)
Carte Réseau (NIC/ Network Interface Card)	Carte d'extension permettant de relier physiquement un ordinateur à un réseau (LAN, WAN, etc.)
USB (Universal Serial Bus)	Port de connexion à chaud, vous permettant de brancher votre périphérique même si votre ordinateur est allumé. A noter que les transferts s'effectuent à haute vitesse.
Firewire	Norme concurrente de l'USB permettant aussi de connecter à chaud divers appareils et permettant des transferts à hautes vitesses.

1.2. Systèmes de numération

Lorsque les ordinateurs ont été créés, ils étaient fort coûteux du fait du nombre de composants qu'ils nécessitaient, en plus de leurs tailles impressionnantes.

Un ordinateur pourrait donc 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.

1.2.1. Représentation des données informatiques

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 (bytes en anglais), qui représente un caractère de données. Pour un ordinateur, un octet représente également un emplacement de mémoire adressable.

Par exemple, la représentation binaire des caractères du clavier et des caractères de contrôle est donnée dans le tableau des codes ASCII (American Standard Code for Information Interchange) dont voici un extrait :

Décimal	Hexadécimal	Octal	Binaire	Char
0	0	000	00000000	NUL
1	1	001	00000001	SOH
2	2	002	00000010	STX
3	3	003	00000011	ETX
4	4	004	00000100	EOT
7	7	007	00000111	BEL

Ce tableau nous présente les équivalences entre différents systèmes de numérotation que nous allons étudier par la suite. Si nous regardons la colonne « binaire », nous voyons que tous les caractères sont exprimés grâce à une combinaison de 8 bits pouvant prendre la valeur 0 ou la valeur 1.

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 =1024 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
Téraoctet (To)	1 téraoctet =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

1.2.2. Systèmes de numération

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

Ce système constitue la base du calcul pour les hommes, principalement parce que ces derniers ont 10 doigts. Nous utiliserons d'ailleurs ce système comme système de référence dans la suite du cours. Cependant, il existe d'autres systèmes de numérotation pouvant représenter des valeurs.

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

Un ordinateur, lui, utilise un système de numération basé sur la représentation du passage de courant, 0 (fermé) ou 1 (ouvert), dans un circuit électrique. Il faut se rappeler qu'à l'époque de l'expansion des ordinateurs, les composants à deux états ont participé à simplifier le traitement pour un ordinateur.

Autre système, 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. Ce système est utilisé pour simplifier les valeurs décimales trop grandes.

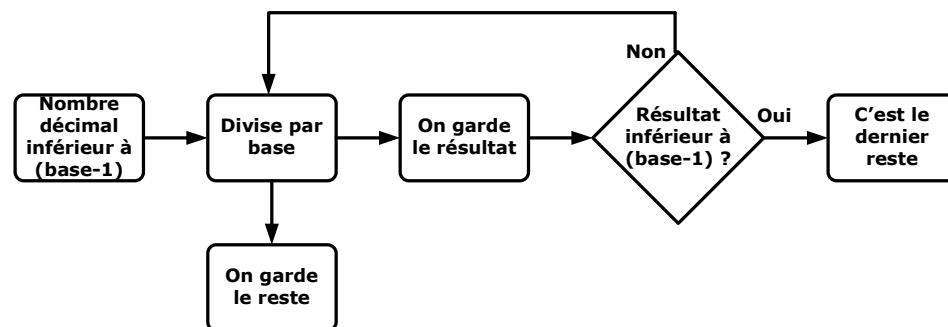
Il est évident ici de l'utilité de disposer de plusieurs systèmes d'informations. Une fois que l'on est familiarisé avec ces différents systèmes, la valeur A2F54B est plus facile à manipuler ou à mémoriser que son équivalent décimal.

1.2.3. Conversions

Entre ces bases il existe des méthodes de conversions :

- Décimal > Binaire
- Décimal > Hexadécimal
- Binaire > Décimal
- Hexadécimal > Décimal
- Binaire > Hexadécimal
- Hexadécimal > Binaire

Pour convertir du décimal vers une autre base, on utilise cette formule :



On divise notre nombre par la base à laquelle on veut le convertir et on continue tant que ce nombre n'est pas inférieur à la base. Il suffit ensuite de prendre les différents restes et de les concaténer du dernier vers le premier (de droite à gauche).

La conversion vers une base décimale se fait en décomposant le nombre en digit (chaque élément de la valeur). Et ensuite on multiplie chaque digit par la puissance de la base en commençant par celui le plus à droite avec une puissance zéro (si le nombre est une valeur hexadécimale alors on multipliera les digits par 16^0 , 16^1 , 16^2 , etc.). C'est donc l'ensemble des valeurs des différents digits ainsi multipliés qui forme la valeur en décimal, comme le montre cette formule

$$\sum_{i=0}^{i=n-1} (base^i \times valeur\ du\ digit)$$

Enfin, pour convertir du binaire vers l'hexadécimal, on prend un groupe de 4 bits et on les convertit en hexadécimal via les puissances de 2. Pour l'inverse, il suffit de faire exactement la même chose en utilisant la première formule comme si l'on convertissait en base 2, en utilisant des groupes de 4 bits ici aussi.

Hexadécimal	Binaire	Hexadécimal	Binaire
0	0000	8	1000
1	0001	9	1001
2	0010	A	1010
3	0011	B	1011
4	0100	C	1100
5	0101	D	1101
6	0110	E	1110
7	0111	F	1111

Tableau de conversion binaire/hexadécimale

1.3. 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.

La première classification de réseau que nous allons faire s'établit sur la base des distances entre les communicants.

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

Ces types de réseaux sont les plus courants, néanmoins il en existe d'autres, à l'instar des MAN (Metropolitan Area Network), qui connectent un ou plusieurs LANs dans une même région géographique. Ce type de réseau est en émergence du fait du développement des réseaux Wireless. On les trouve souvent en ville, situés dans les endroits publics.

Un autre type de réseau est le SAN (Storage Area Network) qui est une zone de stockage et de transfert de données.

Les SANs :

- Utilisent un réseau différent des hôtes afin de ne pas encombrer le trafic (ce type de réseau génère un important trafic).
- Permettent un taux de transfert nettement plus élevé entre serveurs, afin de permettre une réplication ou un mouvement des données plus aisé.
- Permettent de dupliquer des données entre serveurs jusqu'à une distance de 10 km.
- Utilisent diverses technologies qui permettent de ne pas tenir compte du système utilisé.

Un VPN (Virtual Private Network) est un réseau privé qui est construit dans une infrastructure de réseau public tel qu'Internet. Par Internet, un tunnel sécurisé peut être mis en place entre le PC de l'utilisateur et d'un routeur VPN se trouvant au siège social de l'entreprise, afin que celui-ci accède de chez lui au réseau de son entreprise.

1.4. Unités de mesure

La bande passante d'un réseau représente sa capacité, c'est-à-dire la quantité de données pouvant circuler en une période donnée sur le réseau. 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

À 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. Cette différence peut avoir pour raisons :

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

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

- Temps de téléchargement théorique(s) = Taille du fichier / bande passante
- Temps de téléchargement réel (s) = Taille du fichier (b) / débit

2. Modèles OSI et TCP/IP

2.1. Modèle OSI

La première évolution des réseaux informatiques a été des plus anarchiques, chaque constructeur développant sa propre technologie. Le résultat fut une quasi-impossibilité de connecter différents réseaux entre eux.

Pour palier à ce problème d'interconnexions, l'ISO (International Standards Organisation) décida de mettre en place un modèle de référence théorique décrivant le fonctionnement des communications réseaux.

Ainsi fût créé le modèle OSI, à partir des structures réseau prédominantes de l'époque : DECNet (Digital Equipment Corporation's Networking développé par digital) et SNA (System Network Architecture développé par IBM). Ce modèle a permis aux différents constructeurs de concevoir des réseaux interconnectables.

Le modèle OSI est un modèle conceptuel. Il a pour but d'analyser la communication en découpant les différentes étapes en 7 couches, chacune de ces couches remplissant une tâche bien spécifique :

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

Les 7 couches du modèle OSI sont les suivantes :

- **Couche 1 : Couche physique**
La couche physique définit les spécifications du média (câblage, connecteur, voltage, bande passante...).
- **Couche 2 : Couche liaison de donnée**
La couche liaison de donnée s'occupe de l'envoi de la donnée sur le média. Cette couche est divisée en deux sous-couches :
 - La sous-couche MAC (Média Access Control) est chargée du contrôle de l'accès au média. C'est au niveau de cette couche que l'on retrouve les adresses de liaison de donnée (MAC, DLCI).
 - La sous-couche LLC (Layer Link Control) s'occupe de la gestion des communications entre les stations et interagit avec la couche réseau.
- **Couche 3 : Couche réseau**
Cette couche gère l'adressage de niveau trois, la sélection du chemin et l'acheminement des paquets au travers du réseau.
- **Couche 4 : Couche transport**
La couche transport assure la qualité de la transmission en permettant la retransmission des segments en cas d'erreurs éventuelles de transmission. Elle assure également le contrôle du flux d'envoi des données.
- **Couche 5 : Couche session**
La couche session établit, gère et ferme les sessions de communications entre les applications.
- **Couche 6 : Couche présentation**
La couche présentation spécifie les formats des données des applications (encodage MIME, compression, encryptions).

- **Couche 7 : Couche application**

Cette couche assure l'interface avec les applications, c'est la couche la plus proche de l'utilisateur.

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

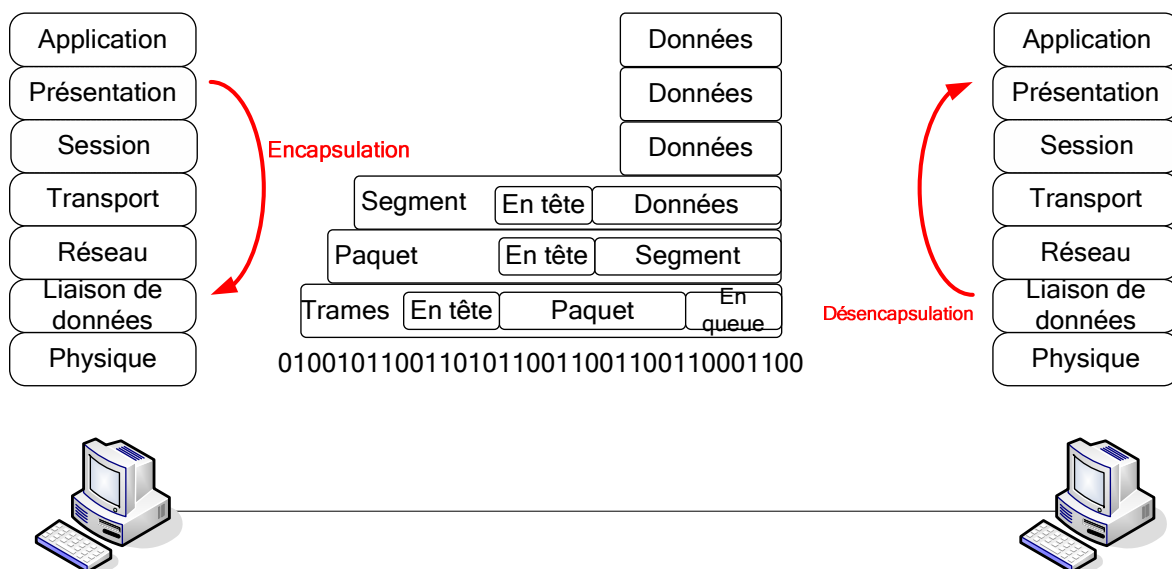
Figure 1- Les 7 couches du modèle OSI

Les avantages de ce modèle sont :

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

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 :



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 l'émetteur construit 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 ou 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 PDU (Unité de données de protocole) est utilisée.

Couche	Designation
7	Données
6	Données
5	Données
4	Segments
3	Paquets
2	Trames
1	Bits

2.2. Modèle TCP/IP

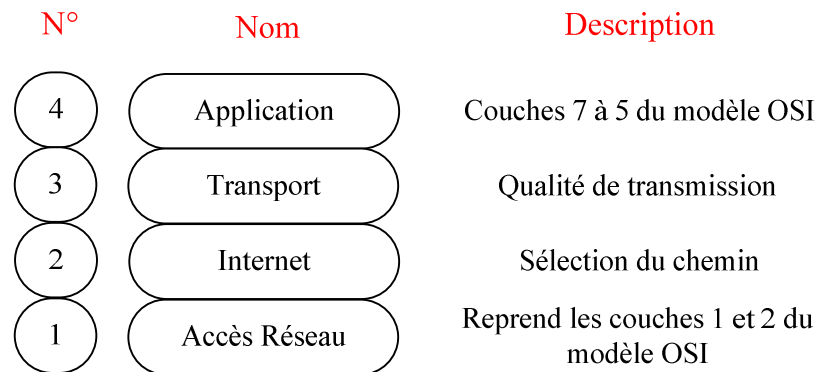
La forme actuelle de TCP/IP résulte du rôle historique que ce système de protocoles a joué dans le parachèvement de ce qui allait devenir Internet. À l'instar des nombreux développements de ces dernières années, Internet est issu des recherches lancées par le DOD (Department Of Defense), département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

À la fin des années 60, les officiels du DOD se rendirent compte que les militaires du département de la défense possédaient une grande quantité de matériel informatique très divers, mais ces machines travaillaient pour la plupart de manière isolée ou encore en réseaux de taille très modeste avec des protocoles incompatibles entre eux, ceci rendant une interconnexion impossible.

Les autorités militaires se sont alors demandées s'il était possible, pour ces machines aux profils très différents, de traiter des informations mises en commun. Habités aux problèmes de sécurité, les responsables de la défense ont immédiatement réalisés qu'un réseau de grande ampleur deviendrait une cible idéale en cas de conflit. La caractéristique principale de ce réseau, s'il devait exister, était d'être non centralisée.

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 Network du DOD), qui allait devenir par la suite le système d'interconnexion de réseau qui régit ce que l'on appelle aujourd'hui Internet : TCP/IP.

TCP/IP est un modèle comprenant 4 couches :



2.3. Comparaison entre le modèle TCP/IP et le modèle OSI

Ces deux modèles 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 :

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

Modèle OSI

Couche	Désignation
Application	Couche Application
Présentation	
Session	
Transport	Couches flux de données
Réseau	
Liaison de données	
Physique	

Modèle TCP/IP

Couche	Désignation
Application	Protocoles
Transport	
Internet	Réseaux
Accès Réseau	

Les modèles OSI et TCP/IP

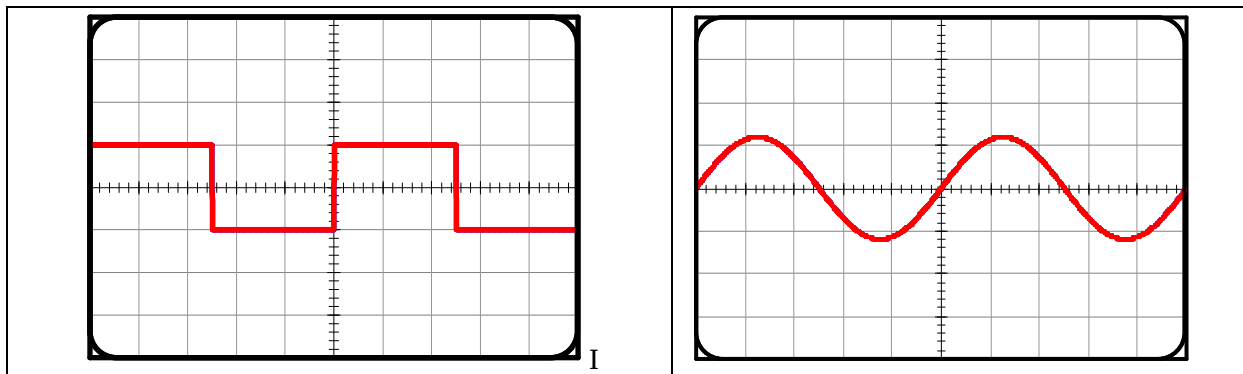
3. Couche 1 : Médias et équipements réseau

Ce chapitre a pour but de vous présenter les différentes connexions physiques entre ordinateurs.

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

3.1.1. Comparaison des signaux analogique et numériques

Lors de l'envoi de données sur un réseau, celles-ci transitent par des liaisons physiques, il convient donc d'observer comment sont-elles représentés dans ces liaisons.



Représentation d'un signal numérique et d'un signal analogique

Signal : tension électrique souhaitée, modèle d'impulsions lumineuses ou encore onde électromagnétique modulée. Il permet d'acheminer les données dans le média.

Le signal numérique dispose d'un graphique de tension que l'on va définir comme « sautillant », il se rapproche d'une onde carrée ou la tension passe quasi instantanément d'un état de basse tension à un état de haute tension.

Le signal analogique présente les caractéristiques suivantes :

- Il oscille
- Son graphique de tension varie constamment en fonction du temps et peut être représenté par une sinusoïde
- Il est utilisé pour les télécommunications depuis le début
 - Exemple : téléphone et radio

Les deux caractéristiques importantes d'une onde sont son amplitude (A), c'est-à-dire sa hauteur et sa longueur, ainsi que sa période. La fréquence de l'onde peut être calculée avec cette formule : $f = 1/T$.

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

Un bloc d'information est un élément binaire, connu sous le nom de bit ou impulsion. Un bit, dans un milieu électrique, est un signal correspondant à un 0 binaire ou à un 1 binaire. Cela peut être aussi simple que 0 (zéro) volts pour un 0 en binaire, et +5 volts pour un 1 binaire, ou un codage plus complexe.

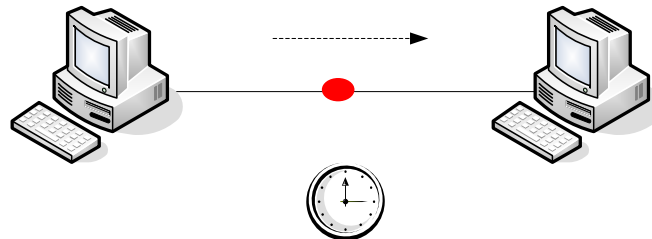
La mise à la terre de référence est un concept important concernant tous les médias de gestion réseau qui emploient des tensions pour diffuser des messages. C'est une masse électrique permettant d'établir une tension zéro dans un graphique de signalisation

3.1.3. Les facteurs pouvant affecter un bit

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

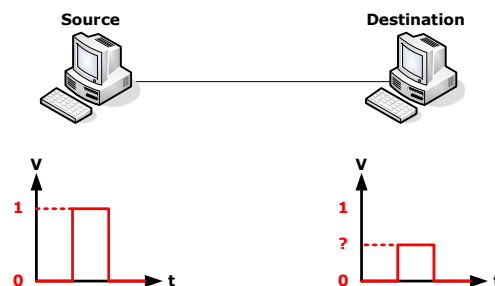
La propagation de signaux réseau :

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.



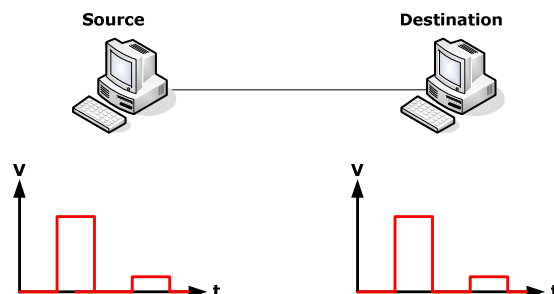
L'atténuation du signal réseau :

Perte de la force du signal. Ce problème est limitable par un bon choix des médias réseau utilisés



La réflexion réseau :

Retour d'énergie causée 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.



Le bruit :

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

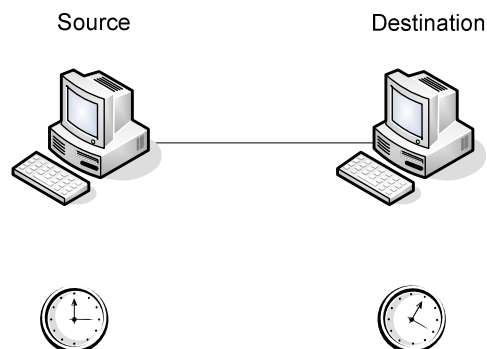
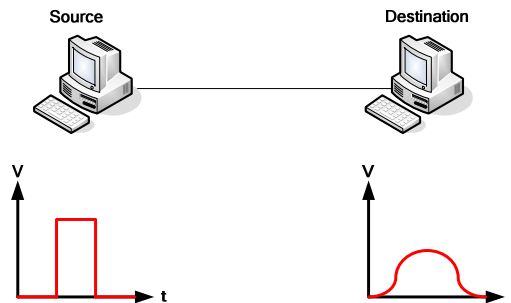
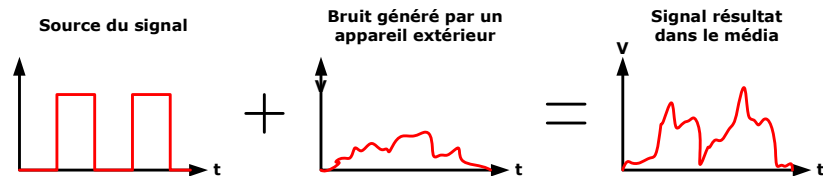
Le bruit peut être causé par des sources d'alimentations externes, des variations thermiques, des interférences électromagnétiques ou encore des interférences de radio fréquences.

La dispersion :

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.

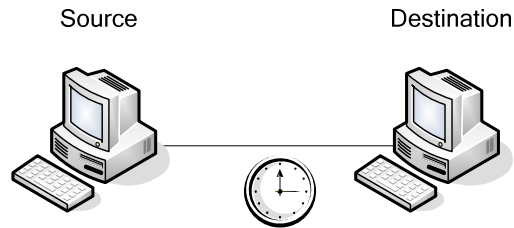
La gigue :

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 ».

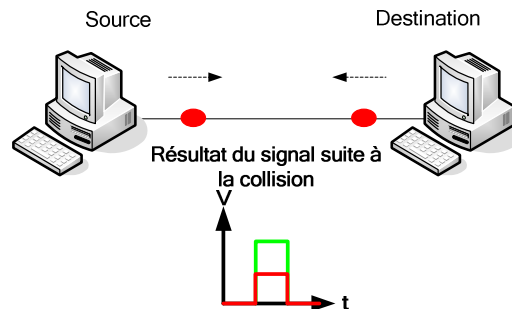


La latence :

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.

**Les collisions :**

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.

Il faut aussi savoir qu'une liaison entre 2 équipements A et B peut être :

- 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.
- 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).
- Full-duplex (bidirectionnelle simultanée) : A et B peuvent émettre et recevoir en même temps (comme dans le cas du téléphone).

3.2. Médias de cuivres

3.2.1. Le câble à paires torsadées non blindées

Le câble UTP est composé de 4 paires de fils torsadés 2 à 2, chacune de ses paires étant isolées 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.

Lorsque le câble à paires torsadées non blindées est utilisé comme média de réseau, il comporte quatre paires de fils de cuivre. La paire torsadée non blindée utilisée comme média de réseau a une impédance de 100 ohms. Ceci la différencie des autres types de câblage à paires torsadées comme ceux utilisés pour le câblage téléphonique.

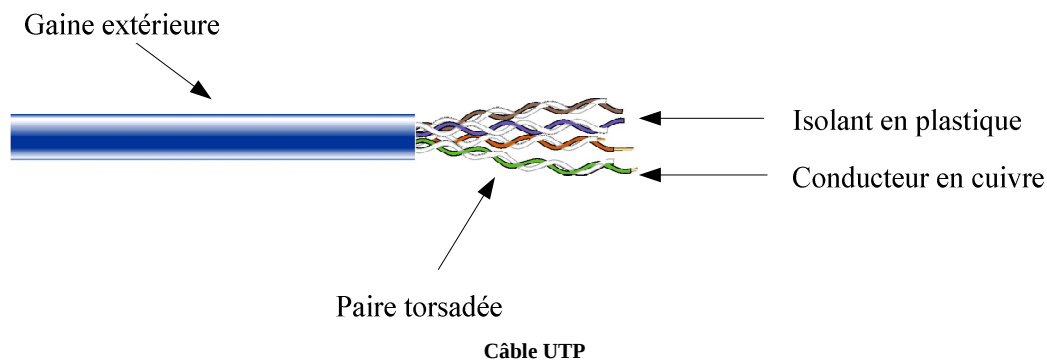
Comme le câble à paires torsadées non blindées a un diamètre extérieur de 0,43 mm et un coût relativement faible, sa petite taille peut s'avérer avantageuse lors d'une installation.

Avantages :

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

Inconvénient :

- Sensible aux interférences

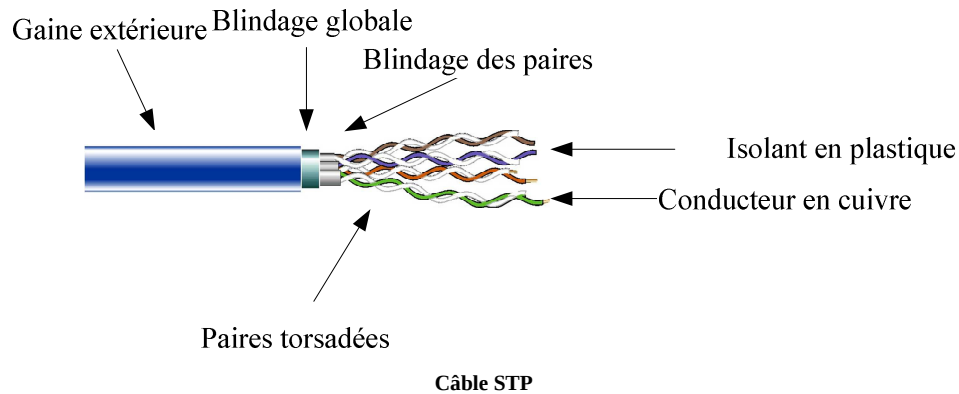


3.2.2. Le câble à paires torsadées blindées

Le câble à paires torsadées et blindées, ou STP, ajoute aux spécifications de l'UTP une méthode de blindage, d'annulation et de torsion de câbles. Comme le précise les spécifications pour les installations de réseau Ethernet, des câbles à paires torsadées blindées de 100 ohms correctement installés offrent une résistance à l'interférence électromagnétique, ainsi qu'à l'interférence de radiofréquences, sans toutefois augmenter sensiblement la taille ou le poids du câble.

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.



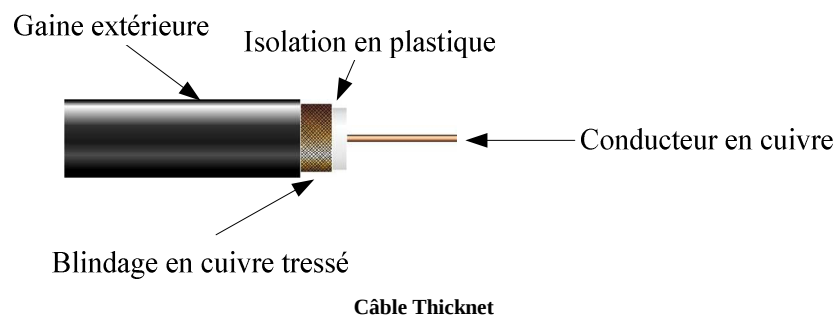
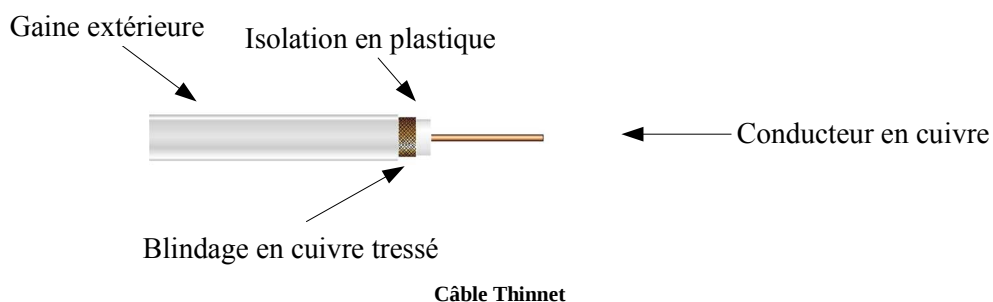
3.2.3. Le câble coaxial

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 :

- **Thicknet** : Epais et raide à cause de son blindage, il est recommandé pour l'installation de câble fédérateur. Sa gaine est jaune.
- **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.
- **Cheapernet** : Version économique et de faible diamètre du câble coaxial.



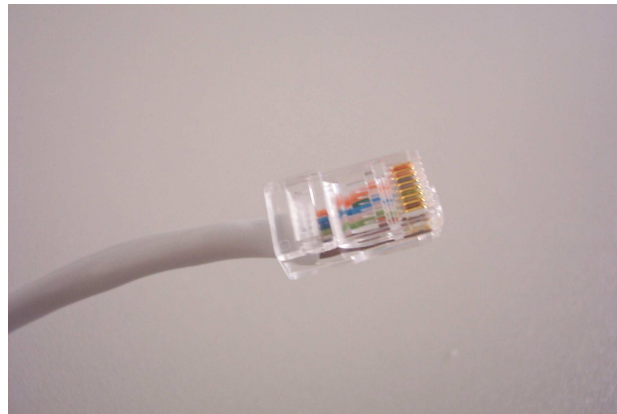
Il importe d'apporter une attention particulière à la mise à la terre. On doit assurer une solide connexion électrique aux deux extrémités du câble. Manquer à ce principe entraîne des parasites électriques qui causent une interférence au niveau de la transmission du signal du média réseau.

3.2.4. Les connecteurs RJ-45

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.

Il s'agit d'un composant réseau passif, car il sert uniquement au passage du courant entre les quatre paires torsadées de câbles torsadés de catégorie 5 et les broches du connecteur RJ-45.

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.



Prise RJ-45 et connecteur RJ-45

Voici un tableau récapitulant les différents types de câbles ainsi que leur débit :

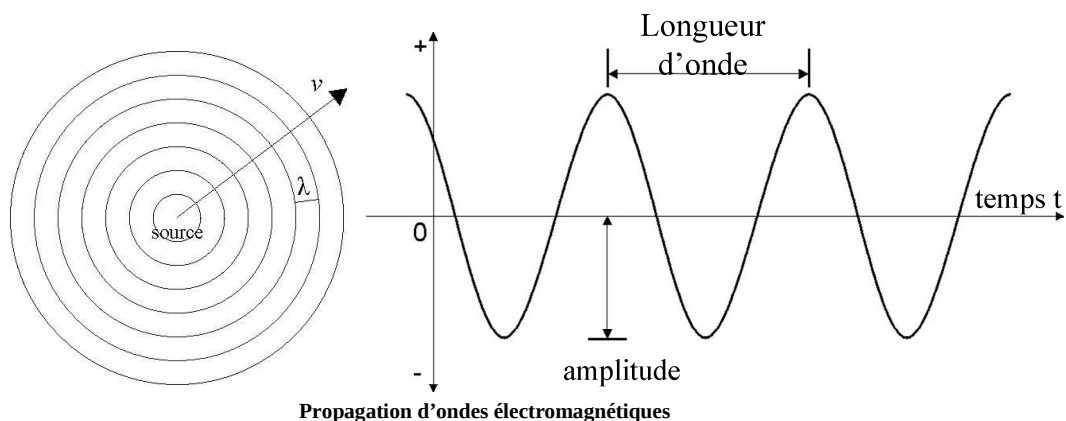
Technologie	Type de câble	Débit théorique	Longueur Max	Connecteur	Coût
10 Base 2 (Thinnet)	Coaxial	10 Mbits/s	200 m	BNC	Peu cher
10 Base 5 (Thicknet)	Coaxial	100 Mbits/s	500 m	BNC	Peu cher
10 Base T	UTP cat 5	10 Mbits/s	100 m	RJ45	Bon marché
100 Base TX	UTP cat 5	100 Mbits/s	100 m	RJ45	Bon marché
10 Base FL	Fibre optique	10 Mbits/s	2000 m	SC	Elevé
100 Base FX	Fibre optique	100 Mbits/s	400 m	SC	Elevé

3.3. Médias optiques

3.3.1. Phénomènes physiques :

Spectre électromagnétique

Les ondes radio, l'infrarouge, les rayons lumineux visibles, ainsi que les rayons gamma et X sont tous des types d'énergie électromagnétique. Cette énergie est créée lorsqu'une source change répétitivement en intensité. Les émissions amplifiées et diminuées créent des ondes, des vibrations qui se déplacent comme des vagues créées par un caillou jeté dans l'eau.



La distance entre les ondes est appelée la longueur d'onde et est désignée par λ . Elle dépend de la fréquence d'altérations de charge. Plus la fréquence d'émission est grande, plus petite est la distance entre les sommums (maximums) d'ondes.

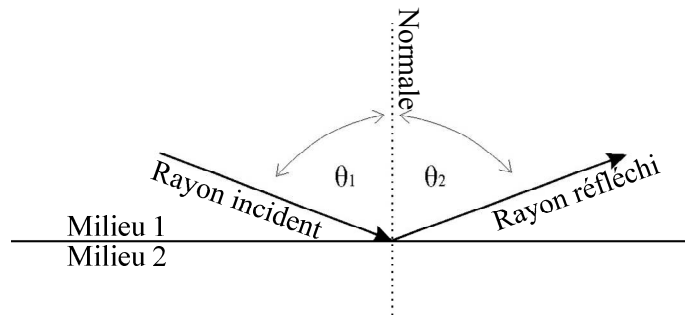
Les ondes électromagnétiques partagent des propriétés similaires. Entre autres, elles se propagent toutes à la vitesse de la lumière c (299 792 458 m/s) quand elles traversent le vide. Quant à un autre environnement, tel que l'air, l'eau ou le verre, leur vitesse v est atténuée.

Lorsqu'on regroupe les ondes électromagnétiques commençant par celles qui ont la plus petite longueur jusqu'aux ondes qui ont la plus grande longueur, on obtient le **spectre électromagnétique**. Les ondes de longueur entre 400 nm et 700 nm constituent la lumière visible. La lumière d'une longueur d'onde supérieure est appelée la lumière infrarouge. Les longueurs couramment utilisées pour le transport d'informations dans la fibre optique sont précisément les longueurs de l'infrarouge : 850 nm, 1310 nm et 1550 nm.

Réflexion

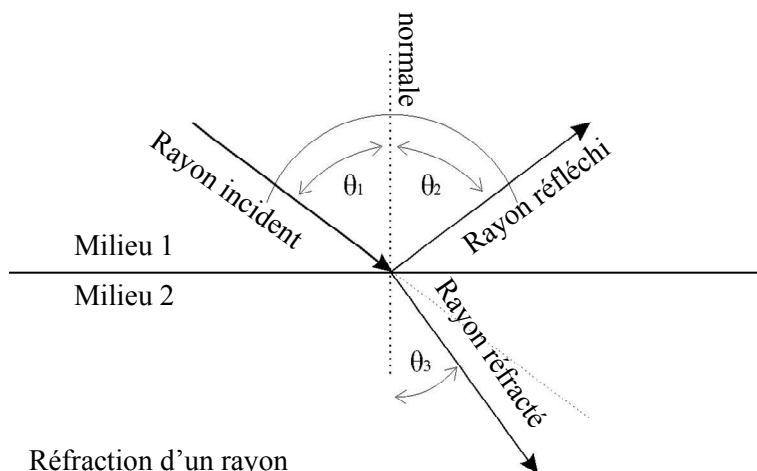
Un rayon passant dans un milieu 1, qui rencontre sur son chemin un autre milieu 2 est appelé **rayon incident**. Une fois arrivé sur la surface de l'autre milieu, le rayon incident se réfléchit. Selon la **loi de réflexion**, l'angle incident θ_1 est égal à l'angle réfléchi θ_2 .

Réflexion d'un rayon où $\theta_1 = \theta_2$



Réfraction

Supposons qu'un rayon incident traverse un milieu transparent, par exemple l'air, et arrive sur la surface d'un autre milieu, également transparent, soit l'eau. Au lieu de se réfléchir complètement, il est possible que le rayon incident traverse la surface qui sépare les deux milieux (le dioptre), ainsi en pénétrant dans l'eau. Lorsque le rayon traverse la surface, son angle s'approche vers la normale. On peut observer ce cas sur le schéma ci-dessous où l'angle θ_1 est supérieur à θ_3 . Ce phénomène est appelé la **réfraction** et l'on dit pour le rayon traversé qu'il est **réfracté**.



Réfraction d'un rayon

Pour qu'un rayon soit réfléchi sans être réfracté, il faut que son angle d'incidence soit plus grand que **l'angle critique** des deux milieux.

Il est important de connaître le facteur qui détermine l'importance de déviation subi par le rayon réfracté. Ce coefficient, nommé **l'indice de réfraction**, est le rapport entre la vitesse de la lumière dans le vide et dans le milieu :

$$n = c / v.$$

Il faut également retenir que l'indice de réfraction dépend de la longueur d'onde λ . Cela veut dire que deux rayons ayant deux différentes longueurs d'ondes ne se comportent pas de la même façon dans un milieu M , à savoir que l'une se déplace plus vite que l'autre. C'est d'ailleurs pour cette raison que l'on a choisi la lumière infrarouge et non pas une autre pour le transport d'informations dans la fibre optique.

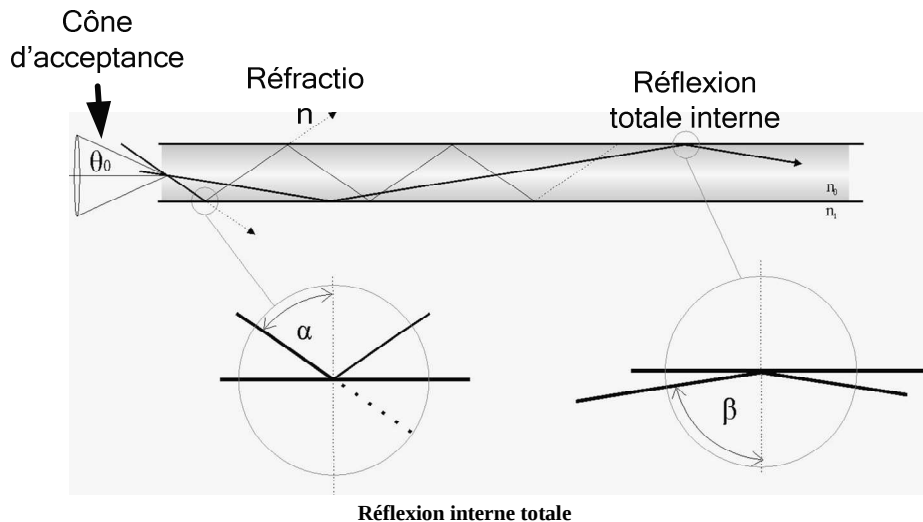
La réflexion interne totale

Dans une fibre optique, les données sont transmises de façon similaire à une transmission réalisée dans un fil électrique : s'il y a de la lumière, l'information traduite en bit 1, sinon en bit 0. L'objectif est évidemment que le rayon, le porteur de l'information, arrive bien de la source jusqu'à destination et

sans être affaibli. Pour ce faire, le rayon doit être guidé dans la fibre sans réfraction, il doit se propager en faisant la **réflexion interne totale**.

Les deux conditions principales pour réaliser la réflexion interne totale sont :

- l'indice de réfraction n_0 du cœur de la fibre doit être **supérieur** à l'indice de réfraction de la gaine n_1 ,
- le rayon entrant doit se situer dans le **cône d'acceptance**.



Sur l'image au dessus l'on voit que le premier rayon entrant est en dehors du cône, avec un angle supérieur à θ_0 . Remarquez sur la première partie agrandie que le rayon est effectivement réfracté et rappelez-vous que dans ce cas, l'angle d'incidence α est bien inférieur à l'angle critique.

Le deuxième rayon, quant à lui, passe bien par le cône, son angle d'incidence β est supérieur à l'angle critique, et il se propage par la réflexion totale interne tout au long de la fibre. C'est un **rayon guidé**.

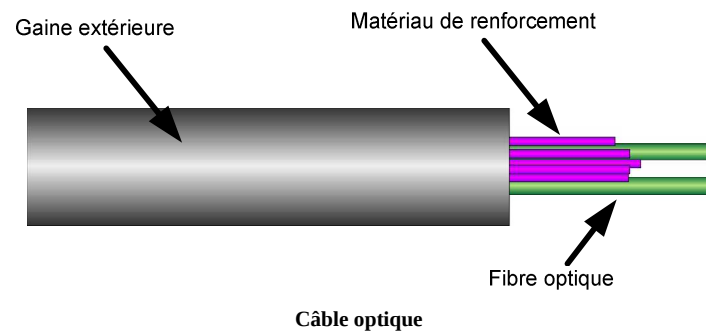
3.3.2. Composants optiques

Fibre optique

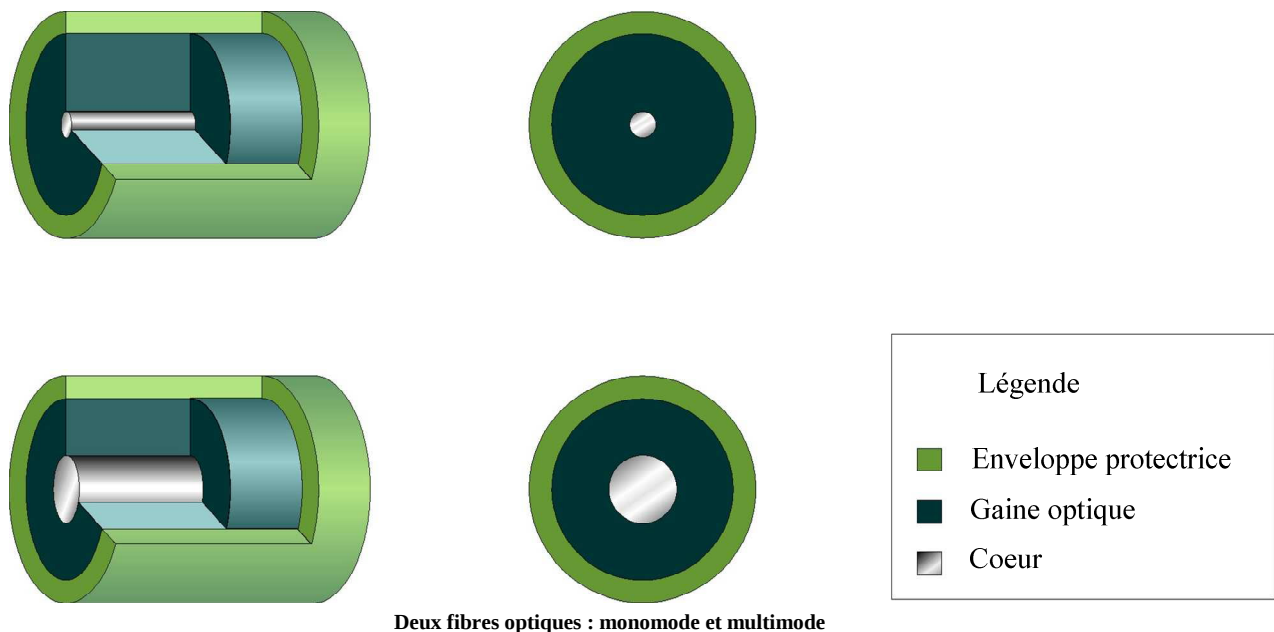
Une fibre optique transmet des données dans un sens seulement. Aussi pour que deux entités communiquent en full duplex, un câble optique doit contenir deux fibres optiques : l'une pour transmission et l'autre pour réception. Un câble peut contenir de 2 jusqu'à 48 fibres. Les fibres réunies ensemble dans un câble ne créent pas de bruit, car elles ne portent pas d'impulsions électriques qui pourraient induire des interférences électromagnétiques. Donc elles n'ont pas besoin d'une protection par blindage, comme les fils en cuivre.



Un câble à fibres optiques est soutenu avec des fils de renforcement en plastique, tel que le Kevlar. Ceci rend un câble plus résistant, assurant ainsi que les fibres optiques ne s'abîment pas lorsqu'elles sont pliées.

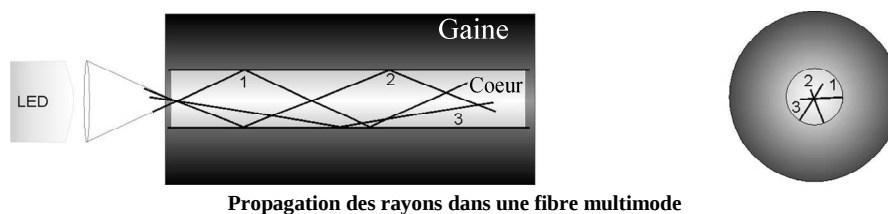


La lumière est guidée dans le centre de la fibre, appelé **cœur**. Le cœur est constitué en majorité de silicium dioxyde (silice), enrichi avec d'autres éléments. Il est entouré par la gaine optique. La gaine est également faite de silice, mais son indice de réfraction est bien inférieur à celui du cœur. Cela permet justement à la lumière de se réfléchir. La gaine optique est protégée par une enveloppe, fabriquée fréquemment en plastique.



Deux fibres optiques : monomode et multimode

Le chemin fait par un rayon est aussi appelé un **mode**. Lorsqu'une fibre optique transmet un seul rayon, elle est appelée fibre **monomode**. La fibre qui transmet plusieurs rayons, elle est appelée fibre **multimode**. Pour transmettre plusieurs rayons, avec des chemins différents, le cœur de la fibre multimode doit être plus grand que celui de la fibre monomode.

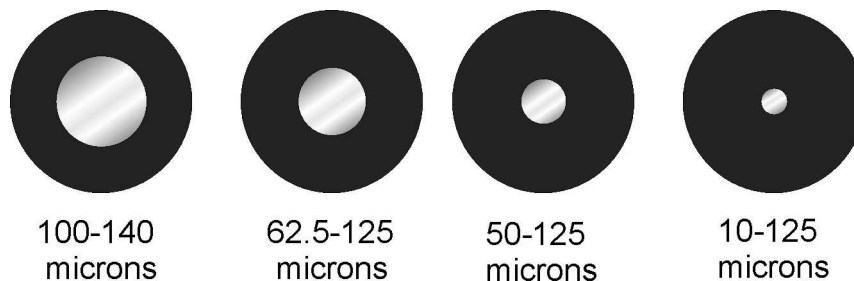


Propagation des rayons dans une fibre multimode

Les sources qui diffusent la lumière dans la fibre ne sont pas les mêmes pour les fibres monomode et multimode. En effet, une fibre multimode utilise la LED (Light Emitting Diode), en français « DEL », Diode Electroluminescente, alors qu'une fibre monomode utilise le laser, qui est en général plus cher. Un laser émet des rayons de longueur plus longue que celle des rayons émis par une LED. De ce fait, la longueur maximale de la fibre multimode est 2000 m. Tandis que la longueur maximale de la fibre monomode est 3000 m. Les fibres monomode sont plus coûteuses et leur utilisation est fréquemment destinée aux liaisons WAN, entre différents bâtiments. Les fibres multimode sont moins chères et plus utilisées dans l'entreprise.



Les diamètres des fibres ont des tailles différentes. Sur le schéma ci-dessous, on peut voir les types multimode et monomode alignés, montrant les diamètres différents en tailles relatives.



La plupart des équipements pour les réseaux locaux transmettent des données en forme électrique. Afin d'intégrer la fibre optique dans un tel réseau, les signaux électriques doivent être transformés en impulsions lumineuses. Pour se faire, il existe des transmetteurs qui transforment, codent et envoient les signaux de lumière. Comme déjà énoncé, il y a deux types de source de lumière :

- **DEL** : diode électroluminescente produit de la lumière infrarouge de longueur de 850 nm, ou 1310 nm.
- **LASER** : (en anglais : Light Amplification by Stimulated Emission Radiation) Amplification de lumière par l'émission de radiation stimulée produit des rayons étroits de lumière infrarouge d'une grande intensité et de longueur d'onde de 1310 nm ou 1550 nm.

A l'autre bout de la fibre se trouve le récepteur. Il transforme les impulsions lumineuses en impulsions électriques qui sont ensuite transférées aux autres équipements réseaux.

Les extrémités de fibre sont attachées aux connecteurs qui se branchent dans les prises des transmetteurs et récepteurs. Les connecteurs de type **SC** (Subscriber Connector) sont le plus souvent utilisés pour les fibres multimode et les connecteurs de type **ST** (Straight Tip) les plus fréquemment utilisés pour les fibres monomode. Le schéma ci-dessous montre les connecteurs ST et SC, respectivement.



Les deux connecteurs de fibre optique : ST et SC (simplex)

Une paire de connecteurs joints dans un emboîtement s'appelle un connecteur duplex. Un connecteur simplex est un connecteur simple, reliant une fibre seulement.

Les câbles optiques qui dépassent leur longueur maximale sont prolongés par des répéteurs, des équipements d'amplification de signaux de lumière.

Signaux et bruit dans les fibres optiques

Malgré le fait que la fibre optique est le meilleur média de transmission, les signaux qui y transitent peuvent être atténués par différents facteurs. Le plus important facteur est la diminution du signal causée par la dispersion. Elle arrive lorsque la fibre est trop pliée ou serrée. L'angle incident d'un rayon peut alors devenir inférieur à l'angle critique faisant ainsi qu'une partie du rayon soit réfractée. L'absorption est une autre forme d'atténuation. Elle arrive lorsqu'un rayon rencontre des impuretés sur son chemin.

Pour contrer les problèmes d'atténuations, on teste les liaisons en fibre optique avec des outils qui mesurent la perte d'énergie et les temps de voyage des signaux.

3.4. Médias sans fil

3.4.1. Fonctionnement d'un réseau sans fil

Les réseaux sans fils ou WLAN (pour Wireless WAN), réussissent à conjuguer tous les avantages d'un réseau filaire traditionnel comme Ethernet mais sans la limitation des câbles.

La mobilité est maintenant l'attrait principal pour les entreprises, la possibilité d'étendre son réseau LAN existant selon les besoins de l'organisation.

Un WLAN a également besoin, tout comme un LAN, d'un média. Au lieu de câbles à paires torsadées, les WLANs utilisent des fréquences radio à 2,4 GHz et 5 GHz.

On parle de "réseaux sans fils" mais la plupart du temps, ces réseaux sont intégrés aux LANs traditionnels, juste considérés comme une extension à l'existant. Aujourd'hui, grâce à des normalisations de l'IEEE et du "Wi-Fi Alliance", les équipements sans fils sont standardisés et compatibles, ce qui explique l'engouement croissant pour ce type de réseau de moins en moins coûteux.

Il faut savoir que la première version d'un réseau sans fil offrait un débit de l'ordre de 1 à 2 Mbps. Grâce à la mobilité rendue possible, cette technologie fut rapidement mise en place.

En effet, tout d'abord pour faciliter certains métiers comme la gestion des stocks dans les entrepôts, rapidement les réseaux sans fils se sont étendus à d'autres secteurs comme dans les hôpitaux, les écoles et universités. Standardiser cette technologie devenait nécessaire, un groupe de travail a donc été mis en place en 1991 par plusieurs constructeurs, le WECA (Wireless Ethernet Compatibility Alliance), plus tard, ce nom changera pour le Wi-Fi (Wireless Fidelity).

En Juin 1997, L'IEEE publie les standards 802.11 pour les réseaux locaux sans fils.

Les réseaux sans fils peuvent fonctionner à deux bandes de fréquences, selon la technologie utilisée. Soit aux alentours de 2400 Mhz (2,4 Ghz) pour le 802.11b et 802.11g soit aux alentours de 5000 Mhz pour le 802.11a.

La bande la plus utilisée pour le moment est l'ISM (Industrial Scientific and Medical) cela correspond à la bande des 2,4 GHz avec une largeur de bande de 83,5 MHz. Soit des fréquences allant de 2,4 GHz à 2,4835 GHz.

Tableau récapitulatif des fréquences et débits :

	802.11b	802.11a	802.11g
Bande de fréquence	2,4 Ghz	5 Ghz	2,4 Ghz
Débit maximum	11 Mbps	54 Mbps	54 Mbps

Les lois de la radio :

- Débit plus grand = Couverture plus faible
- Puissance d'émission élevée = Couverture plus grande mais durée de vie des batteries plus faible
- Fréquences radio élevées = Meilleur débit, couverture plus faible

Pour qu'un réseau sans fil fonctionne, il faut au moins 2 périphériques au minimum, comme un point d'accès (AP) et une carte sans fil pour le client. Voici les différents composants que l'on peut trouver dans un WLAN :

- **Les adaptateurs du client :**
 - o PCMCIA : Utilisé sur les ordinateurs portables en externe, antenne intégrée
 - o LM : Identique au PCMCIA, même bus, mais sans antenne
 - o PCI : Utilisé pour les ordinateurs fixes
 - o Mini PCI : Utilisé sur les ordinateurs portables en interne, nécessite une antenne supplémentaire
- **Les points d'accès (AP) :** Les modèles Cisco Aironet 1100 et 1200 sont les plus utilisés pour un accès aux utilisateurs
- **Les ponts, ou Wireless bridges (BR) :** Périphérique principalement utilisé pour relier deux réseaux filaires
- **Les antennes :**
 - o Directionnelles
 - o Omnidirectionnelles
- **Les périphériques sans fil natifs :**
 - o PDA
 - o Ordinateur portable
 - o Téléphones IP
 - o Imprimantes

3.4.2. Authentification et sécurité

Avec la venue du 802.11 et des réseaux sans fil, le problème de la sécurité s'est bien évidemment posé. Bien évidemment la propagation des ondes fut le premier souci, la solution matérielle des antennes directionnelles ainsi que la pose de filtres sur les vitres de manière à ne pas laisser passer les ondes fut une des solutions, mais trop onéreuse pour beaucoup d'entreprises. Plusieurs solutions logicielles ont donc vu le jour.

La première repose sur l'utilisation d'un SSID (Service Set Identifier) qui permet de se connecter au réseau si l'on connaît le SSID. Cette solution est tout de même peu sécurisée du fait qu'un logiciel permettant de capturer des trames peut facilement récupérer ce SSID.

Une autre sécurisation peut agir sur l'adresse MAC de la carte directement. Cette méthode est tout de même un peu plus sécurisée puisque se basant sur les adresses MAC enregistrées comme ayant accès au réseau. Néanmoins cette méthode reste statique est chaque nouvel utilisateur doit être validé dans la base d'adresses MAC. Pour les grandes entreprises cela représenterait une charge importante de travail. Cette solution est à réserver pour de petits réseaux (PME ou LAN).

Une troisième solution consiste en une clé de chiffrement qui crypte les transferts. Cette clé est nécessaire pour se connecter à l'AP et pour maintenir la connexion. On parle de clé WEP (Wired Equivalent Privacy). Le cryptage se fait sur 64 ou 128 bits. La norme WPA (Wi-Fi Protected Access) met en place un système de clé dynamique.

Il est bien évident que le jumelage de ces différentes solutions peut augmenter la sécurité de votre réseau, mais cela reste encore inférieur à un réseau filaire. Il faut donc attendre la spécification 802.11i ou l'application de la norme WPA 2 pour en théorie enfin avoir un niveau sécurité acceptable pour un réseau de grande envergure.

3.4.3. Modes d'implémentations

Considérons deux stations équipées chacune d'une carte Wi-Fi.

Nous avons deux possibilités de connecter ces stations entre elles :

- Soit en les connectant directement l'une à l'autre (comme on pourrait le faire avec un câble croisé et deux cartes réseau Ethernet)
- Soit en passant d'abord par une borne (comme on pourrait le faire avec un concentrateur et une paire de câbles Ethernet droits).

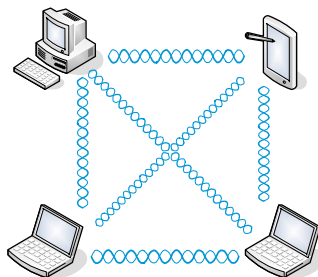
Dans le cas du Wi-Fi, ce n'est pas le média qu'il faut modifier afin de choisir la méthode de connexion, mais la configuration de la carte.

En effet, une carte Wi-Fi ne se configure pas de la même façon suivant que l'on veuille établir une connexion en mode Ad-Hoc (connexion directe d'une station à l'autre) ou en mode Infrastructure (en utilisant une borne).

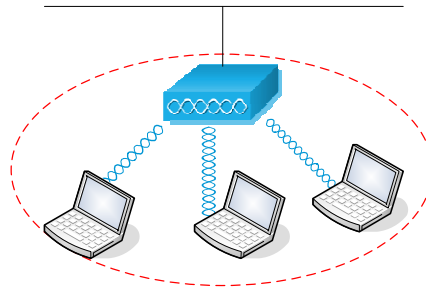
Le mode Ad-Hoc apporte l'avantage de la mobilité. En effet, on peut mettre en réseau deux stations mobiles tant que chacune d'elles se situe dans la zone de couverture de l'autre, on peut donc facilement se déplacer tout en conservant la connectivité par exemple dans une salle de réunion.

Le mode infrastructure, quant à lui, permet de connecter un réseau Wi-Fi à un réseau filaire (internet, ou d'entreprise par exemple). Cependant la mobilité d'une telle configuration est limitée à la zone de couverture de la/ les borne(s) reliée(s) au réseau filaire.

Nota : contrairement à l'Ethernet, il est possible de connecter plusieurs stations entre elles en mode Ad-Hoc, cependant, il arrive fréquemment que l'on perde la porteur, ce qui rend le service instable. Pour des raisons de performances et de qualité de connexion, il est déconseillé de connecter plus de 4 stations en mode Ad-Hoc :



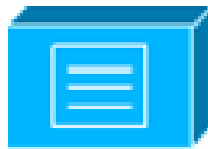
- Infrastructure : connexion en passant par une borne (équivalent au concentrateur Ethernet).



3.5. Equipements de couche 1

3.5.1. Répéteur

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



Symbole d'un répéteur

3.5.2. Concentrateur

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, resynchronisé et ré émis au travers de tous les autres ports.



Symbole d'un Concentrateur 10 Base T



Symbole d'un concentrateur 100 base T

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.

3.5.3. Emetteur/récepteur

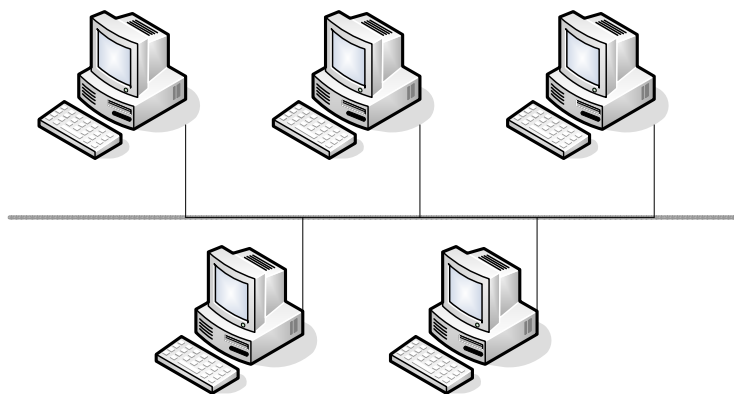
Un émetteur-récepteur (transceiver) convertit un signal en un autre. Il est souvent intégré aux cartes réseau.

3.6. 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.

3.6.1. La topologie en bus

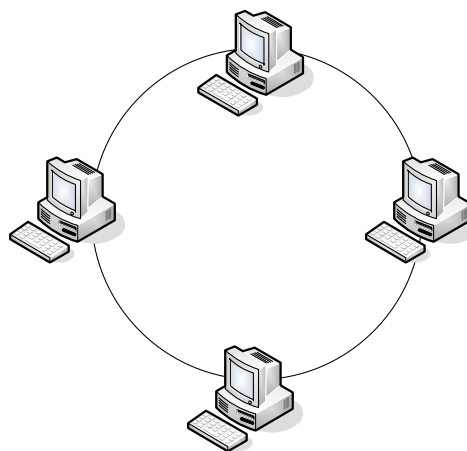
- **Perspective physique** : Tous les hôtes sont connectés directement à une liaison
- **Perspective logique** : Tous les hôtes voient tous les signaux provenant de tous les autres équipements



Topologie en bus

3.6.2. La topologie en anneau

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

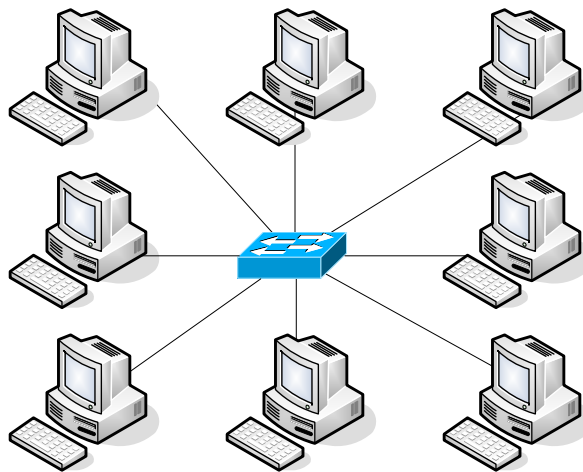


Topologie en anneau

Une variante de cette topologie est le double anneau ou 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.

3.6.3. La topologie en étoile

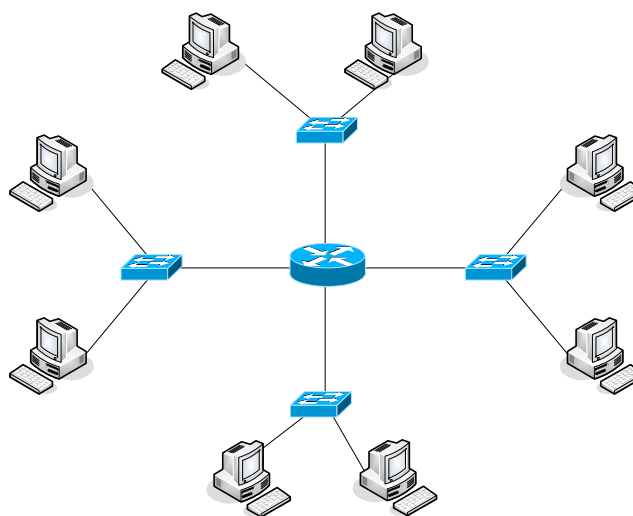
- **Perspective physique :** Cette topologie comporte un nœud central d'où partent toutes les liaisons avec les autres nœuds.
- **Perspective logique :** Toutes les informations passent par un seul équipement, par exemple un concentrateur



Topologie en étoile

3.6.4. La topologie en étoile étendue

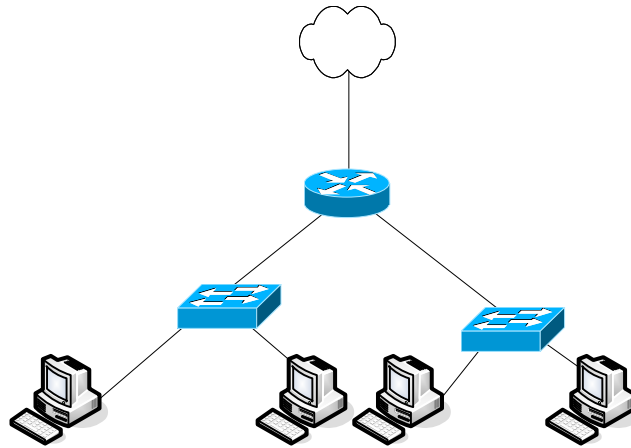
Cette topologie est identique à la topologie en étoile si ce n'est que chaque nœud connecté au nœud central est également le centre d'une autre étoile.



Topologie en étoile étendue

3.6.5. La topologie hiérarchique

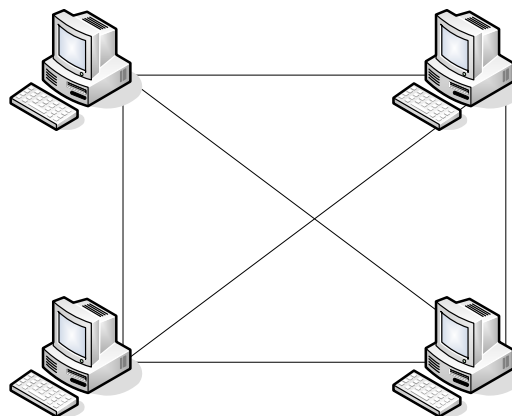
- **Perspective physique** : Cette topologie ressemble à une topologie en étoile sauf qu'elle n'utilise pas de nœud central. Elle utilise un nœud de jonction à partir duquel elle se branche vers d'autres nœuds.
- **Perspective logique** : Le flux d'informations est hiérarchique



Topologie hiérarchique

3.6.6. La topologie complète (maillée)

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



Topologie complète

4. Couche 2 : Technologies Ethernet

4.1. Introduction aux technologies LAN

Un LAN (Local Area Network) est un réseau local, il a donc une taille géographiquement limitée (quelques milliers de mètres maximum).

Un LAN permet un accès multiple aux médias à large bande tout en assurant une connectivité continue aux services locaux (ressources et accès Internet partagés, messagerie, etc.). Son but est de relier physiquement des terminaux réseaux proches (stations de travail, serveurs, imprimantes, etc.) par une liaison physique.

Ils sont caractérisés par un haut débit et un faible pourcentage d'erreurs dues à l'atténuation. Ils relient les différents périphériques, terminaux et stations de travail entre eux.

4.2. Introduction à Ethernet

Ethernet est la technologie de base des réseaux LAN la plus utilisée actuellement. Le principe repose sur le fait que toutes les machines sont reliées à une même ligne de communication. L'institut IEEE l'a normalisé et adapté dans son modèle IEEE 802.3. Ces deux technologies sont très similaires (elles diffèrent sur un champ de trame seulement).

4.2.1. Ethernet et le modèle OSI

La technologie Ethernet opère au niveau de la couche physique et de la couche liaison de données (la couche MAC seulement).

Lorsque plusieurs terminaux communiquent par le biais d'un média partagé, les données passent le plus souvent par un répéteur (accessoirement multi ports). Toutes les stations connectées à ce même média « voient » donc ce trafic. Elles communiquent entre elles également par ce même média. Des collisions se créent alors, car elles utilisent ce média en concurrence. On peut donc assimiler un domaine de collision à un environnement partagé.

4.2.2. Spécifications et normes

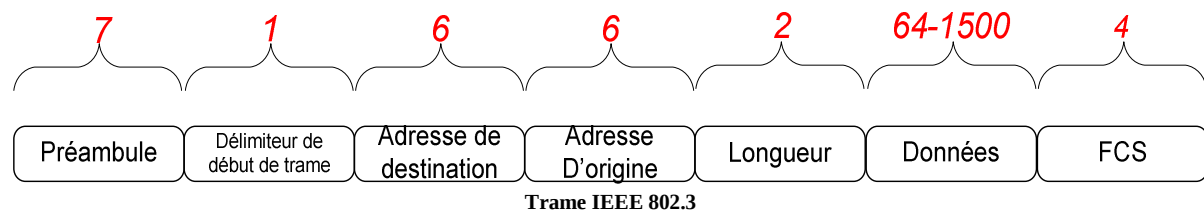
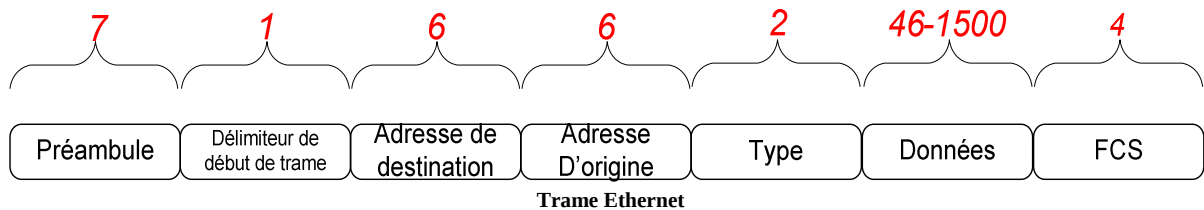
Chaque désignation de technologie utilise une normalisation qui permet d'identifier ses caractéristiques. Celles-ci sont de la forme : vitesse en Mbps – type de signal – type de câble. (ex : 100 Base TX)

- Deux types de signalisation existent : Baseband (transmission numérique) ou Broadband (utilisation de porteuse : transmission par ondes par exemple).
- Le type de câble utilisé : cuivre à paires torsadées non blindé (Unshielded Twisted Pairs), ou de type fibre optique (Fiber).
- On exprime aussi sa capacité à supporter le Full Duplex par un X. (à l'exception du 10 Base T qui supporte tout de même le mode Full Duplex).

L'IEEE a défini des normes pour les différentes technologies Ethernet :

Norme	Appellation	Débit	Média utilisé
802.3	Ethernet	10 Mbps	Coaxial / UTP / fibre optique
802.3u	Fast Ethernet	100 Mbps	UTP / Fibre optique
802.3z	Gigabit Ethernet	1000 Mbps	Fibre optique
802.3ab	Gigabit Ethernet	1000 Mbps	Câble UTP
802.3ae	10 Gigabit Ethernet	10 000 Mbps	Fibre Optique

4.2.3. 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 à 1 consécutifs, servant à synchroniser les portions de réception des trames de toutes les stations.
- **Champ d'adresse de destination** : peut être de type unicast, multicast ou broadcast.
- **Champ d'adresse d'origine** : toujours de type unicast.
- **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.
 - C'est sur cette partie que diffèrent les trames 802.3 et Ethernet : la valeur du champ permet de déterminer le type de trame : 802.3 ou Ethernet.
 - La trame est de type 802.3 si la valeur hexadécimale du champ est strictement inférieure à 0X600 ; La trame est de type Ethernet si la valeur hexadécimale du champ est égale à 0X600.
- **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 ici aussi avoir recours au remplissage.
- **FCS** : Séquence de contrôle de trame. Cette séquence contient un code de redondance cyclique permettant à l'unité réceptrice de vérifier l'intégrité des données transmises.

4.3. Fonctionnement d'Ethernet

4.3.1. MAC

Le principe utilisé pour partager l'accès à des ressources communes est appelé MAC pour Media Access Control (à ne pas confondre avec l'adresse MAC).

Dans un environnement où plusieurs hôtes se partagent un média unique de communication, un problème de priorité doit être résolu. Le problème est le même que dans une situation courante : lors d'une discussion à l'intérieur d'un groupe de personnes, une seule personne parle à la fois si elle veut être comprise par son ou ses interlocuteurs.

Dans un environnement Ethernet, c'est au niveau de la sous-couche MAC que l'on va utiliser un processus de détection des collisions : plusieurs hôtes émettent en même temps sur le même média. Ethernet et 802.3 utilisent un principe d'accès au média non déterministe : CSMA/CD (Carrier Sense Multiple Access with Collision Detection)

Les hôtes se partagent donc 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.

4.3.2. Erreurs possibles

Pendant une transmission de données, de nombreux facteurs peuvent entraîner une corruption de celle-ci.

Le but est de détecter ces erreurs correctement pour déterminer quelles trames doivent être retransmises afin de récupérer des données intègres.

Collisions

Dans un environnement partagé, la première corruption rencontrée est de type collision. Lorsque deux hôtes ou plus émettent un signal au même instant sur le média, il se produit un survolage qui ne signifie plus rien en terme de données. Ces collisions ne se produisent que dans un environnement Half-Duplex. (car dans un environnement Full-Duplex, chaque paire torsadée n'est utilisée qu'entre deux hôtes dans un seul sens de transmission.). L'algorithme CSMA/CD permet de détecter ces collisions et de les éviter.

Il existe trois types de collision :

- Collision locale
- Collision distante
- Collision de retard

La collision locale est de type survolage, comme vu dans l'exemple précédent.

Une collision distante résulte d'une trame ayant une longueur inférieure au minimum ou d'un FCS incorrect. Elle est souvent rencontrée à une certaine distance d'environnement répété (hub ou répéteur) mais n'a pas de problème de survoltage. Il peut s'agir de fragments de collision non détruits par un équipement de type répéteur par exemple.

Une collision de retard n'est pas détectée par la couche liaison de données. En effet, elle est caractérisée par une erreur dans les données à partir du 64^{ème} octet. Contrairement aux deux autres types de collision, une collision de retard ne déclenche pas une réémission directe de la trame (car elle n'a pas été détectée par la couche de liaison). La station réceptrice analyse d'abord cette trame avec une couche supérieure (qui détecte l'erreur dans la trame) puis demande un renvoi de cette trame.

Trames longues

Ce type d'erreur est un simple dépassement de la taille maximale d'une trame.

La taille du champ « Données » (variable) d'une trame ne doit pas excéder 1500 octets. Une trame a donc une taille maximale de 1526 octets. Une trame de taille supérieure est donc considérée comme fausse.

Trames courtes

Comme pour les trames longues, l'erreur se situe au niveau du champ « données » qui doit avoir une taille minimale de 46 octets (ou 64 pour IEEE 802.3). Les trames courtes se caractérisent donc par une taille inférieure à 72 octets (ou 90 octets pour IEEE 802.3) mais avec un FCS valide : sinon elle serait considérée comme un fragment de trame, détruit lui aussi.

Autres types d'erreur

D'autres erreurs peuvent survenir du fait de la mauvaise qualité du média (ou d'interférences extérieures) :

- FCS incorrect : le résultat du FCS est faux quant aux données transmises
- le champ longueur ne concorde pas avec la taille du champ « données »
- longueur de champ incorrecte : le préambule ne fait pas 7 octets, ...

Une fois qu'une erreur de ce type est détectée, la couche supérieure (de la station réceptrice) va demander un renvoi de cette trame à la station émettrice, jusqu'à obtenir une trame valide.

5. Couche 2 : Commutation Ethernet

5.1. Domaine de collision

On appelle domaine de collision la partie d'un réseau comprenant un environnement partagé. C'est dans ce domaine que les hôtes vont accéder en concurrence à une ressource. De ce fait, des collisions vont se créer sur cette partie du réseau. Le domaine de collision s'étend sur la plus grande partie du réseau contenant des équipements de couche 1 interconnectés.

5.2. Segmentation

Les domaines de collision posent des problèmes, proportionnellement à leur taille. En effet, plus un domaine de collision est grand (mesuré en nombre d'hôtes), plus la bande passante par hôte est faible, et plus le nombre d'erreurs est grand.

Pour diminuer ces effets néfastes, il suffit de segmenter un domaine en plusieurs, de tailles inférieures. On aura alors moins de collisions par segment, donc une plus grande fiabilité et une meilleure bande passante.

Le principe de la segmentation est de n'envoyer des données que sur la portion de réseau concernée. On va ainsi réduire le trafic inutile, ainsi que le nombre d'utilisateurs concurrents du même média. Pour la segmentation, des équipements de couche 2 sont nécessaires. C'est à ce niveau que l'on peut prendre des décisions d'adressage (sur quel média transmettre une trame).

5.2.1. Segmentation par ponts

Les ponts permettent de segmenter un réseau en n'envoyant les données que sur la partie du réseau concernée. Après avoir appris sur quelle portion se trouvent les hôtes (par leur adresse mac), un pont filtrera le trafic suivant l'adresse de destination. Il laissera donc transiter les données vers la partie du réseau qui contient l'adresse de destination, et bloquera les paquets qui ne sont pas destinés à cette même partie.

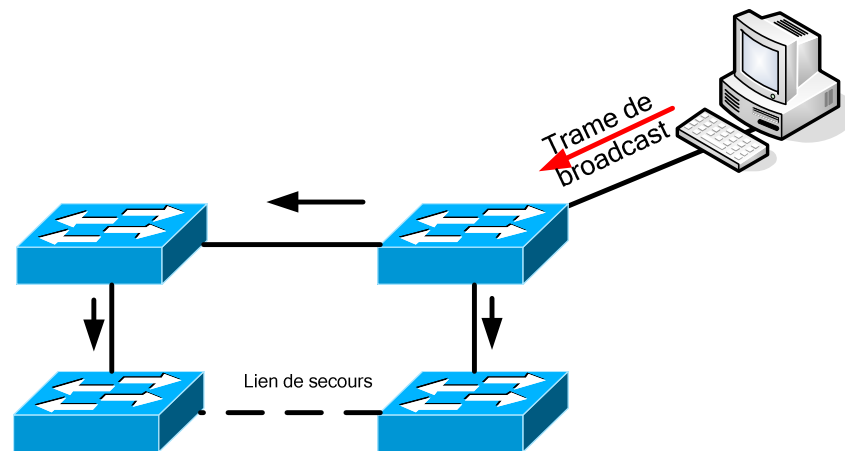
5.2.2. Segmentation par commutateurs

Les commutateurs sont l'équivalent de répéteurs multi ports intelligents. Chaque hôte ou groupe d'hôtes connecté à un port du commutateur veut envoyer des données. Au lieu de retransmettre ces données sur chaque port, le commutateur ne va renvoyer que sur le port où se trouve la partie du réseau contenant le(s) destinataire(s).

Pour se faire, le commutateur va apprendre les adresses MAC de chaque hôte connecté à ses ports. Il saura ainsi quels hôtes se trouvent sur chacun de ses ports. Il stocke ces données dans une table d'adresses MAC.

Les commutateurs fonctionnent beaucoup plus vite que les ponts et créent des domaines sans collisions entre 2 ports en interne (par l'utilisation de circuits virtuels).

5.2.3. Spanning Tree



Dans un réseau utilisant de nombreux commutateurs, des chemins redondants sont souvent utilisés afin d'établir une connectivité fiable et tolérante aux pannes. Un problème se pose alors, car du fait de ces chemins redondants, des boucles de commutation peuvent apparaître. Des tempêtes de broadcast peuvent alors se produire, entraînant une congestion du réseau.

Le protocole Spanning Tree a été développé dans le but de contrer ce problème de boucles de commutation.

Chaque commutateur utilisant le protocole Spanning Tree, envoie des datagrammes BPDU (Bridge Protocol Data Units) à ses compères pour indiquer sa présence. Chaque commutateur calcule alors les routes optimales suivant la topologie et élimine les chemins redondants inutiles grâce à l'algorithme STA (Spanning Tree Algorithm).

Lors de l'utilisation de Spanning Tree, un port de commutateur peut prendre 5 états différents :

- Blocage : aucune trame acheminée, unités BPDU entendues
- Ecoute : aucune trame acheminée, écoute des trames
- Apprentissage : aucune trame acheminée, apprentissage des adresses
- Acheminement : trames acheminées, apprentissage d'adresses
- Désactivation : Aucune trame acheminée, aucune unité BPDU entendue

Le protocole Spanning Tree permet donc de créer un réseau sans liaisons redondantes sans les éliminer. Ces chemins sont alors utilisables en cas de nécessité : si une liaison n'est plus disponible, l'algorithme Spanning Tree recalcule un arbre de chemins permettant de remplacer la liaison manquante.

6. Couche 3 : Protocole IP

6.1. Protocoles routables

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

Un protocole routable définit la notion d'adressage hiérarchique : un hôte est défini par une adresse unique sur un segment de réseau unique.

Un protocole de routage (à ne pas confondre avec protocole routable), grâce à la structure du protocole routé, a toutes les informations nécessaires pour envoyer un paquet sur le segment spécifié à l'hôte spécifié.

6.1.1. Protocoles orientés connexion et non orientés connexion

Un protocole non orienté connexion ne définit pas de chemin unique pour acheminer les paquets d'un hôte source vers un hôte de destination. Les paquets peuvent alors emprunter des chemins différents suivant la topologie réseau existante entre ces deux hôtes. Cela implique une durée de trajet différente pour chaque paquet et donc un ordre d'arrivée différent de celui d'émission. L'hôte de destination ne peut pas réordonner les paquets.

Le protocole IP est un protocole non orienté connexion.

Un protocole orienté connexion définit un chemin unique entre l'hôte source et l'hôte de destination. Les paquets empruntent alors le même chemin et arrivent donc dans le même ordre. Pour ce faire, l'hôte source établit en premier lieu une connexion avec l'hôte de destination. Une fois cette connexion établie, chaque paquet est envoyé par ce seul chemin. On appelle ce processus « commutation de circuits ».

Le protocole TCP est un protocole orienté connexion.

6.1.2. Protocoles routés

Protocole routé : c'est un protocole de communication de couche 3. Il définit le format des paquets, et notamment la manière de désigner le destinataire du paquet. Un protocole routé peut être **routable** ou **non routable**.

- **Routable :** les messages envoyés à l'aide de ce protocole peuvent sortir de leur réseau (via un routeur). En effet, le format du paquet comprend une distinction entre la partie hôte et la partie réseau.
- **Non routable :** les messages envoyés à l'aide de ce protocole ne peuvent pas sortir de leur réseau. En effet, le format du paquet ne comprend pas de mécanisme permettant à un élément réseau de faire suivre ces paquets au travers de différents réseaux.

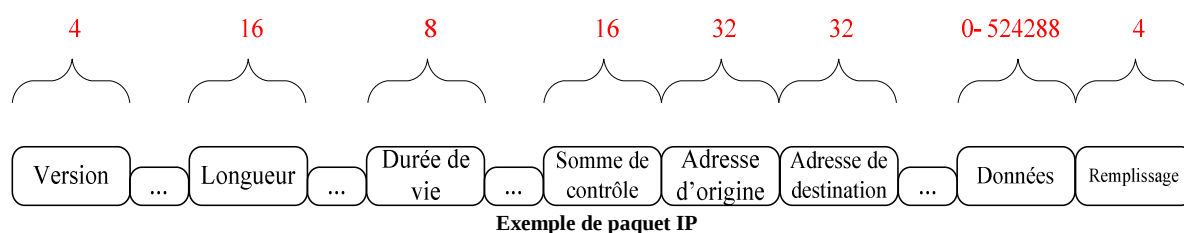
La liste des protocoles routés suivante présente les protocoles les plus connus :

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

6.2. Protocole IP

6.2.1. Paquet IP

Les informations provenant de la couche 4 sont encapsulées dans le PDU de couche 3 : le paquet, dont voici les principaux éléments :

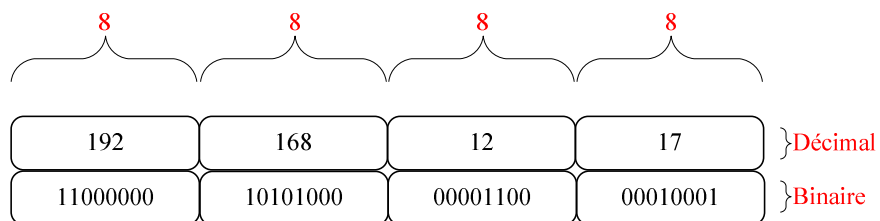


Champs	Description
Version	Indique la version de protocole IP utilisée (4 bits).
Longueur totale	Précise la longueur du paquet IP en entier, y compris les données et l'en-tête, en octets (16 bits).
Durée de vie	Un compteur qui décroît graduellement, par incréments, jusqu'à zéro. À ce moment, le datagramme est supprimé, ce qui empêche les paquets d'être continuellement en boucle (8 bits).
Somme de contrôle	Assure l'intégrité de l'en-tête IP (16 bits).
Adresse d'origine	Indique le nœud émetteur (32 bits).
Adresse de destination	Indique le nœud récepteur (32 bits).
Données	Cet élément contient des informations de couche supérieure (longueur variable, maximum 64 Ko).
Remplissage	Des zéros sont ajoutés à ce champ pour s'assurer que l'en-tête IP soit toujours un multiple de 32 bits.

6.2.2. Adressage IP

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

- Une partie désignant le réseau (on l'appelle netID)
- Une partie désignant les hôtes (on l'appelle host-ID)



Exemple d'adresse IP

Les hôtes situés sur un réseau ne peuvent communiquer qu'avec des hôtes situés sur le même réseau, même si des stations se trouvent sur le même segment. C'est ce même numéro qui permet au routeur d'acheminer le paquet au destinataire.

6.2.3. Classes d'adresses IP

L'organisme chargé d'attribuer les adresses IP publiques est l'InterNIC (Internet Network Information Center).

On appelle « Bits de poids fort », les premiers bits de l'octet le plus à gauche. Les adresses IP sont réparties en plusieurs classes, en fonction des bits qui les composent :

Classe	Bits de poids fort	Plage	Masque par défaut
A	0	1 à 126	255.0.0.0
B	10	128 à 191	255.255.0.0
C	110	192 à 223	255.255.255.0
D	1110	224 à 239	Aucun
E	1111	240 à 255	Aucun

Dans la classe A, il existe 2 adresses réservées, la plage 0.0.0.0 qui est inutilisable car non reconnue sur les réseaux, ainsi que la plage 127.0.0.0 qui est réservée pour la boucle locale.

Dans toute adresse IP, il existe 2 parties, la partie réseau et la partie hôte. Ces parties sont délimitées grâce au masque de sous réseau associé.

Les bits à 1 représentant la partie réseau et les bits à 0 la partie hôte.

Par exemple la partie réseau d'une classe C sera les 3 premiers octets et la partie hôte le dernier octet.

Il existe 2 adresses IP particulières et réservées dans un réseau, la toute première adresse IP appelée adresse réseau qui caractérise le réseau lui-même et la toute dernière de la plage est l'adresse de broadcast qui est définie par une adresse IP pouvant atteindre toutes les machines du réseau.

Pour une adresse réseau, tous les bits de la partie hôte seront à 0.

Pour une adresse broadcast, tous les bits de la partie hôte seront à 1.

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 passerelle).

Dans ce cas, seul l'ordinateur relié à Internet a besoin de réserver une adresse IP auprès de l'InterNIC. On caractérise cette adresse d'adresse publique. Toutefois, les autres ordinateurs ont tout de même besoin d'une adresse IP pour pouvoir communiquer ensemble de façon interne. Ce sont des adresses privées.

Ainsi, l'InterNIC a réservé trois plages 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 public. Il s'agit des plages d'adresse suivantes :

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

6.2.4. IPv4 et IPv6 (IPng / IP next generation)

Le protocole IPv4, le standard actuel, était censé avoir une taille suffisante pour fournir des adresses IP (2^{32} , soit 4 294 967 296 adresses possibles). Néanmoins cette limite est en passe d'être atteinte. Pour palier à cela, en 1992, l'organisme IETF (*Internet Engineering Task Force*) a alors décidé de « moderniser » le système d'adressage IP afin d'éviter cette pénurie.

Différentes solutions ont été mises en place, dans un premier temps afin de réduire cette consommation d'IP.

IPv6 emploie 128 bits à la place des 32 bits actuellement utilisés par IPv4. IPv6 emploie des nombres hexadécimaux pour représenter une adresse, alors qu'IPv4 utilise des nombres décimaux. IPv6 fournit $3,4 \times 10^{38}$ adresse IP (2^{128}). Cette version d'IP devrait donc fournir assez d'adresses pour les futurs besoins des nouveaux pays développés.

Exemple d'une adresse IP v4 :

Valeur : 34.208.123.12

Nombre d'octets utilisés : 4

Exemple d'une adresse IP v6 :

Valeur : 21DA:00D3:0000:2F3B:02AA:00FF:FE28:9C5A

Valeur simplifiée: 21DA:D3::2F3B:2AA:FF:FE28:9C5A

Nombre d'octets utilisés : 16

On peut noter que ces nouvelles adresses seront bien plus difficiles à retenir que les adresses IP actuelles : aussi l'organisme en charge de cette version a aussi créé une méthode permettant de simplifier ces IPs : on retire les 0 de chaque début de bloc et, si cela supprime un bloc, on le remplace par « :: ».

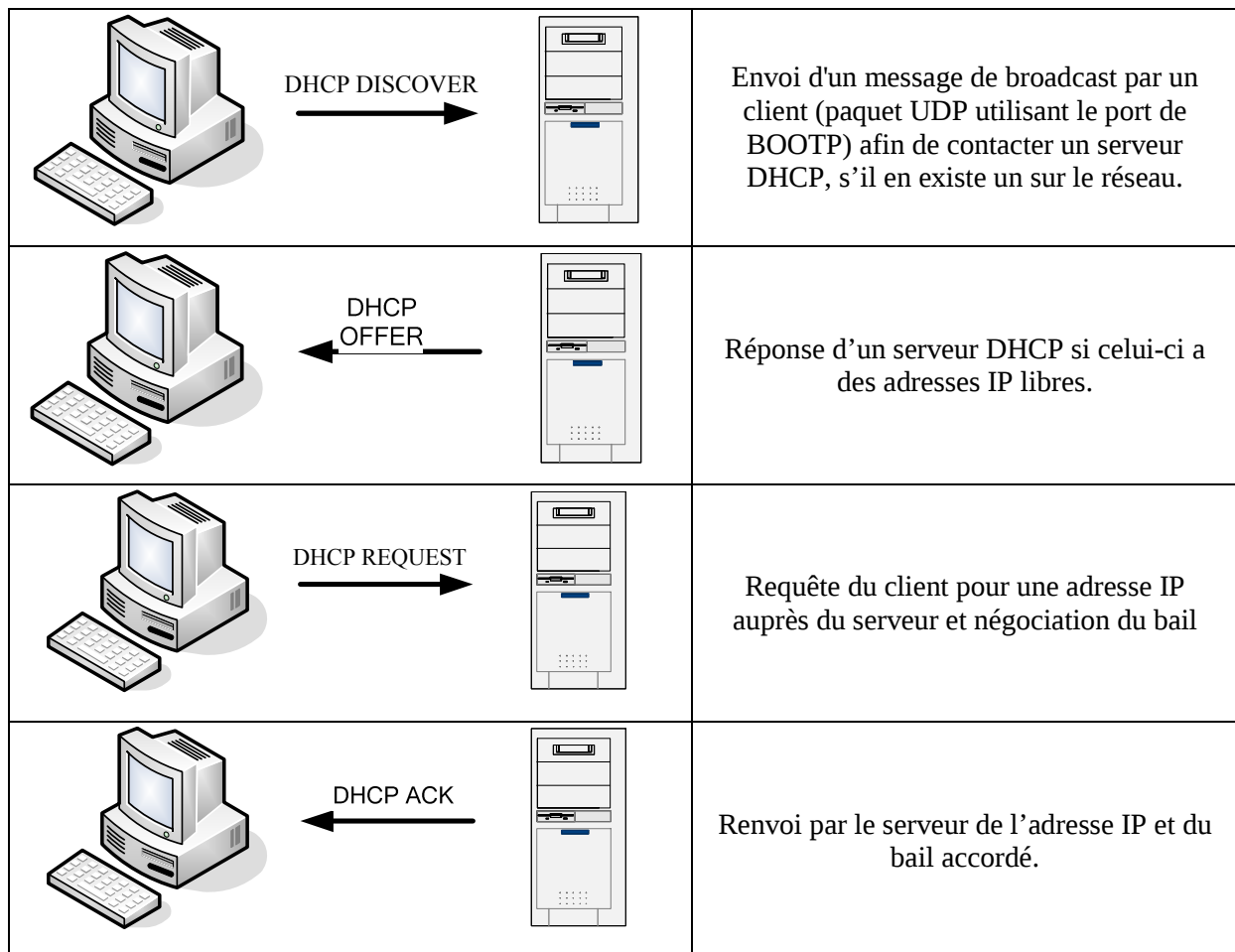
6.3. Gestion des adresses IP

6.3.1. Méthodes d'obtention

On distingue 2 méthodes d'attribution d'adresses IP pour les hôtes :

- **Statique** : chaque équipement est configuré manuellement avec une adresse unique
- **Dynamique** : On utilise des protocoles qui attribuent des IP aux hôtes
 - **RARP** : Protocole associant les adresses MAC aux adresses IP. Il permet à des stations sans disque dur local connaissant leur adresse MAC de se voir attribuer une IP.

- o **BOOTP** : Ce protocole permet à un équipement de récupérer son adresse IP au démarrage. L'émetteur envoie un message de broadcast (255.255.255.255) reçu par le serveur qui répond lui aussi par un broadcast contenant l'adresse MAC de l'émetteur ainsi qu'une IP.
- o **DHCP** : Remplaçant de BOOTP, il permet l'obtention dynamique d'IP. Lorsqu'un ordinateur entre en ligne, il communique avec le serveur qui choisit une adresse et un masque de sous réseau et l'attribue à l'hôte. Il permet de plus d'obtenir des serveurs DNS, la passerelle par défaut ainsi qu'optionnellement les adresses des serveurs WINS.



6.3.2. Résolution d'adresses

- Le protocole ARP

Le protocole ARP permet d'identifier l'adresse physique d'un hôte (adresse MAC unique) à partir de son adresse IP. ARP signifie Address Resolution Protocol.

Chaque machine connectée au réseau possède une adresse physique de 48 bits. Ce numéro unique est en fait encodé dans chaque carte réseau dès la fabrication de celle-ci en usine (adresse MAC). Toutefois, la communication sur un réseau ne se fait pas directement à partir de ce numéro car cette adresse n'est pas hiérarchique. On ne peut donc pas déterminer l'appartenance d'un hôte à un réseau à partir de cette adresse. Pour cela on utilise une adresse dite logique : 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 l'émetteur qui va stocker le couple d'adresses dans la table de correspondance et la communication sera possible.

- Le protocole RARP

Le protocole RARP (Reverse Address Resolution Protocol) permet de connaître l'adresse IP d'un hôte, à partir de son adresse physique.

Lorsqu'une machine ne connaît que l'adresse physique d'un dispositif, elle peut émettre une requête RARP afin d'avoir son adresse IP.

6.3.3. Le protocole ICMP

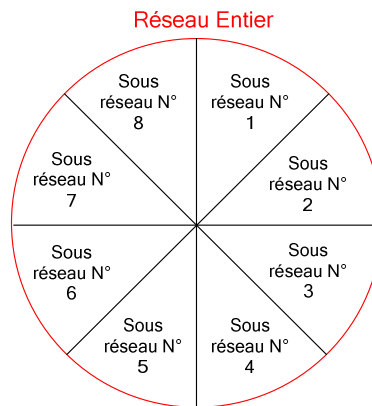
Le protocole ICMP (Internet Control Message Protocol) est un protocole qui permet de gérer les informations relatives aux erreurs générées au sein d'un réseau IP. Etant donné le peu de contrôles que le protocole IP réalise, il permet, non pas de corriger ces erreurs, mais de faire part de ces erreurs.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

Un exemple typique d'utilisation du protocole ICMP est la commande ping. Lors de l'exécution de cette commande, des informations précises peuvent être obtenues : le temps mis par un paquet pour atteindre une adresse, ou bien un éventuel problème de routage pour atteindre un hôte.

7. Couche 3 : Subnetting

7.1. Intérêt du Subnetting

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.

Il faut savoir qu'il y a une perte d'adresses quand on utilise le mécanisme de création de sous réseaux :

- Tout d'abord au niveau des sous réseaux eux-mêmes, le premier sous réseau et le dernier doivent être enlevés. En effet, La première adresse sera l'adresse de réseau : ce sera l'adresse réseau pour la globalité du réseau. La dernière plage ayant l'adresse de broadcast pour le réseau tout entier. Il faut donc enlever les deux plages entières pour éviter toute confusion. On aura donc $N-2$ sous réseaux utilisables.
- Pour les hôtes également, il y a une perte d'adresses, sans faire de sous réseaux, on avait une seule adresse réseau et une seule adresse broadcast, avec les sous réseaux, on va avoir une adresse de sous réseau à chaque sous réseau et une adresse de broadcast de sous réseau à chaque sous réseau. Il faut donc également penser à la règle des $N-2$ pour les hôtes.

7.2. Méthodes de calcul

7.2.1. Méthode classique

On entend par méthode classique le fait de procéder sans formule spécifique, par la méthode calculatoire.

Cette méthode se détaille en 6 étapes :

- Empruntez le nombre de bits suffisants
- Calculez le nouveau masque de sous réseau
- Identifiez les différentes plages d'adresses IP
- Identifiez les plages d'adresses non utilisables
- Identifiez les adresses de réseau et de broadcast
- Déterminez les plages d'adresses utilisables pour les hôtes.

▪ **Empruntez le nombre de bits suffisant**

Il faut tout d'abord déterminer le nombre de bits que l'on va emprunter à la partie réseau.

On détermine tout d'abord le nombre d'hôtes ou de sous réseaux maximums que l'on désire, car suivant ce nombre, on n'utilisera pas les même plages d'adresses (254 hôtes maximum pour une plage de classe C, 65534 pour une plage de classe B et 16 777 216 pour une plage de classe A)

On écrit en binaire le chiffre souhaité de sous-réseaux ou d'hôtes ce qui nous donne le nombre de bits à emprunter ou à laisser. Il faut penser à la règle des N-2, on cherche des plages utilisables. Il faut donc penser à additionner 2 aux hôtes ou aux sous réseaux utilisables que l'on cherche à avoir.

Pour les sous réseaux nous allons emprunter des bits à la partie hôte (allonger le masque) et pour les hôtes nous allons laisser les bits à 0 pour le nombre d'hôtes souhaités

▪ **Calculez le nouveau masque de sous réseau**

Maintenant que l'on sait combien de bits l'on va emprunter, on calcule le nouveau masque de sous réseau auquel on emprunte les bits à la partie hôte. Pour cela on prend le masque de la plage que l'on veut utiliser, on le convertit en binaire, puis on emprunte le nombre de bits nécessaires à 1 pour la création des sous réseaux.

Ou bien on laisse le nombre suffisant de bits à 0 pour les hôtes.

▪ **Identifiez les différentes plages d'adresses**

A l'aide du masque de sous réseau on calcule les différentes plages d'adresses possibles. Pour cela il suffit d'écrire chaque possibilité binaire sur les bits que l'on a empruntés pour la création des sous réseaux.

▪ **Identifiez les plages d'adresses non utilisables**

On retire maintenant la première et la dernière plage d'adresse des différents choix que l'on a. La première adresse sera l'adresse de réseau : ce sera l'adresse réseau pour la globalité du réseau. La dernière plage ayant l'adresse de broadcast pour le réseau tout entier.

▪ **Identifiez les plages de réseau et de broadcast**

Des plages d'adresses qui restent, on retire aussi les premières et dernières adresses. La première servira d'adresse réseau pour la plage d'adresse. La dernière servira d'adresse de broadcast pour la plage spécifiée.

▪ **Déterminez les plages d'adresses Hôtes.**

Maintenant qu'il ne nous reste plus que les plages d'adresses utilisables, on a donc les plages d'adresses IP utilisables par les hôtes pour communiquer sur le sous réseau.

7.2.2. Méthode du nombre magique

Cette méthode permet d'aller plus vite dans le calcul, elle est basée sur la formule que voici :

$256 = \text{Masque de sous réseau} + \text{Taille du sous réseau}$

Cette formule va vous permettre de calculer rapidement :

- Un masque de sous réseau
- Un nombre d'hôtes par sous réseau

Cette formule est propre à l'octet modifié avec le masque de sous réseau.

Elle permet de trouver le nombre d'hôtes par sous réseaux très vite, dès que l'on a le masque.

Il suffit de soustraire au nombre magique la valeur de l'octet du masque modifié, le résultat ainsi donné est la taille du sous réseau par rapport à cet octet.

Exemple :

On vient de faire du Subnetting sur une classe C, on a donc un masque résultant en 255.255.255.224.

On applique le nombre magique, $256 - 224 = 32$, il va donc y avoir 32 hôtes par sous réseau (30 utilisables).

On peut également extrapoler, et ce résultat indique donc que les plages de sous réseau seront espacées de 32.

En annexe, on peut également utiliser une formule logique afin de simplifier la création de sous réseaux :

$$256 = \text{Taille du sous réseau} * \text{Nombre de sous Réseaux}$$

Exemple :

On désire savoir le nombre d'hôtes sur 5 sous réseaux avec une classe C on aura donc un masque de type 255.255.255.X

La puissance de 2 la plus proche et supérieur à 5 est donc 8.

On prend la formule :

$$256 = \text{Taille du sous réseau} * \text{Nombre de sous Réseau}$$

Et on l'applique :

$$256 = \text{Taille du sous réseau} * 8$$

$$\text{Taille du sous réseau} = 256 / 8 = 32$$

En enlevant les 2 adresses (celle du sous réseau et celle de broadcast) on a un total de 30 adresses utilisables par sous réseau.

Cela donnera donc un masque de 255.255.255.224 ($256-32 = 224$)

Et donnera une donc une configuration de type :

Adresse de début du sous réseau : 192.168.0.32

Adresse de fin du sous réseau : 192.68.0.63

Adresse de début du sous réseau : 192.168.0.64

Adresse de fin du sous réseau : 192.68.0.95

Adresse de début du sous réseau : 192.168.0.96

Adresse de fin du sous réseau : 192.68.0.127

Et ainsi de suite

En utilisant ces 2 formules, il est donc beaucoup plus rapide de calculer un masque de sous réseau ou un nombre d'hôte. Néanmoins il vaut mieux bien comprendre la méthode de base avant d'utiliser celle-ci, afin de ne pas faire d'erreur lorsque vous les utilisez, toujours garder à l'esprit que ces formules sont valides uniquement pour l'octet modifié par la création de sous réseaux.

8. Couche 3 : Introduction au routage

8.1. Principes fondamentaux

Avant de commencer ce chapitre, il convient de définir commutation de trames et commutation de paquets (routage). Car, si au premier abord il pourrait sembler que ces 2 termes désignent la même chose, ce n'est pas du tout le cas. La première distinction vient du fait que la commutation de trames s'effectue au niveau de la couche 2 du modèle OSI, alors que le routage s'effectue au niveau de la couche 3 du modèle OSI. Cela indique donc que les routeurs et les commutateurs ne prennent pas leur décision avec les mêmes informations.

Pour joindre les hôtes non locaux, une machine va faire une requête ARP pour avoir l'adresse MAC de la station de destination, si la destination n'est pas locale la requête ARP va échouer, la station enverra alors la trame à sa passerelle par défaut, c'est-à-dire au routeur.

Le routeur examine l'adresse de destination de la couche 3 du paquet, effectue un ET logique binaire avec le masque de sous réseau pour identifier le réseau de destination et prendre la bonne décision de commutation.

De la même manière qu'un commutateur garde une table des adresses MAC connues, un routeur garde une table des adresses réseaux dans sa table de routage. Il va ainsi être capable de commuter les paquets vers un réseau spécifique.

8.2. Domaine de broadcast

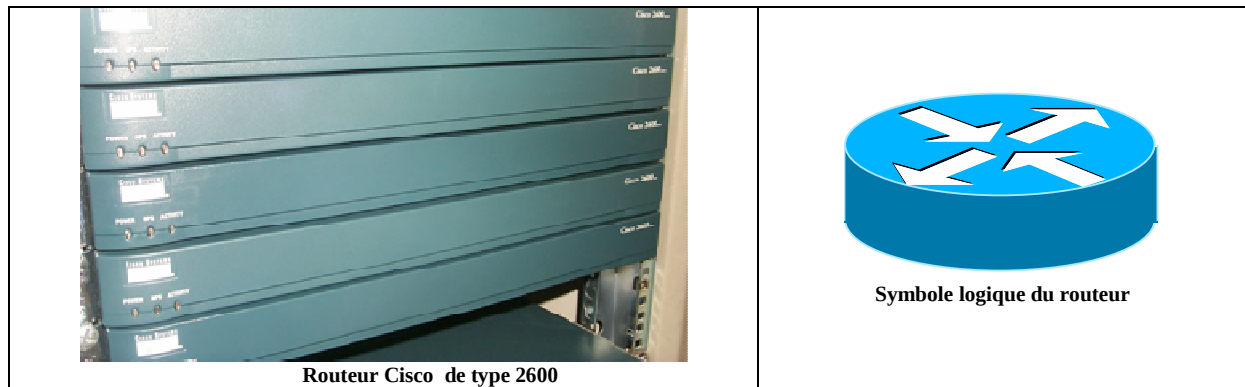
Un domaine de broadcast est un domaine logique ou n'importe quels hôtes connectés à un réseau peuvent envoyer des données à une autre machine sans passer par des services de routage.

Plus spécifiquement c'est un segment réseau composé d'hôtes et de dispositifs pouvant être atteint en envoyant un paquet à l'adresse de broadcast. Ces domaines de broadcast sont toujours séparés par des dispositifs de couche 3.

8.3. Les équipements de couche 3 : les routeurs

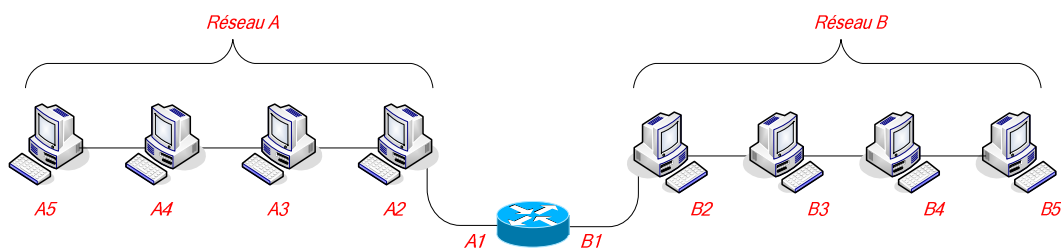
Routeur :

Équipement de couche 3 permettant d'interconnecter deux réseaux ou plus en se basant sur les adresses de couche 3. Le routeur permet également une segmentation des domaines de broadcast et des domaines de collisions.



Le routeur dispose d'une interface (une carte réseau) le reliant au réseau local. Celle-ci dispose d'une adresse IP.

Par exemple, sur le schéma ci-dessous, les adresses des hôtes sont A5, A4, A3 et A2, faisant partie du réseau A. On attribue A1 à l'interface du routeur, lui permettant ainsi de se connecter au réseau A. Un autre réseau, B, est lui aussi connecté au routeur. Ce dernier dispose donc d'une interface ayant pour IP B1 afin de pouvoir communiquer avec le réseau.



Supposons maintenant que l'on souhaite envoyer des données de A vers B :

- Le routeur reçoit la trame de couche 2, supprime l'en tête de liaison de données
- Il examine l'adresse de couche 3 afin de déterminer le destinataire
- Il effectue un ET logique entre l'adresse IP et le masque de sous réseau afin de déterminer le réseau de destination
- Il consulte sa table de routage pour déterminer l'interface par laquelle les données doivent être envoyées.

C'est pour cela que chaque interface du routeur doit être sur un réseau différent. Sinon le routeur ne pourra pas déterminer par quelle interface envoyer les informations. C'est le principe de commutation de paquets ou routage.

8.4. Détermination 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 :

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

8.5. Systèmes autonomes, IGP et EGP

Un système autonome est un réseau ou un ensemble de réseaux sous un contrôle administratif commun. Un système autonome est composé de routeurs ayant les mêmes règles et fonctions.

Deux familles des protocoles de routage sont les protocoles IGP (Interior Gateway Protocol) et les protocoles EGP (Exterior Gateway Protocol).

Les IGP routent les données dans un système autonome, comme nous venons de le voir :

- RIP and RIPv2
- IGRP
- EIGRP
- OSPF
- IS-IS

EGP route les données entre les réseaux autonomes. Un exemple d'EGP est BGP.

8.6. Routage statique et dynamique

Il existe différents protocoles de routage permettant de trouver le meilleur chemin. Chaque protocole utilise différents systèmes, différents algorithmes pour fournir au routeur les informations nécessaires à la mise en place de la table de routage.

Voici un tableau récapitulatif de ces différents protocoles avec leurs descriptions :

Nom du protocole	Type (IGP ou EGP)	Algorithme	Métriques	Mise à jour	Remarque
RIP	IGP	Vecteur de distance	15 sauts maximums	30 sec	15 sauts maximums
RIP v2	IGP	Vecteur de distance	15 sauts maximums	30 sec	Inclus des préfixes de routage et les masques de sous réseau dans les informations de routage
IGRP	IGP	Vecteur de distance	Délais, charge, bande passante, fiabilité	90 secondes	Choisi le meilleur chemin selon différent critères. Propriétaires Cisco.
EIGRP	IGP	Hybride	Délais, charge,	Instantanée	Propriétaire Cisco.

			bande passante, fiabilité	à chaque changement topologique	Meilleur convergence et moins de bande passante utilisée.
OSPF	IGP	Etat de lien	Le coût de la route	Instantanée à chaque changement topologique	Utilisé pour les réseaux à grandes échelles
IS-IS	IGP	Etat de lien	Poids du lien	Instantanée à chaque changement topologique	Supporte de multiples protocoles routés tel qu'IP.
BGP	EGP	Vecteur de chemin	Politique réseau, Attribut de chemin		Protocole utilisé par la plupart des ISP et les grandes compagnies.

9. Couche 4 : Couche transport

9.1. Introduction

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.

9.2. TCP et UDP

La pile de protocoles TCP/IP comprend 2 protocoles de couche 4 : TCP et UDP

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

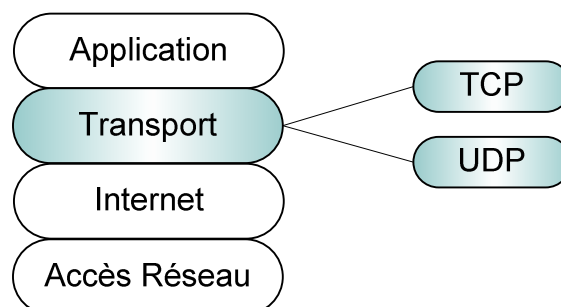
- Fiabilité
- Division des messages sortants en segments
- Ré assemblage des messages au niveau du destinataire
- Ré envoi de toute donnée non reçue

Segments : PDU de couche 4

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.



9.2.1. 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	Secure SHell
Telnet	23	Telnet
SMTP	25	Simple Mail Transfer
DNS	53	Domain Name System
HTTP	80	World Wide Web HTTP
POP3	110	Post Office Protocol - Version 3
NNTP	119	Network News Transfer Protocol
IMAP2	143	Interactive Mail Access Protocol v2
NEWS	144	News
HTTPS	443	Protocole HTTP sécurisé (SSL)

Numéros de ports

Les ports sont attribués de la manière suivante :

Plage de ports

Utilisation

0 à 1023

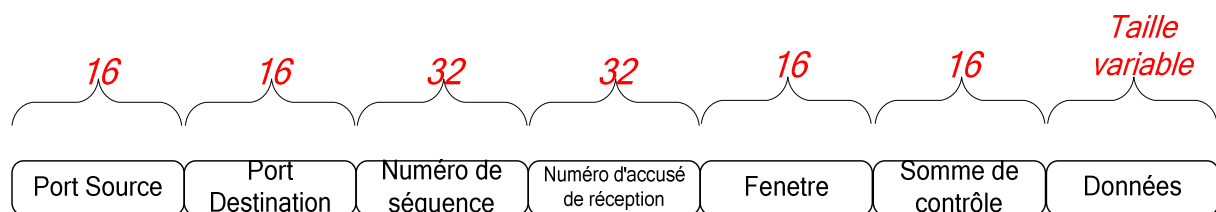
réservés aux applications publiques

1023 à 65535

attribué aux entreprises pour les applications commerciales et utilisé par le système d'exploitation pour l'attribution dynamique des ports source.

9.2.2. Structures d'un segment TCP

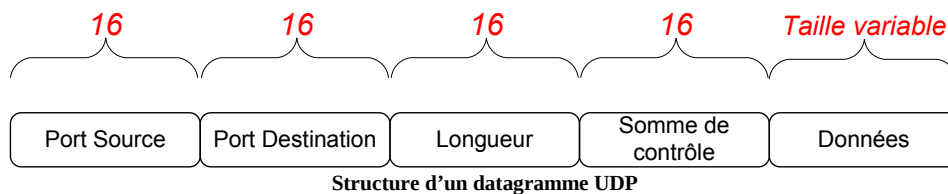
Le protocole TCP encapsule les informations provenant de la couche supérieure dans des segments dont voici les principales informations :



Champs	Descriptions
Port source	Numéro du port appelant
Port de destination	Numéro du port appelé
Numéro de séquence	Numéro utilisé pour assurer le séquençage correct des données entrantes
N° d'accusé de réception	Prochain octet TCP attendu
Somme de contrôle	Somme de contrôle calculée des champs d'en-tête et de données
Données	Données du protocole de couche supérieure

9.2.3. Structure d'un datagramme UDP

UDP étant un protocole non orienté connexion, il dispose d'un en-tête de taille réduite par rapport aux en-têtes des segments TCP :



Le protocole UDP est conçu pour les applications ne devant pas assembler de séquences de segments. Il laisse aux protocoles de la couche application le soin d'assurer la fiabilité.

9.3. Méthode de connexion TCP

Un service orienté connexion comportent 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

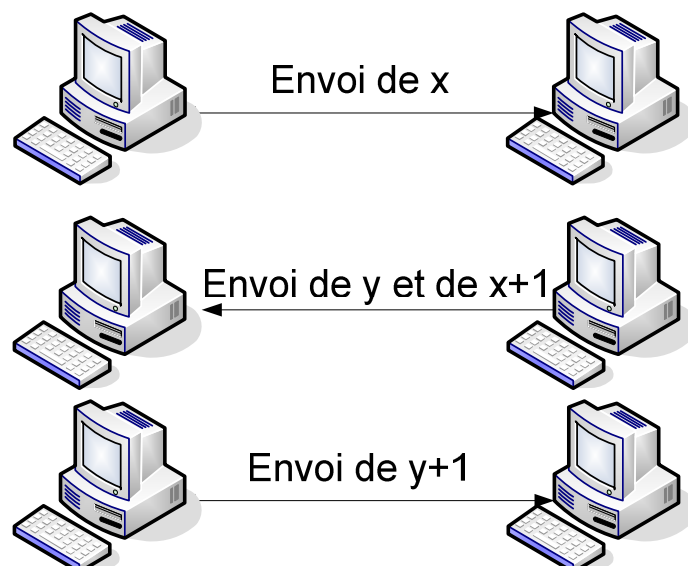
9.3.1. Connexion ouverte/échange en 3 étapes

Les hôtes TCP établissent une connexion en 3 étapes, appelée 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.

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).

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



Il existe également des méthodes garantissant la fiabilité des protocoles

9.3.2. 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 plus lentement et un nouveau compteur est déclenché.

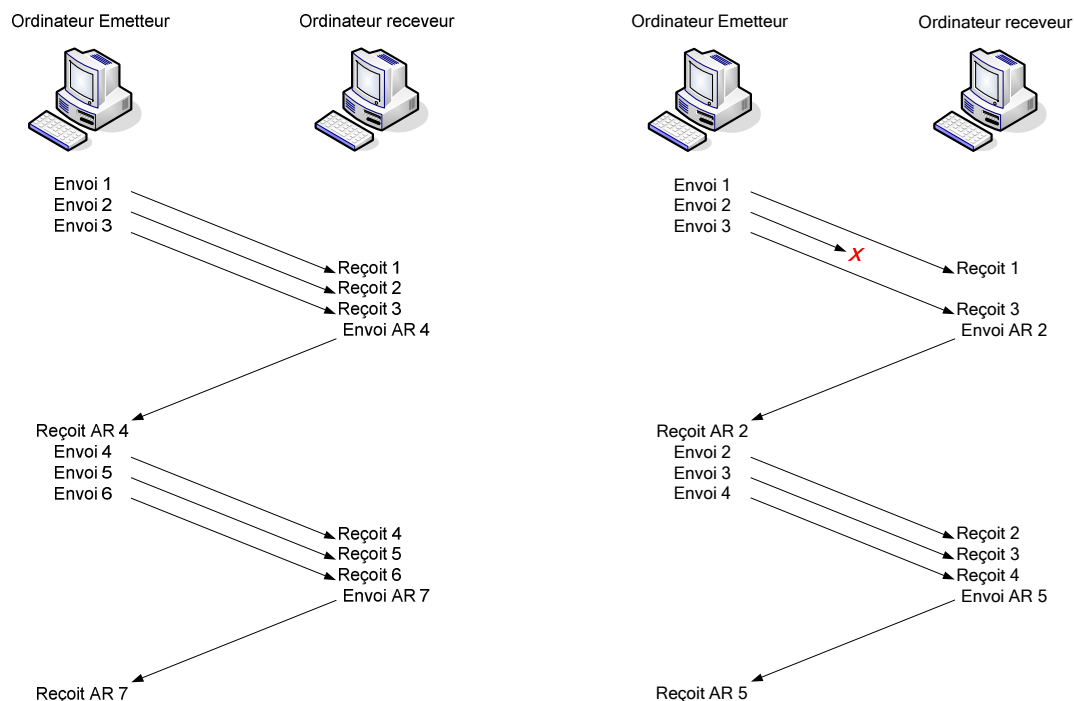
Cependant, cette technique est consommatrice de bande passante, c'est alors qu'intervient le mécanisme de fenêtrage.

9.3.3. 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



Transmission sans perte de paquets

Transmission avec perte de paquets : ici le paquet 2 est renvoyé (le 3 aussi même s'il a été reçu).

10. Couche 5 : Couche session

Comme nous l'avons vu précédemment, 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 qu'elle prend en compte :

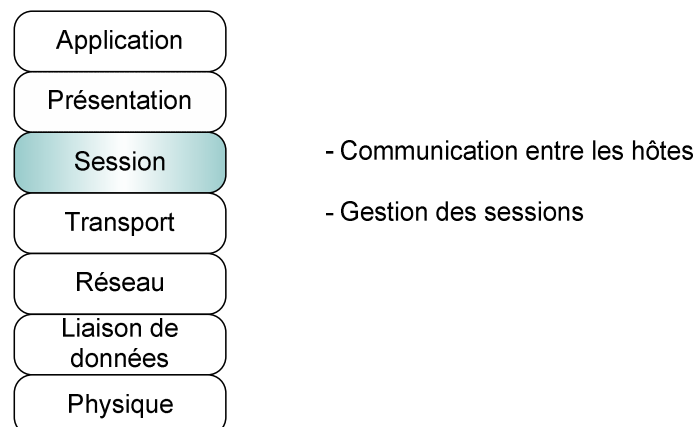
- le lancement des sessions
- la resynchronisation du dialogue
- l'arrêt des sessions

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

Une communication entre ordinateurs suppose de nombreuses conversations courtes (commutation de paquets comme nous l'avons vu précédemment) 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).

Le contrôle du dialogue consiste en l'identification des rôles de chacun à un moment donné.



10.1. Contrôle du dialogue

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

- Si la communication bidirectionnelle simultanée est permise :
 - La gestion de la communication est assurée par d'autres couches des ordinateurs en communication.
- Si ces collisions au sein de la couche session sont intolérables, le contrôle de dialogue dispose d'une autre option : la communication bidirectionnelle alternée
 - Ce type de communication est rendu possible par l'utilisation d'un jeton de données au niveau de la couche session qui permet à chaque hôte de transmettre à tour de rôle.

10.2. Synchronisation du dialogue

Cette étape est des plus importantes, elle permet aux hôtes communicants au travers d'un réseau de marquer une pause pour par exemple sauvegarder la communication en cours et resynchroniser le dialogue.

Pour cela ils utilisent un « point de contrôle », envoyé par l'un des interlocuteurs à l'autre pour enregistrer la conversation, vérifier l'heure de la dernière portion de dialogue effectuée, comme si vous aviez un double appel avec votre cellulaire. Ce processus est appelé la synchronisation du dialogue.

Comme dans le langage humain, il est important dans une discussion de montrer à son interlocuteur le début d'une conversation (« allo » dans le cas d'une conversation téléphonique) ainsi que de signifier que l'on se prépare à mettre fin à la conversation (« au revoir »). C'est pour cela que les deux contrôles principaux sont :

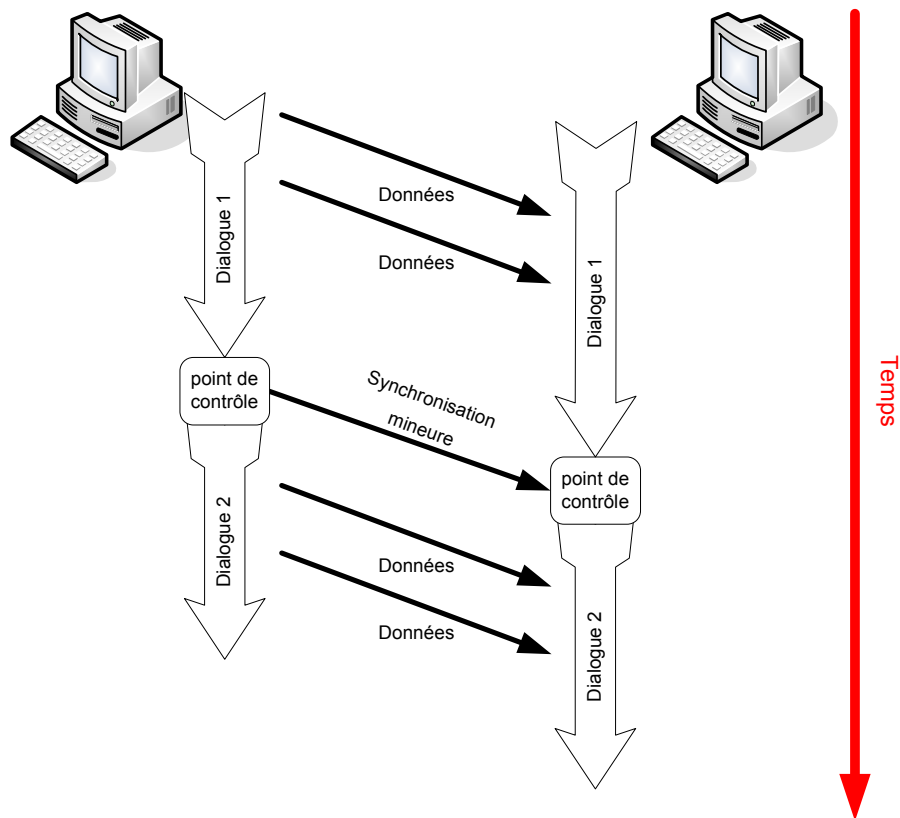
- Lancement ordonné de la communication
- Fin de la communication

10.3. Division du dialogue

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

Notre schéma représente une petite synchronisation. Au niveau du point de contrôle, la couche session de l'hôte A envoie un message de synchronisation à l'hôte B, et les deux hôtes exécutent la séquence qui suit :

- Sauvegarder les fichiers donnés.
- Sauvegarder les paramètres réseau.
- Sauvegarder les paramètres de synchronisation.
- Noter le point d'extrémité de la conversation.



Les points de contrôle fonctionnent comme un logiciel de traitement de texte lorsqu'il fait une pause d'une seconde pour effectuer une sauvegarde automatique d'un document. Ces points de contrôle servent toutefois à séparer les parties d'une session, préalablement appelées dialogues.

Nous venons de voir comment les hôtes s'organisent autour de la communication, nous allons maintenant voir comment les données sont générées pour que les hôtes se comprennent.

11. Couche 6 : Couche présentation

Afin que deux hôtes communiquant puissent se comprendre, il est nécessaire qu'il parle le même langage : c'est à cette tâche qu'est dévolue la couche présentation.

11.1. Fonctions et normes

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 peut être comparée à un traducteur lors d'une conférence internationale : elle s'occupe de « traduire » les données de manière à ce que l'hôte récepteur soit en mesure de comprendre.

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

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

Après avoir reçu les données de la couche application, la couche présentation exécute certaines ou toutes ces fonctions avant d'acheminer les données à la couche session.

Au niveau de la station de réception, la couche présentation reçoit les données de la couche session et exécute les fonctions nécessaires avant de les faire suivre à la couche application.

Les normes de la couche 6 définissent également la présentation des graphiques. Les trois principaux formats graphiques sont :

- BMP (BitMaP) est un format ancien encore largement répandu, il est maintenant supplanté par le JPEG, qui fournit des fichiers avec un meilleur taux compression/taille
- JPEG (Joint Photographic Experts Group) - Format graphique le plus utilisé pour la compression des images fixes complexes et des photographies.
- PNG (Portable Network Graphics) est un format graphique en émergence sur Internet qui compresse les textures.

D'autres normes de la couche 6 concernent la présentation des sons et des séquences animées. Les normes suivantes appartiennent à cette catégorie:

- MPEG (Motion Picture Experts Group) - Format de compression et de codage de vidéo animée pour CD ou tout autre support de stockage numérique.
- MP3 (MPEG Layer 3) - Format de compression de musique le plus utilisé pour le moment. Il utilise l'étude de l'oreille humaine ainsi des algorithmes de compression.
- Divx (MPEG 4) format de compression créé à partir du format MPEG 4 développé par Microsoft et permettant une compression bien meilleure que le MPEG 1 ou 2 (exemple : faire tenir un film sur un CD au lieu d'un DVD).

**Représentation des données :**

- o Lisibilité des données par le destinataire
- o Formatage des données
- o Contrôle de la syntaxe

Les normes de la couche présentation établissent donc des standards de formats de fichier afin que les hôtes soient en mesure de comprendre les informations.

11.2. Le cryptage des données

Le cryptage est défini par l'utilisation d'algorithmes permettant d'encoder le message de manière à ce que seul l'hôte à qui on l'adresse puisse le comprendre.

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

Une clé de cryptage peut être utilisée pour crypter les données à la source en encodant les données avec elle, ce qui obligera l'hôte récepteur à posséder cette clé pour les décrypter. Un algorithme est donc utilisé pour rendre ces données incompréhensibles à quiconque ne disposant pas de la clé.

11.3. 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 »"

On peut aussi utiliser un dictionnaire pour remplacer certains mots trop long : ils sont constitués des mots ou des séquences revenant le plus souvent ainsi que des séquences de remplacement, de manière à réduire considérablement les fichiers.

12. Couche 7 : Couche application

12.1. Introduction:

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 :

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

Dans le modèle OSI, la couche application est la plus proche du système terminal (ou la plus proche des utilisateurs).

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.

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.

Voici en détails les principaux protocoles utilisés par la couche transport :

12.2. DNS

12.2.1. Présentation du protocole DNS

Chaque station possède une adresse IP propre. Cependant il est nettement plus simple de travailler avec des noms de stations ou des adresses plus explicites comme par exemple <http://www.labo-cisco.com>, qu'avec des adresses IP.

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.

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

Aux origines de TCP/IP, étant donné que les réseaux étaient très peu étendus, c'est-à-dire que le nombre d'ordinateurs connectés à un même réseau était faible, les administrateurs réseau créaient des fichiers appelés tables de conversion statique (fichiers généralement appelé hosts ou hosts.txt), associant sur une ligne, grâce à des caractères ASCII, l'adresse IP de la machine et le nom littéral associé, appelé nom d'hôte.

Ce système à l'inconvénient majeur de nécessiter la mise à jour des tables de tous les ordinateurs en cas d'ajout ou modification d'un nom de machine. Ainsi, avec l'explosion de la taille des réseaux, et de leur interconnexion, il a fallu mettre en place un système plus centralisé de gestion des noms. Ce système est nommé Domain Name System, traduisez Système de nom de domaine.

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.

Les machines appelées serveurs de nom de domaine permettent d'établir la correspondance entre le nom de domaine et l'adresse IP sur les machines d'un réseau. Chaque domaine possède ainsi, un serveur de noms de domaines, relié à un serveur de nom de domaine de plus haut niveau. Ainsi, le système de nom est une architecture distribuée, c'est-à-dire qu'il n'existe pas d'organisme ayant à charge l'ensemble des noms de domaines. Par contre, il existe un organisme (l'InterNIC pour les noms de domaine en .com, .net, .org et .edu par exemple). Le système de noms de domaine est transparent pour l'utilisateur, néanmoins il ne faut pas oublier les points suivants.

Chaque ordinateur doit être configuré avec l'adresse d'une machine capable de transformer n'importe quel nom en une adresse IP. Cette machine est appelée Domain Name Server.

L'adresse IP d'un second Domain Name Server (secondary Domain Name Server) peut également être introduite : il peut relayer le premier en cas de panne.

12.2.3. Codes des domaines internet

La classification du domaine, parfois appelées TLD (Top Level Domain, soit domaines de plus haut niveau), correspond généralement à une répartition géographique.

Toutefois, il existe des noms, créés pour les Etats-Unis à la base, permettant de classifier le domaine selon le secteur d'activité, par exemple :

- .com correspond aux entreprises à vocation commerciales (désormais ce code de domaine ne rime plus à grand chose et est devenu international)
- .edu correspond aux organismes éducatifs
- .gov correspond aux organismes gouvernementaux
- .net correspond aux organismes ayant trait aux réseaux
- .org correspond aux entreprises à but non lucratif
- .biz correspond aux entreprises en générale
- .info réservé aux sites d'informations

12.3.FTP et TFTP

12.3.1. FTP

FTP est un protocole fiable et orienté connexion qui emploie TCP pour transférer des fichiers entre les systèmes qui supportent ce protocole. Le but principal du ftp est de transférer des fichiers à partir d'un ordinateur à un autre en copiant et/ou en déplaçant des fichiers des serveurs aux clients, et des clients vers les serveurs. Le protocole FTP est assigné au port 21 par défaut.

Quand des fichiers sont copiés d'un serveur, FTP établit d'abord une connexion de contrôle entre le client et le serveur. Alors une deuxième connexion est établie, qui est un lien entre les ordinateurs par lequel les données sont transférées. Le transfert de données peut se faire en mode Ascii ou en mode binaire. Ces modes déterminent le codage utilisé pour le fichier de données, qui dans le modèle OSI est une tâche de couche présentation, comme nous l'avons vu précédemment.

Après que le transfert de fichiers ait fini, la connexion de transfert de données se coupe automatiquement. La connexion de contrôle est fermée quand l'utilisateur se déconnecte et clôt la session.

12.3.2. TFTP

TFTP est un service non orienté connexion qui emploie UDP. TFTP (Trivial FTP) est employé sur un routeur pour transférer des dossiers de configuration et des images d'IOS de Cisco et aussi pour transférer des fichiers entre les systèmes qui supportent TFTP. TFTP est conçues pour être léger et simple à utiliser. Néanmoins TFTP peut lire ou écrire des fichiers sur un serveur à distance mais il ne peut pas lister les répertoires et ne supporte pas une authentification utilisateur. Il est utile dans certains LANs parce qu'il fonctionne plus rapidement que le ftp.

12.4.HTTP

Le protocole de transfert hypertexte (HTTP) fonctionne avec le World Wide Web, qui est la partie la plus utilisée et la plus importante d'Internet. Une des raisons principales de cette croissance extraordinaire est la facilité avec laquelle il permet l'accès à l'information.

Un navigateur web est une application client/serveur, qui implique l'existence d'un client et d'un serveur, composant spécifique installé sur les 2 machines afin de fonctionner.

Un navigateur web présente des données dans un format multimédia, c'est-à-dire un contenu réagissant aux actions de l'utilisateur. Le contenu peut être du texte, des graphiques, du son, ou de la vidéo.

Les pages web sont écrites en utilisant l'HTML (*HyperText Markup Language*) : un navigateur web reçoit la page au format HTML et l'interprète de manière à afficher la page d'une manière beaucoup plus agréable qu'un document texte.

Pour déterminer l'adresse IP d'un serveur HTTP distant, le navigateur utilise le protocole DNS pour retrouver l'adresse IP à partir de l'URL. Les données qui sont transférées au serveur HTTP contiennent la localisation de la page Web sur le serveur.

Le serveur répond à la requête par l'envoi au navigateur du code html ainsi que des différents objets multimédia qui agrémentent la page (son, vidéo, image) et qui sont indiqués dans les instructions de la page HTML. Le navigateur rassemble tous les fichiers pour créer un visuel de la page Web, et termine la session avec le serveur. Si une autre page est demandée, le processus entier recommence.

12.5.SMTP

Les serveurs d'email communiquent entre eux en employant le *Simple Mail Transfer Protocol (SMTP)* pour envoyer et recevoir du courrier. Le protocole SMTP achemine des messages email dans le format Ascii en utilisant TCP. On l'utilise souvent en tant que protocole d'envoi de mail, rarement en tant que protocole de récupération d'email, car il est peu sécurisé et surtout n'offre aucune authentification.

12.6.SNMP

Le *Simple Network Management Protocol (SNMP)* est un protocole de la couche application qui facilite l'échange d'information de gestion entre les dispositifs d'un réseau. Le SNMP permet à des administrateurs réseau de contrôler l'état du réseau, détecter et résoudre des problèmes de réseau, et de prévoir le développement du réseau, si jamais celui-ci arrive à saturation. Le SNMP emploie le protocole UDP en tant que protocole de couche transport.

Un réseau contrôlé par SNMP comprend les trois composants clés suivants:

- **Système de gestion de réseau (NMS / Network Management System)** : NMS exécute les applications qui supervisent et contrôlent les dispositifs gérés. Un ou plusieurs NMS doivent exister sur n'importe quel réseau géré.
- **Dispositifs managés** : Les dispositifs managés sont des nœuds du réseau qui contiennent un agent SNMP et qui résident sur un réseau managé. Les dispositifs managés rassemblent et stockent des informations de gestion et rendent cette information disponible à NMS à l'aide des dispositifs SNMP. Les dispositifs managés, parfois appelés éléments de réseau, peuvent être des routeurs, des serveurs d'accès, des commutateurs, et des ponts, des concentrateurs, des ordinateurs hôtes, ou des imprimateurs.
- **Agents** : Les agents sont des modules de logiciel réseau - gestion qui résident dans des dispositifs managés. Un agent a la connaissance locale d'information de gestion et traduit cette information en un format compatible avec SNMP.

12.7.Telnet

12.7.1. Présentation du protocole Telnet

Le protocole Telnet est un protocole standard d'Internet permettant l'interfaçage de terminaux et d'applications à travers Internet. Ce protocole fournit les règles de base pour permettre de relier un client à un interpréteur de commande (côté serveur).

Le protocole Telnet s'appuie sur une connexion TCP pour envoyer des données au format ASCII codées sur 8 bits entre lesquelles s'intercalent des séquences de contrôle Telnet. Il fournit ainsi un système orienté communication, bidirectionnel alterné (half-duplex), codé sur 8 bits facile à mettre en œuvre.

Le protocole Telnet repose sur trois concepts fondamentaux :

- Le paradigme du terminal réseau virtuel (NVT)
- Le principe d'options négociées
- Les règles de négociation

Ce protocole est un protocole de base, sur lequel s'appuient certains autres protocoles de la suite TCP/IP (FTP, SMTP, POP3, etc.).

Les spécifications de Telnet ne mentionnent pas d'authentification, car Telnet est totalement séparé des applications qui l'utilisent (le protocole FTP définit une séquence d'authentification au-dessus de Telnet).

En outre, le protocole Telnet est un protocole de transfert de données non sûr, c'est-à-dire que les données qu'il véhicule circulent en clair sur le réseau (de manière non chiffrée). Lorsque le protocole Telnet est utilisé pour connecter un hôte distant à la machine sur lequel il est implémenté en tant que serveur, ce protocole est assigné au port 23.

Hormis les options et les règles de négociation associées, les spécifications du protocole Telnet sont basiques. La transmission de données à travers Telnet consiste uniquement à transmettre les octets dans le flux TCP (le protocole Telnet précise tout de même que les données doivent par défaut, c'est-à-dire si aucune option ne précise le contraire, être groupées dans un tampon avant d'être envoyées. Plus exactement cela signifie que par défaut les données sont envoyées ligne par ligne). Lorsque l'octet 255 est transmis, l'octet suivant doit être interprété comme une commande. L'octet 255 est ainsi nommé IAC (Interpret As Command, traduisez Interpréter comme une commande).

12.7.2. La notion de terminal virtuel

Aux débuts d'Internet, le réseau (ARPANET) était composé de machines dont les configurations étaient très peu homogènes (claviers, jeux de caractères, résolutions, longueur des lignes d'affichage). D'autre part, les sessions des terminaux possédaient également leur propre façon de contrôler les flux de données en entrée/sortie.

Ainsi, au lieu de créer des adaptateurs pour chaque type de terminal afin qu'il puisse y avoir une interopérabilité de ces systèmes, il a été décidé de mettre au point une interface standard, appelée NVT (Network Virtual Terminal, traduisez Terminal réseau virtuel), fournissant une base de communication standard, composée de :

- Caractères ASCII 7 bits auxquels s'ajoutent le code ASCII étendu
- Trois caractères de contrôle
- Cinq caractères de contrôle optionnels
- Un jeu de signaux de contrôle basique

Le protocole Telnet consiste ainsi à créer une abstraction du terminal, permettant à n'importe quel hôte (client ou serveur) de communiquer avec un autre hôte sans connaître ses caractéristiques.

Table des matières

1. Réseaux WAN	5
1.1. Définition	5
1.2. Dispositifs WAN	5
1.3. Normes WAN	6
1.3.1. Normes WAN de la couche physique	6
1.3.2. Normes WAN de la couche liaison de données	7
1.4. Technologies WAN	8
1.4.1. Services à commutation de circuits	8
1.4.2. Services à commutation de paquets/cellules	9
1.4.3. Services dédiés	10
1.4.4. Autres services	10
2. Introduction aux routeurs.....	11
2.1. Présentation d'un routeur Cisco	11
2.1.1. Composants internes	11
2.1.2. Composants externes	13
2.2. Branchements.....	14
2.2.1. Interfaces LAN et WAN.....	14
2.2.2. Accès pour configuration	14
2.3. Système d'exploitation Cisco IOS	15
2.3.1. Principes et spécifications	15
2.3.2. Modes de commandes	16
2.3.3. Système d'aide	17
2.3.4. Commandes d'édition avancée.....	17
2.3.5. Historique des commandes.....	18
2.3.6. Fichiers de configuration.....	18
3. Configuration de base d'un routeur	20
3.1. Commandes de visualisation d'état	20
3.2. Date et heure	21
3.3. Nom d'hôte et résolution de noms	21
3.4. Descriptions et bannière de connexion	23
3.5. Mots de passe.....	23
3.6. Serveur HTTP.....	24
3.7. Configuration des interfaces	25
3.7.1. Interfaces Loopback	25
3.7.2. Interfaces Ethernet/IEEE 802.3	25
3.7.3. Interfaces série.....	26
4. Informations et accès aux autres dispositifs.....	27
4.1. CDP.....	27
4.1.1. Théorie	27
4.1.2. Configuration	28
4.1.3. Visualisation et résolution de problèmes.....	28
4.2. Telnet.....	28
4.2.1. Théorie	28
4.2.2. Commandes et utilisation	29

5. Gestion d'IOS et processus de démarrage.....	30
5.1. Processus de démarrage	30
5.1.1. Séquence d'amorçage.....	30
5.1.2. Commandes boot system.....	31
5.1.3. Registre de configuration	31
5.1.4. Mode SETUP	32
5.2. Gestion d'IOS	33
5.2.1. Informations générales	33
5.2.2. Gestion des systèmes de fichiers	34
5.2.3. Mode RXBoot	34
6. Routage	36
6.1. Principes fondamentaux.....	36
6.1.1. Fonctions de routage et de commutation.....	36
6.1.2. Processus de transmission	37
6.1.3. Table(s) de routage.....	38
6.2. Routage statique et dynamique	40
6.2.1. Caractéristiques et comparatif.....	40
6.2.2. Caractéristiques des protocoles de routage.....	40
6.3. Convergence, boucles de routage et solutions	41
6.3.1. Convergence.....	41
6.3.2. Boucles de routage	41
6.3.3. Métrique de mesure infinie.....	42
6.3.4. Split Horizon	42
6.3.5. Route Poisoning	42
6.3.6. Mises à jour déclenchées.....	43
6.3.7. Compteurs de retenue.....	43
6.4. Routage à vecteur de distance.....	44
6.5. Routage à état de liens	45
6.6. Systèmes autonomes, protocoles de routage intérieurs et extérieurs	46
6.7. Configuration par défaut, routage statique et visualisation d'état	47
7. Protocole RIP	49
7.1. Théorie.....	49
7.2. Configuration.....	50
7.2.1. Commandes.....	50
7.2.2. Procédure de configuration	51
7.3. Vérification	51
8. Protocole IGRP	52
8.1. Théorie.....	52
8.2. Configuration.....	54
8.2.1. Commandes.....	54
8.2.2. Procédure de configuration	55
8.3. Vérification	56

9. Protocole ICMP	57
9.1. Théorie	57
9.2. Messages ICMP	58
9.2.1. Types de messages	58
9.2.2. Echo Request/Reply	58
9.2.3. Destination Unreachable	59
9.2.4. Parameter Problem	59
9.2.5. Source Quench	60
9.2.6. Redirect/Change Request	60
9.2.7. Timestamp Request/Reply	61
9.2.8. Information Request/Reply	61
9.2.9. Address Mask Request/Reply	61
9.2.10. Router Discovery/Solicitation	61
10. Résolution de problèmes	62
10.1. Commandes de vérification	62
10.2. Erreurs courantes et modèle OSI	63
10.3. Débogage	63
10.4. Procédure de récupération des mots de passe d'un routeur	64
11. ACL	65
11.1. Théorie	65
11.1.1. Principe fondamental	65
11.1.2. Masque générique	66
11.2. ACL standard	67
11.3. ACL étendue	67
11.4. ACL nommée	68
11.5. Mise en place et vérification des ACLs	69

1. Réseaux WAN

1.1. Définition

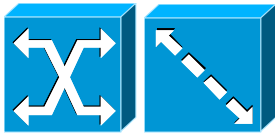
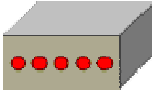
Par définition, un réseau WAN est :

- Un réseau longue distance.
- Un réseau qui interconnecte des réseaux LAN qui sont généralement 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.

1.2. Dispositifs WAN

	Routeur
	Commutateur (ATM, Frame Relay, RNIS)
	Modem et unité CSU/DSU (modem analogique, modem câble, unité CSU/DSU pour T1/E1, TA et NT1 pour RNIS)
	Serveur de communication (PABX)

Les dispositifs WAN les plus couramment utilisés sont les suivants :

- **Routeur** : Dispositif de couche 3 basant ses décisions d'acheminement sur les adresses de la couche réseau (IP, IPX, etc.). Il offre des interfaces LAN et WAN permettant l'interconnexion des réseaux locaux au réseau mondial (Internet).
- **Commutateur** : Dispositif de couche 2 qui assure la commutation du trafic WAN. Ce dispositif est présent au cœur d'un réseau WAN.
- **Modem et unité CSU/DSU** : Unité de couche 1 agissant au niveau de la forme du signal électrique. Ce dispositif se place aux extrémités des liaisons WAN, adaptant ainsi les signaux au format désiré pour chaque côté.
- **Serveur de communication** : Il concentre les communications utilisateur entrantes et sortantes.

1.3. Normes WAN

Les principaux organismes définissant et gérant les normes WAN sont les suivants :

- **UIT-T** (Union Internationale des Télécommunications - secteur de normalisation des Télécommunications)
- **ISO** (International Organization for Standardization)
- **IETF** (Internet Engineering Task Force)
- **EIA** (Electrical Industries Association)
- **TIA** (Telecommunications Industry Association)

On peut classer les normes WAN en fonction de la couche du modèle OSI correspondante. On obtient donc ceci :

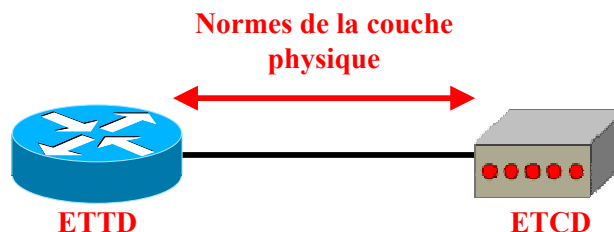
- Normes de la couche physique.
- Normes de la couche liaison de données.

1.3.1. Normes WAN de la couche physique

Les normes WAN de la couche physique décrivent comment fournir des connexions électriques, mécaniques, opérationnelles et fonctionnelles pour les services WAN.

Elles décrivent notamment :

- L'équipement terminal de traitement des données (ETTD, ou DTE en Anglais).
 - o L'ETTD est la partie client d'une liaison WAN. C'est lui qui gère les données.
- L'équipement de terminaison de circuit de données (ETCD, ou DCE en Anglais).
 - o L'ETCD est la partie fournisseur de services de la liaison WAN. Il a pour but d'acheminer les données fournies par l'ETTD.

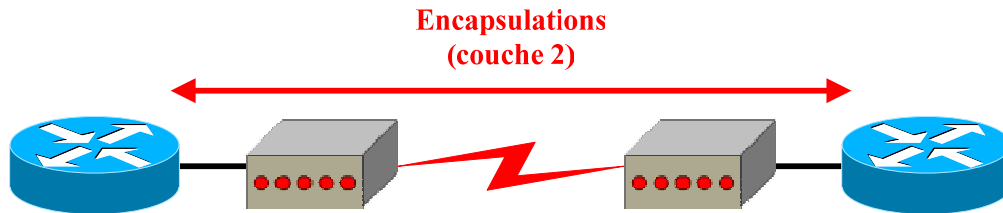


La couche physique d'un réseau WAN décrit principalement l'interface entre l'ETTD et l'ETCD :

- EIA/TIA-232
- EIA/TIA-449
- EIA/TIA-612/613
- V.24
- V.35
- X.21
- G.703

1.3.2. Normes WAN de la couche liaison de données

Les normes WAN de la couche liaison de données décrivent la façon dont les trames sont transportées entre des systèmes par une liaison unique. Elles définissent donc le mode d'encapsulation et les caractéristiques de transmission de ces données.



Les encapsulations les plus couramment utilisées sont les suivantes :

- **HDLC :**
 - o Encapsulation par défaut pour les interfaces WAN d'un routeur Cisco.
 - o Incompatibilité possible entre les différents constructeurs, due aux différences d'implémentation.
 - o Dérivé et remplaçant de SDLC.
- **PPP :**
 - o Comprend un champ identifiant le protocole de couche réseau.
 - o Gestion de l'authentification grâce aux protocoles PAP et CHAP.
 - o Remplace le protocole SLIP due à sa polyvalence.
- **Frame Relay :**
 - o Encapsulation simplifiée, dérivée de LAPB.
 - o Dépourvue de mécanismes de correction d'erreurs.
 - o Prévue pour des unités numériques haut de gamme.
- **LAPB :**
 - o Utilisée sur les réseaux X.25.
- **LAPD :**
 - o Utilisée sur les canaux D des liaisons RNIS.

1.4. Technologies WAN

Les technologies WAN sont classifiées en fonction des catégories suivantes :

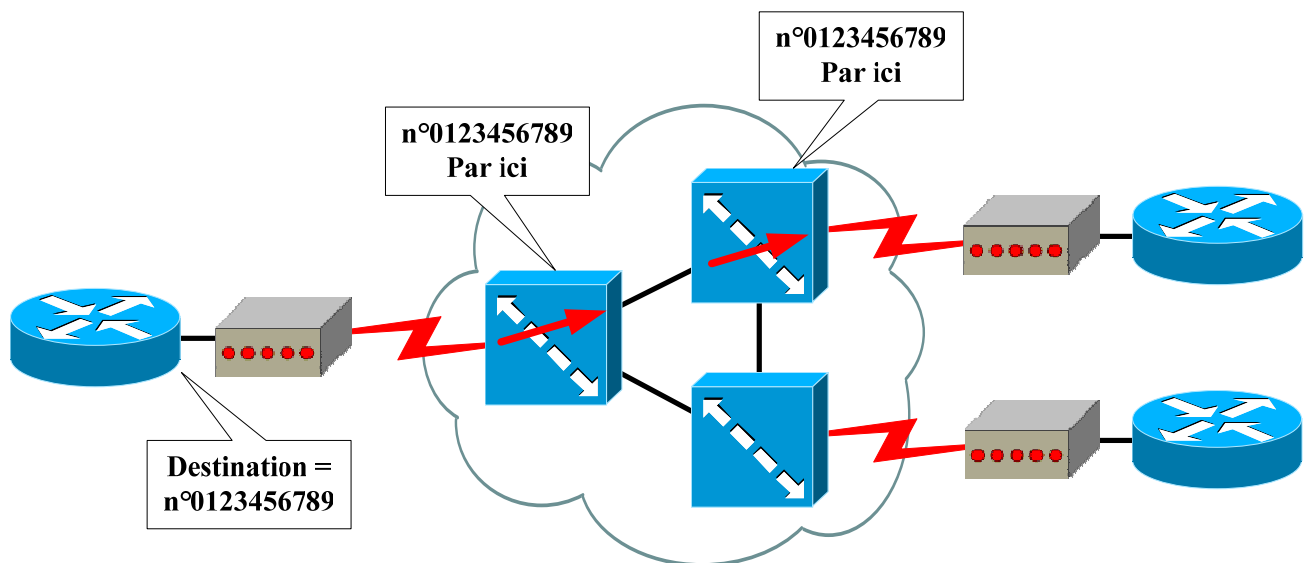
- Services à commutation de circuits.
- Services à commutation de paquets.
- Services à commutation de cellules.
- Services dédiés.
- Autres services.

Les liaisons WAN doivent toujours être des liaisons point-à-point entre les équipements d'extrémité. Ceci peut être obtenu de deux manières :

- Utilisation d'une liaison physique distincte (services à commutation de circuits ou services dédiés).
- Utilisation d'un circuit virtuel au travers d'un environnement commuté (services à commutation de paquets/cellules).

1.4.1. Services à commutation de circuits

Les services à commutation de circuits se servent du réseau téléphonique (analogique ou numérique) pour créer une liaison dédiée non permanente entre la source et la destination.



La liaison est établie grâce à un identifiant, à savoir un numéro de téléphone, pour indiquer au réseau téléphonique la destination avec laquelle on souhaite créer une liaison. Après établissement de l'appel, la liaison dédiée est établie. Il s'agit donc d'une commutation physique des différents centraux téléphoniques.

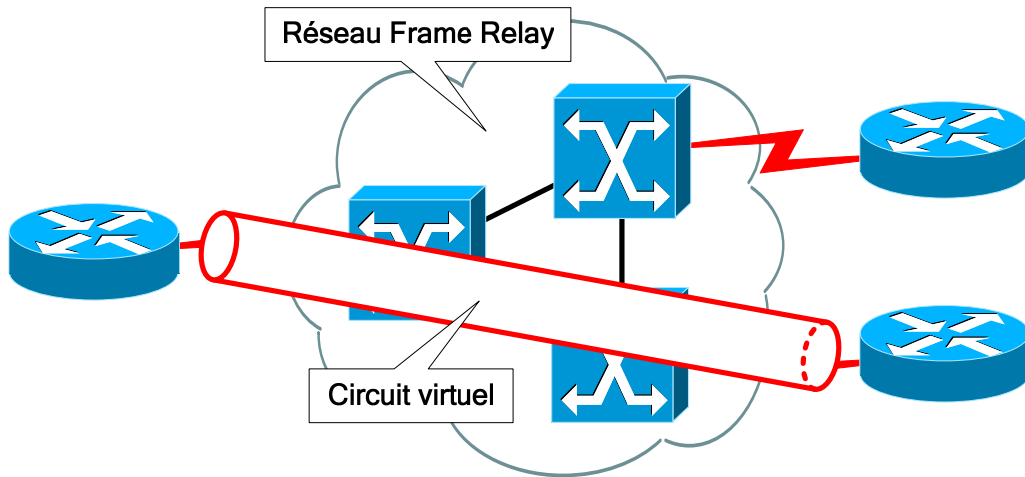
Les technologies basées sur ce type de services offrent la bande passante maximale du lien mais uniquement pour la durée de l'appel.

Les exemples de services à commutation de circuits sont :

- **POTS** (Plain Old Telephone Service).
- **RNIS** (Réseau Numérique à Intégration de Services).

1.4.2. Services à commutation de paquets/cellules

Le principe de base est de fournir une connectivité au travers de commutateurs. On a par conséquent la possibilité d'accéder à toutes les destinations possibles via des liaisons point-à-point ou point-à-multipoint (aussi appelé plus simplement multipoint).



L'utilisation de circuits virtuels par dessus un réseau commuté permet de respecter le principe de connexion point-à-point entre la source et la destination. Le résultat est donc d'avoir un circuit virtuel par destination.

La différence entre les services à commutation de paquets et de cellules est sur la taille des trames ainsi que sur leur traitement :

- Pour la commutation de paquets, les trames ont une taille variable et le traitement est logiciel.
- Pour la commutation de cellules, les trames ont une taille fixe et réduite permettant un traitement matériel.

Les technologies basées sur ces services offrent une bande passante partagée entre les différents trafics de façon permanente.

Les technologies basées sur le service à commutation de paquets sont :

- **Frame Relay**
- **X.25**

Le seul exemple de service à commutation de cellules est :

- **ATM** (Asynchronous Transfer Mode)

1.4.3. Services dédiés

Les services dédiés offrent, comme leur nom l'indique, un lien physique dédié entre chaque source et destination. Le nombre de liens nécessaires s'accroît donc en fonction du nombre de clients à interconnecter.

Parmi les technologies existantes proposant un service dédié, on peut citer les suivants :

- **T1, T3, E1, E3**
- **SDH** (Synchronous Digital Hierarchy)

1.4.4. Autres services

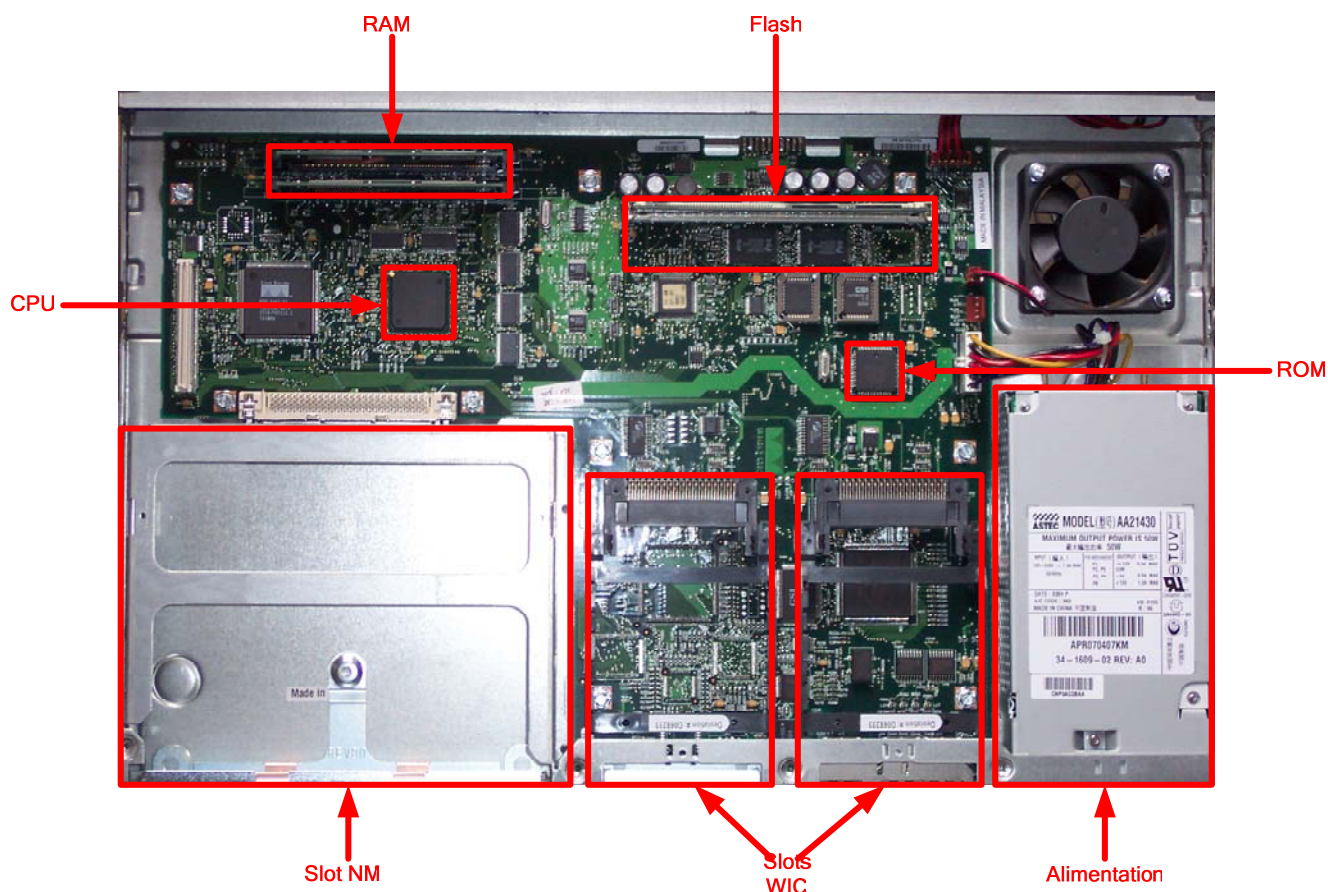
Toutes les technologies non référencées dans les catégories de services précédentes, principalement de nouvelles technologies, sont présentes dans cette dernière :

- **Modem câble**
- **Satellite**
- **Sans fil**

2. Introduction aux routeurs

2.1. Présentation d'un routeur Cisco

2.1.1. Composants internes



Vue interne d'un routeur Cisco 2620XM

La connaissance exacte de l'emplacement de chaque composant interne d'un routeur n'est pas fondamentale. Il peut tout de même être utile de savoir reconnaître les différents slots pour les barrettes de RAM et de Flash au cas où une mise à jour serait à effectuer.

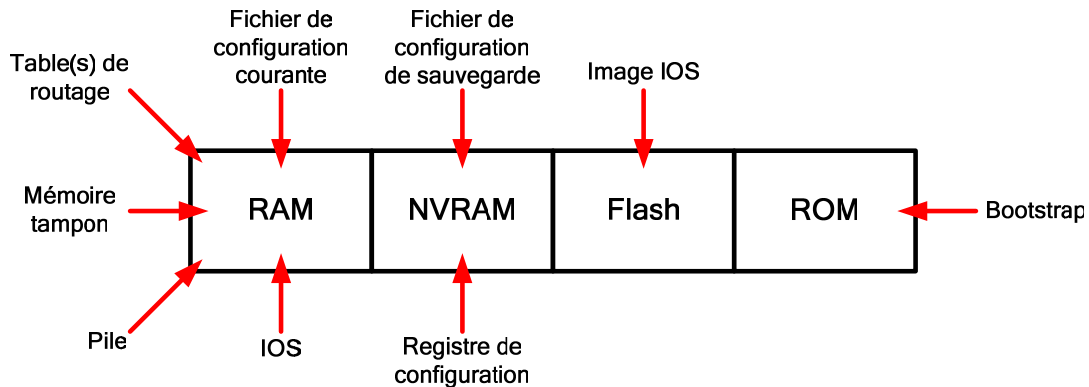
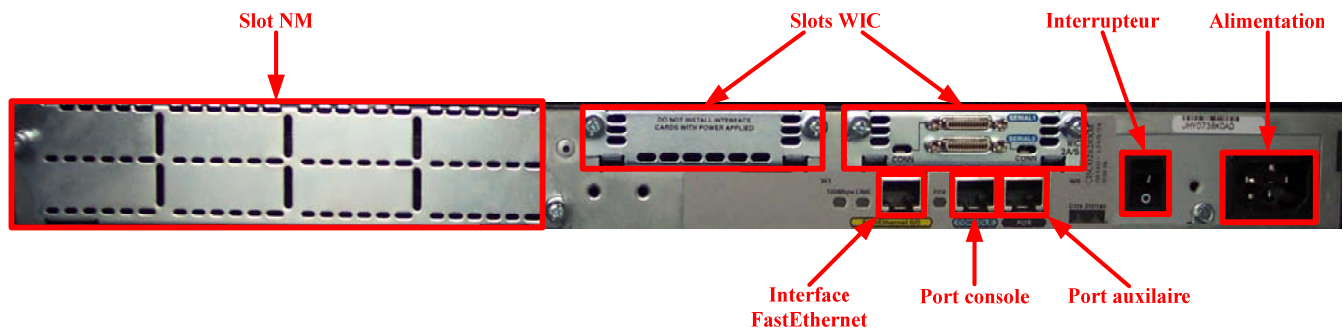


Schéma des mémoires d'un routeur Cisco

Schématiquement, les composants internes qui nous intéressent principalement sont les différentes mémoires utilisées :

- **RAM** : C'est la mémoire principale de travail du routeur. Elle contient entre autres le système d'exploitation une fois chargé, le fichier de configuration active, la ou les tables de routage, ainsi que les mémoires tampon utilisées par les interfaces et la pile utilisée par les processus logiciels. Sa taille varie en fonction du modèle de routeur (64 ou 96 Mo sur un 2620XM). Le contenu de cette mémoire est effacé lors de la mise hors tension ou du redémarrage.
- **NVRAM (Non-Volatile RAM)** : Cette mémoire est non volatile, c'est-à-dire que son contenu n'est pas effacé lorsque l'alimentation est coupée. Sa très petite capacité de stockage (32 Ko sur un 2620XM) ne lui permet pas de stocker autre chose que le registre de configuration et le fichier de configuration de sauvegarde.
- **Flash** : C'est la mémoire de stockage principale du routeur. Elle contient l'image du système d'exploitation Cisco IOS (32 Mo sur un 2620XM). Son contenu est conservé lors de la mise hors tension et du redémarrage.
- **ROM** : Elle contient le bootstrap ainsi que la séquence d'amorçage du routeur. Celle-ci est donc uniquement utilisée au démarrage du routeur.

2.1.2. Composants externes



Vue arrière d'un routeur Cisco 2620XM

Un routeur Cisco peut offrir plusieurs types de connectiques parmi les suivantes :

- **Port console** : Accès de base pour configuration.
- **Port auxiliaire** : Accès pour configuration au travers d'une ligne analogique et modems interposés.
- **Interface(s) LAN**
- **Interface(s) WAN**
- **Slot(s) NM** (Network Module)
- **Slot(s) WIC** (WAN Interface Card)

2.2. Branchements

2.2.1. Interfaces LAN et WAN

Les interfaces réseaux fournies par un routeur Cisco peuvent être de divers types et sont classifiées en fonction du type de réseau à connecter (LAN ou WAN).

Elles peuvent être fixées au châssis ou livrées sous la forme de cartes (WIC ou NM) pour les routeurs modulaires.

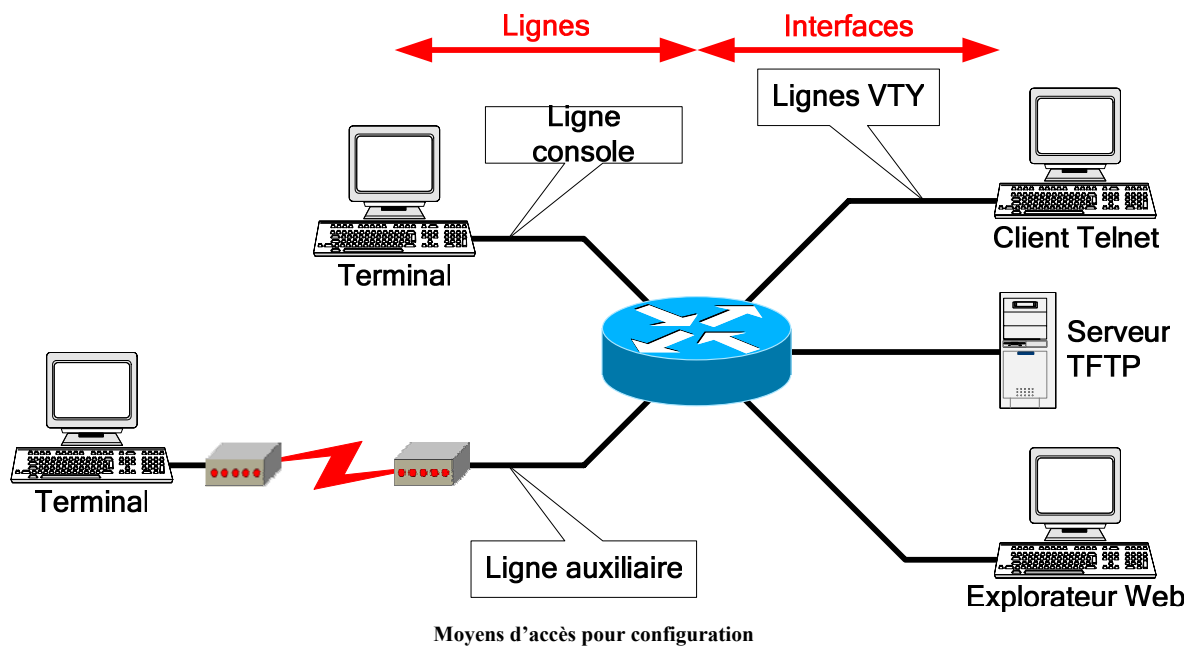
Ces interfaces seront utilisées par les protocoles de couche 3 du modèle OSI pour le routage.



Carte WIC-2A/S

2.2.2. Accès pour configuration

La configuration d'un routeur se fait par l'intermédiaire de lignes.

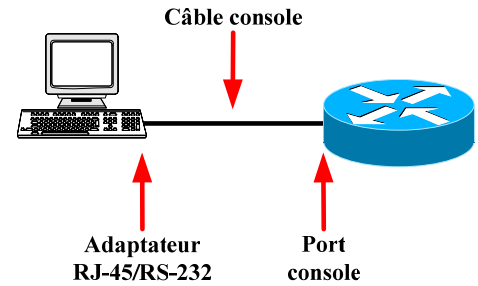


Un routeur peut être configuré à partir des sources externes suivantes :

- **Ligne console** : Accès primaire, à utiliser si aucun autre accès de configuration n'est disponible.
- **Ligne auxiliaire** : Accès à distance via une liaison RTC et modems interposés.
- **Ligne(s) VTY** : Accès via un client Telnet (5 ou 16 lignes disponibles par routeur en fonction du modèle).
- **Explorateur Web** : Accès utilisant le serveur HTTP interne du routeur.
- **Serveur TFTP** : Import/export de fichiers de configuration.
- **Serveur FTP** : Import/export de fichiers de configuration.

La ligne console est l'accès de configuration à utiliser lorsque aucune configuration n'est chargée ou si cette dernière ne permet pas l'accès par un autre moyen (Telnet, etc.).

Il faut connecter le port console du routeur à un port série (RS-232) en utilisant un câble console (rollover).

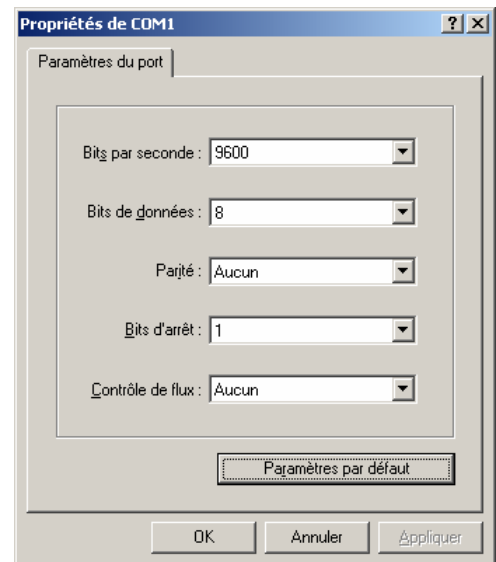


Un émulateur de terminaux (exemple : HyperTerminal sous Windows) permet l'accès à l'interface de configuration du routeur.

Les paramètres à utiliser sont les suivants :

- Vitesse : 9600 bauds
- Bits de données : 8
- Parité : Aucun
- Bits d'arrêt : 1
- Contrôle de flux : Aucun

Sous HyperTerminal, le bouton "Paramètres par défaut" permet de spécifier automatiquement ces paramètres.



Paramètres de connexion pour HyperTerminal

2.3. Système d'exploitation Cisco IOS

2.3.1. Principes et spécifications

IOS (Internetwork Operating System) est le système d'exploitation propriétaire Cisco utilisé sur la plupart des dispositifs Cisco. Ce système d'exploitation offre une CLI (Command Line Interface).

Le programme d'exécution des commandes, ou EXEC, est l'un des composants de la plateforme logicielle Cisco IOS. EXEC reçoit et exécute les commandes entrées dans la CLI.

Pour arrêter l'exécution d'une commande, il faut utiliser une des combinaisons de touches suivantes :

- **CTRL+MAJ+6**
 - Pour toutes les commandes.
- **CTRL+C**
 - Fonctionne avec les commandes **show** et pour le mode SETUP.

EXEC transmet des messages de notification sur le terminal ainsi que les messages de déboguage. Par défaut, ces messages arrivent uniquement sur le terminal connecté via la ligne console. Pour activer ou désactiver l'affichage de ces messages, il faut utiliser la commande **terminal [no] monitor** depuis le mode utilisateur ou privilégié.

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

2.3.2. Modes de commandes

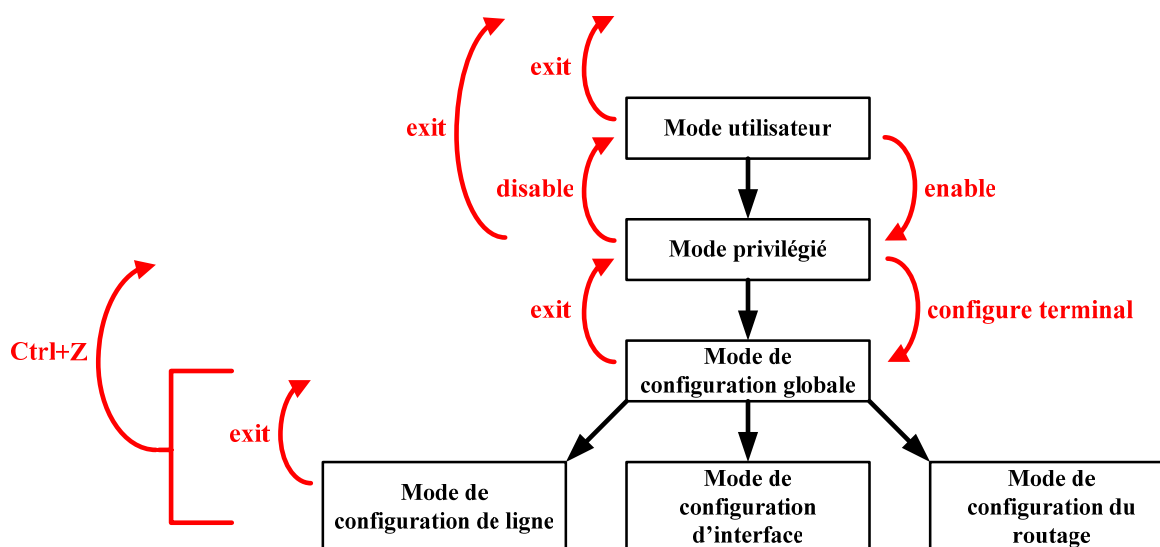
Il existe une multitude de modes différents accessibles en CLI sur un routeur Cisco :

- **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 ne dispose que de commandes de visualisation d'état sur le 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 que dans chaque mode spécifique des commandes ayant une portée localisée au composant du routeur spécifié par ce mode.
- **Mode SETUP** : Mode affichant un dialogue interactif, 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 dans lequel on est en repérant l'invite de commande que nous fournit l'interpréteur de commandes EXEC :

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

Nous allons maintenant voir les commandes et les combinaisons de touches permettant de naviguer dans ces différents modes d'IOS :



Hiérarchie et navigation dans les modes d'IOS

Les commandes à utiliser pour passer dans un mode de configuration spécifique sont les suivantes :

- **line {type} {numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'une ligne
- **interface {type} {numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'interface
- **router {protocole} [option]**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration du routeur

Pour les lignes et les interfaces, la numérotation commence à 0.

2.3.3. Système d'aide

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

- **Le caractère ?** : Ce caractère peut être utilisé de 3 façons différentes. Seul, ce caractère indique au routeur de nous fournir une liste complète des commandes accessibles depuis le mode dans lequel on se trouve. Collé à une chaîne de caractères, il fournit la liste des mots clé commençant par cette chaîne. Enfin, après un mot clé, il fournit la liste des options pour ce dernier.
- **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.

2.3.4. Commandes d'édition avancée

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
ECHAP+B	Reculé d'un mot
CTRL+B ou ←	Reculé d'un caractère
CTRL+E	Va à la fin de la ligne de commande
CTRL+F ou →	Avance d'un caractère
ECHAP+F	Avance d'un mot
terminal [no] editing	Active/désactive les commandes d'édition avancée

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 ou en fin de ligne d'écran lorsque la commande en elle-même fait plus d'une ligne écran. Ceci indique donc qu'une partie de la ligne de commande est masquée.

2.3.5. Historique des commandes

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 ↑	Rappel de la commande précédente
CTRL+N ou ↓	Rappel de la commande suivante
show history	Affiche la liste des commandes en mémoire
terminal history size {taille}	Définit la taille de la mémoire de commandes (valeur maximale de 256)
terminal [no] history	Active/désactive les fonctions d'historique

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

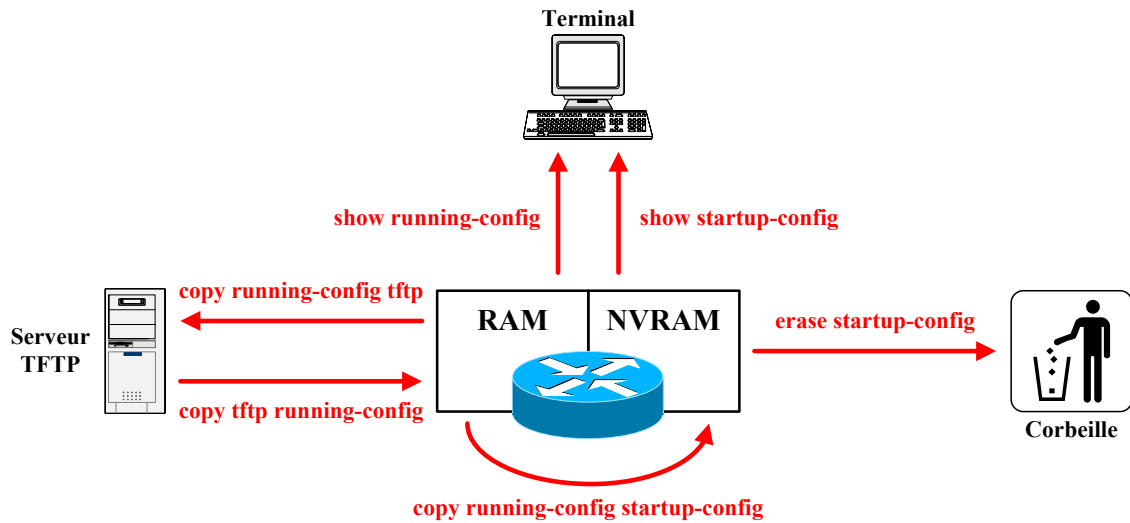
2.3.6. Fichiers de configuration

Les deux fichiers de configuration d'un routeur Cisco sont les fichiers de configuration active (dans la RAM) et de sauvegarde (dans la NVRAM). Ils régissent respectivement la configuration en cours d'utilisation par le routeur et la configuration utilisée lors du démarrage du routeur.

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

- 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 noms.
- Chaque interface avec sa configuration spécifique.
- Toutes les informations de routage.
- Chaque ligne et sa configuration spécifique.

Les différentes commandes (IOS >= 11) associées aux fichiers de configuration sont les suivantes :



- **show running-config** : Affiche la configuration courante
- **show startup-config** : Affiche la configuration de sauvegarde
- **copy running-config startup-config** : Sauvegarde la configuration courante dans la NVRAM
- **copy running-config tftp** : Exporte la configuration courante vers un serveur TFTP
- **copy tftp running-config** : Importe une configuration dans la RAM depuis un serveur TFTP
- **copy startup-config tftp** : Exporte la configuration de sauvegarde vers un serveur TFTP
- **copy tftp startup-config** : Importe une configuration dans la NVRAM depuis un serveur TFTP
- **erase startup-config** : Supprime le fichier de configuration de sauvegarde

3. Configuration de base d'un routeur

3.1. Commandes de visualisation d'état

IOS propose une panoplie importante de commandes permettant la visualisation de l'état. Ces commandes commencent toutes par le mot clé **show**. Les commandes de visualisation d'état à connaître en premier lieu sont les suivantes :

- **show running-config** : Affiche le fichier de la configuration active.
- **show startup-config** : Affiche le fichier de la configuration de sauvegarde.
- **show version** : Affiche la configuration matérielle système, la version d'IOS, le nom et la source de l'image IOS d'amorçage, ainsi que la valeur du registre de configuration.
- **show processes** : Affiche des informations sur les processus actifs.
- **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.
- **show buffers** : Fournit des statistiques sur les mémoires tampon des interfaces du routeur.
- **show arp** : Affiche les entrées ARP connues.
- **clear arp** : Vide les entrées dynamiques de la table ARP.
- **show hosts** : Affiche la table de résolution de noms.
- **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 interfaces [{type} {numéro}]** : Affiche les informations de configuration ainsi que des statistiques de trafic pour chaque interface configurée sur le routeur (couches 2 et 3).
- **show controllers [{type} {numéro}]** : Affiche les informations de couche 1 des interfaces.
- **show ip interface [{type} {numéro}] [brief]** : Affiche les informations IP pour les interfaces
- **clear counters [{type} {numéro}]** : Permet de mettre à zéro toutes les statistiques des interfaces du routeur.
- **show ip route** : Affiche la table de routage IP.
- **show protocols** : Affiche le nom et l'état de tous les protocoles configurés de couche 3.
- **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.
- **show sessions** : Affiche la liste des sessions en cours.
- **show users** : Affiche la liste des utilisateurs actuellement connectés au routeur.
- **show clock** : Affiche la date et l'heure actuelle.
- **show history** : Affiche la liste des commandes en mémoire.

3.2. Date et heure

Comme pour tout système informatique, la date et l'heure peuvent être configurés. Ceci peut s'avérer utile lorsque l'on utilise les fonctions de log ou de débogage, en fournissant la date et l'heure exacte des événements survenus.

Les commandes utilisées pour le système de date et d'heure sont les suivantes :

- **show clock**
 - Affiche la date et l'heure du système
- **clock set {hh:mm:ss} {jour} {mois} {année}**
 - Mode privilégié
 - Permet de configurer l'heure sur le routeur
 - **hh:mm:ss** correspond à l'heure (de 0 à 23), aux minutes et aux secondes.
 - **jour** est un nombre (de 1 à 31).
 - **mois** est le nom du mois.
 - **année** est l'année avec 4 chiffres.

Cette configuration est manuelle, et est nécessaire à chaque redémarrage du routeur. Il est possible d'utiliser le protocole NTP (Network Time Protocol), afin de maintenir synchronisé le routeur avec un serveur de temps.

3.3. Nom d'hôte et résolution de noms

Les noms d'hôtes sont très utiles. En effet, ils permettent d'identifier un hôte avec un nom facile à retenir plutôt que d'utiliser des adresses de couches réseau. Pour pouvoir utiliser ces noms d'hôtes, il faut un système de résolution de noms, sachant que cette résolution de noms a une portée locale.

Les noms d'hôtes et les résolutions de noms ne sont pas transmis de routeur à routeur. Cela signifie qu'il faut configurer la résolution de noms sur tous les dispositifs réseau sur lesquels on souhaite utiliser des noms d'hôtes pour la communication réseau.

Il est possible de configurer :

- Le nom d'hôte du routeur
- La résolution de noms statique
- La résolution de noms dynamique grâce au protocole DNS

Les commandes à utiliser sont les suivantes :

- **hostname {nom}**
 - o Mode de configuration globale
 - o Attribution du nom d'hôte du routeur
 - o Ce nom est affiché par l'invite de commandes
 - o La valeur par défaut est "Router"
- **ip host {nom} [tcp_port_number] {IP1} [{IP2}...]**
 - o Mode de configuration globale
 - o Création d'une entrée statique de résolution de noms dans la table d'hôtes
 - o **tcp_port_number** permet de spécifier le port TCP à utiliser avec cet hôte pour un accès Telnet
 - o Il est possible de spécifier plusieurs adresses IP pour un seul hôte. Dans ce cas, seule la commande **telnet** utilisera les adresses autres que la première si les précédentes ne répondent pas
- **[no] ip domain-lookup**
 - o Mode de configuration globale
 - o Active/désactive la résolution dynamique de noms (DNS)
- **ip name-server {DNS1} [{DNS2}...]**
 - o Mode de configuration globale
 - o Permet de spécifier le ou les serveurs DNS avec lesquels nous effectuerons les résolutions d'adresses
 - o On peut préciser jusqu'à 6 serveurs DNS différents
- **ip domain-name {préfixe}**
 - o Mode de configuration globale
 - o Précise le préfixe DNS par défaut à utiliser pour la résolution d'adresses dynamique

La commande **show hosts** permet d'afficher la table des correspondances entre les noms d'hôte et leur(s) adresse(s) de couche 3. Les champs de cette table sont les suivants :

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 valide
EX	Entrée obsolète, expirée
Age	Temps (en heures) écoulé depuis que le logiciel a appris l'entrée
Type	Champ identifiant le protocole de couche 3
Address(es)	Adresses logiques associées au nom de machine

3.4. Descriptions et bannière de connexion

Les descriptions d'interface et la bannière de connexion sont très utiles pour fournir des informations quant à l'utilité ou la fonction de chaque routeur et de chacune de ces interfaces.

La bannière de connexion s'affiche lors de la connexion d'un terminal à une ligne et permet de transmettre un message aux utilisateurs du routeur. Ceci peut être utile pour les avertir d'un arrêt imminent du routeur ou pour faire passer un message publicitaire.

Pour définir cette bannière, il faut utiliser la commande :

- **banner motd {caractère d'encapsulation} {message} {caractère d'encapsulation}**
 - o Mode de configuration globale
 - o Le message doit être encapsulé entre un caractère quelconque qui ne doit pas exister dans le 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 une interface. Pour cela, il faut utiliser la commande :

- **description {texte}**
 - o Mode de configuration d'interface
 - o Le texte de description ne peut pas excéder 80 caractères sur les anciens modèles (exemple : Routeur 2500) ou 240 caractères sur les modèles plus récent (exemple : Routeur 2600).
 - o Cette description est visible en utilisant la commande **show interfaces**.

3.5. 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 ainsi que sur l'accès au mode privilégié.

Pour configurer une protection par mot de passe sur une ligne, il faut utiliser les commandes suivantes :

- **line {console | aux | vty} {{numéro} | {premier} {dernier}}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration de la ou des lignes voulues
 - o Il est possible d'accéder à plusieurs lignes en même temps. Pour cela, il suffit de préciser non pas le numéro mais la plage de numéros. Par exemple, pour accéder directement dans le mode de configuration des 5 lignes VTY, il suffit d'utiliser la commande **line vty 0 4**
- **password {mot de passe}**
 - o Mode de configuration de ligne
 - o Permet de spécifier le mot de passe pour la ligne courante
 - o Le mot de passe est écrit par défaut en clair dans le fichier de configuration
- **login**
 - o Mode de configuration de ligne
 - o Précise qu'aucun login ne sera demandé lors de la connexion
 - o Cette commande ne peut être utilisée que si un mot de passe est déjà configuré sur la ligne.

Les mots de passe pour les lignes console et auxiliaire ne sont pris en compte qu'après le redémarrage du routeur. Les lignes auxiliaire et VTY ne sont pas opérationnelles si elles n'ont pas de mot de passe configuré. Cela signifie qu'aucun accès autre que par la ligne console n'est faisable sans configuration préalable.

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

- **enable password {mot de passe}**
 - o Mode de configuration globale
 - o Le mot de passe est écrit en clair dans le fichier de configuration
- **enable secret {mot de passe}**
 - o Mode de configuration globale
 - o Le mot de passe est crypté dans le fichier de configuration en utilisant l'algorithme MD5.
 - o Cette commande est prioritaire par rapport à **enable password** si elles sont toutes deux configurées

Malheureusement, tous les mots de passe, à l'exception du **enable secret**, sont écrits en clair dans le fichier de configuration. Ceci implique une plausible faille de sécurité (sauvegarde d'un fichier de configuration sur un serveur TFTP non sécurisé, etc.).

Pour y remédier, il faut utiliser la commande **service password-encryption** depuis le mode de configuration globale. Cette commande permet de crypter tous les mots de passe écrits en clair dans le fichier de configuration en utilisant un algorithme propriétaire Cisco.

3.6. Serveur HTTP

IOS fournit un serveur HTTP. Ce serveur fournit un moyen d'accès pour configuration.

La commande à utiliser pour contrôler l'état de ce serveur HTTP est :

- **[no] ip http server**
 - o Mode de configuration globale
 - o Active/désactive le serveur HTTP interne du routeur
 - o Actif par défaut

Pour accéder au service HTTP fourni par le routeur, il faut utiliser un explorateur Web et y accéder en indiquant l'adresse IP d'une interface.

Lors de la connexion, la page Web demande un nom d'utilisateur et un mot de passe. Les valeurs par défaut ne correspondent à aucun nom d'utilisateur et au mot de passe du mode privilégié.

Ce serveur HTTP faisant l'objet de beaucoup d'exploits et de failles de sécurité, il est recommandé de le désactiver lorsque l'on n'en a plus/pas besoin.

3.7. Configuration des interfaces

Les interfaces sont utilisées par les routeurs pour transférer les paquets de données entre différents réseaux de couche 3.

Ces interfaces peuvent être de différents types. Dans ce cours, nous étudierons uniquement les interfaces suivantes :

- Loopback
- Ethernet
- Serial

La commande **show interfaces** permet l'affichage de l'état des interfaces du routeur. On peut déterminer :

- L'adresse IP et le masque de sous-réseau
- L'adresse de couche 2
- L'encapsulation utilisée
- Les statistiques sur le trafic transitant au travers de chaque interface
- L'état de l'interface, qui correspond à ceci :

Interface (couche 1)	Line protocol (couche 2)
Administratively down (shutdown)	Down (problème de couche 2)
Down (problème de câble)	Up (réception des "keepalive")
Up (média fonctionnel)	

3.7.1. Interfaces Loopback

Les interfaces Loopback sont généralement utilisées pour simuler des interfaces réelles.

Pour leur configuration, on dispose des commandes suivantes :

- **interface loopback {numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'interface
- **ip address {IP} {masque} [secondary]**
 - o Mode de configuration d'interface
 - o Permet d'attribuer une adresse IP à cette interface
 - o Le paramètre **secondary** précise qu'il s'agit d'une adresse IP secondaire

3.7.2. Interfaces Ethernet/IEEE 802.3

Les interfaces de type Ethernet/IEEE 802.3 peuvent être de type :

- Ethernet (IEEE 802.3)
- Fast Ethernet (IEEE 802.3u)
- Gigabit Ethernet (IEEE 802.3ab ou IEEE 802.3z)
- 10-Gigabit Ethernet (IEEE 802.3ae)

Les interfaces Gigabit ou 10-Gigabit ne seront pas étudiées dans ce cours.

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

- **interface {Ethernet | FastEthernet} {numéro | slot/numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'interface
- **ip address {IP} {masque} [secondary]**
 - o Mode de configuration d'interface
 - o Permet d'attribuer une adresse IP à cette interface
 - o Le paramètre **secondary** précise qu'il s'agit d'une adresse IP secondaire
- **[no] keepalive**
 - o Mode de configuration d'interface
 - o Active/désactive les "keep alive" sur l'interface
 - o Utile pour rendre une interface opérationnelle sans avoir à brancher un média
- **[no] shutdown**
 - o Mode de configuration d'interface
 - o Active/désactive administrativement l'interface

3.7.3. Interfaces série

Les interfaces série sont classifiées en fonction de leur mode de transmission qui peut être :

- Synchrone
- Asynchrone
- Synchrone/asynchrone (par défaut en mode synchrone)

Elles sont le plus souvent présentes sous la forme de cartes WIC à insérer dans des slots de routeurs modulaires.

Les commandes utilisées par ces interfaces sont les suivantes :

- **interface {serial | async} {numéro | slot/numéro}**
 - o Mode de configuration globale
 - o Permet de passer dans le mode de configuration d'interface
 - o Le mot clé **async** n'est utilisable que pour les interfaces de type asynchrone
- **clock rate {vitesse}**
 - o Mode de configuration d'interface
 - o Spécifie la vitesse de fonctionnement de la liaison WAN
 - o A faire uniquement sur une interface ETCD
 - o Le paramètre **vitesse** est exprimé en bits par seconde
- **ip address {IP} {masque} [secondary]**
 - o Mode de configuration d'interface
 - o Permet d'attribuer une adresse IP à cette interface
 - o Le paramètre **secondary** précise qu'il s'agit d'une adresse IP secondaire
- **[no] shutdown**
 - o Mode de configuration d'interface
 - o Active/désactive administrativement l'interface

4. Informations et accès aux autres dispositifs

4.1. 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.

4.1.1. 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.)

4.1.2. 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

4.1.3. 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

4.2. Telnet

4.2.1. Théorie

Telnet est un protocole faisant partie intégrante de la pile de protocole TCP/IP et fonctionnant au niveau de la couche application du modèle OSI. Il offre un moyen d'accès distant aux dispositifs réseaux sous la forme d'un terminal virtuel (VTY).

La communication réseau s'effectue à l'aide du protocole TCP sur le port 23.

Telnet est utilisé à la fois pour l'accès distant pour configuration ainsi qu'à des fins de tests et de résolution de problèmes. Ce dernier point sera étudié dans le chapitre correspondant.

4.2.2. Commandes et utilisation

L'accès Telnet s'effectue au travers d'une ligne VTY. Un tel accès est par conséquent possible que si au moins une ligne VTY est correctement configurée et libre d'accès.

Pour rappel, chaque routeur Cisco dispose d'un total de 5 ou 16 lignes VTY (dépend du modèle et de l'IOS).

Les commandes et combinaisons de touches liées à l'utilisation de Telnet sont les suivantes :

telnet {IP nom} [tcp_port_number]	Etablir une session Telnet avec l'hôte correspondant à l'IP ou au nom précisé (tcp_port_number permet d'explicitier le numéro de port TCP à utiliser)
connect {IP nom}	Identique à telnet
{IP nom}	Identique à telnet
exit	Fermeture de la session Telnet avec déconnexion (déconnecté par défaut après 10 minutes d'inactivité)
disconnect	Identique à exit
CTRL+MAJ+6 puis X	Suspendre la session Telnet en cours et la mettre en tâche de fond (reprise avec la touche Entrée si une seule session est en tâche de fond, sinon utiliser la commande resume)
show sessions	Afficher la liste des sessions en cours
resume {numéro}	Reprend la session Telnet précisée (numéro correspond à celui fournit par la commande show sessions)

La combinaison de touches **CTRL+MAJ+6** ne fonctionne qu'avec un clavier QWERTY.

On peut observer qu'une erreur dans l'écriture d'une ligne de commande quelconque depuis le mode privilégié pourrait faire croire à IOS que l'on tente d'établir une session Telnet vers un hôte ayant pour nom notre commande erronée.

Cela aurait pour impact de lancer une résolution DNS, qui pourrait durer jusqu'à expiration du timeout, pour obtenir l'adresse IP de cet hôte fictif. L'une des solutions pour remédier à ce problème est de désactiver le service DNS sur le routeur si on ne l'utilise pas.

5. Gestion d'IOS et processus de démarrage

5.1. Processus de démarrage

Le processus de démarrage d'un routeur Cisco est important à connaître, malgré le fait que l'on ne fasse pas de modifications sur ce processus à long terme. Cela devient en revanche primordial lorsqu'il faut mettre à jour l'image d'IOS actuellement en place sur le routeur ou lorsqu'un problème survient.

Cette partie portera sur :

- **La séquence d'amorçage** : Quelles sont les étapes de l'amorçage d'un routeur Cisco ?
- **Les commandes boot system** : Où le routeur peut trouver une image d'IOS ?
- **Le registre de configuration** : Comment doit démarrer le routeur ?

5.1.1. Séquence d'amorçage

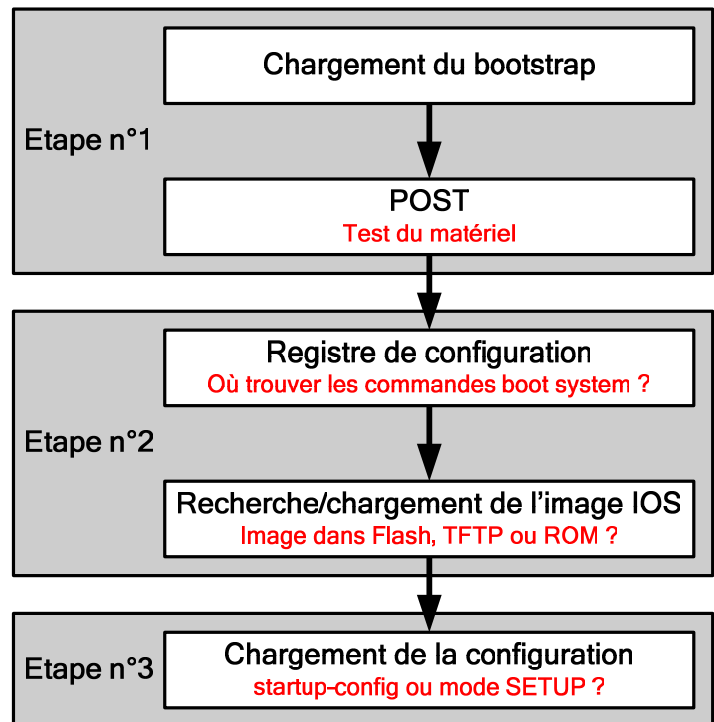
La séquence d'amorçage d'un routeur est découpée en 3 étapes :

- Etape n°1 : POST (Power On Self Test)
- Etape n°2 : Chargement d'IOS
- Etape n°3 : Chargement de la configuration

L'étape n°1 se résume au chargement du bootstrap, microcode contenu dans la ROM du routeur, qui va se charger de tester le matériel.

L'étape n°2 consiste à trouver une image d'IOS fonctionnelle afin de la charger en RAM. Ceci se fait en 2 phases.

La première phase consiste à analyser la valeur du registre de configuration, afin de déterminer si le routeur doit utiliser la séquence de recherche d'image IOS par défaut ou celle précisée dans le fichier de configuration de sauvegarde.



La deuxième phase correspond à la recherche de l'image d'IOS à proprement parler, en utilisant ces séquences de recherche. Si la séquence de recherche d'image IOS précisée dans le fichier de configuration de sauvegarde ne permet pas de trouver une image valide ou si elle est ignorée, le routeur tentera de démarrer en utilisant la première image présente en Flash.

Si aucune image IOS n'a pu être trouvée, le démarrage du routeur s'arrêtera au mode RXBoot.

L'étape n°3 consiste à charger une configuration. Par défaut, le routeur importera le fichier de configuration de sauvegarde dans le fichier de configuration courante.

Si le fichier de configuration de sauvegarde n'est pas chargé, car inexistant ou ignoré, la configuration initiale est chargée et le mode SETUP est automatiquement lancé, afin de procéder à la configuration basique du routeur.

5.1.2. Commandes boot system

Les commandes **boot system**, aussi appelées options bootstrap, servent à indiquer au routeur l'emplacement d'une image IOS et peuvent désigner trois types d'emplacements différents :

- **Flash** : C'est l'espace de stockage standard pour les images IOS.
- **TFTP** : Le serveur TFTP permet de stocker de nombreux fichiers. Il est généralement utilisé à des fins de mise à jour du système.
- **ROM** : Sur les anciens routeurs, tels que les Cisco 2500, la ROM contenait une image IOS minimaliste. Celle-ci était utilisée comme solution de secours. Sur les nouveaux routeurs, ceci n'est plus utile car le mode RXBoot est beaucoup plus performant et permet à lui seul de récupérer une image IOS depuis un serveur TFTP.
- **boot system flash {nom du fichier}**
 - o Mode de configuration globale
 - o Permet de spécifier le nom du fichier dans la Flash contenant l'image IOS
- **boot system tftp {nom du fichier} {IP du serveur TFTP}**
 - o Mode de configuration globale
 - o Précise le nom du fichier ainsi que l'adresse IP du serveur TFTP stockant l'image IOS
- **boot system rom**
 - o Mode de configuration globale

Les commandes ci-dessus permettent donc de préciser l'emplacement ainsi que l'ordre de recherche de l'image IOS pour la séquence d'amorçage. L'emplacement est explicité par la commande elle-même, alors que l'ordre de recherche est défini par l'ordre dans lequel on a entré les commandes.

5.1.3. Registre de configuration

Le registre de configuration est un registre de 16 bits qui se trouve dans la mémoire NVRAM. Sa valeur est exprimée en hexadécimal et sa valeur par défaut est 0x2102. 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	Passer par le mode moniteur de mémoire ROM et attendre que l'utilisateur tape la commande b ou boot pour démarrer
0x---1	Démarrer avec la première image présente en Flash ou en utilisant l'image minimaliste présente en ROM (anciens routeurs)
0x---2 à 0x---F	Demander d'utiliser les commandes boot system présentes dans la configuration de sauvegarde. Si aucune commande boot system ne permet d'atteindre une image IOS valide, le routeur tentera de démarrer avec la première image disponible en Flash.

Les 12 autres bits du registre de configuration ont une signification qui ne sera pas étudiée dans ce cours. Il faudra par conséquent toujours garder la valeur par défaut sauf si l'on en connaît l'effet.

Les commandes liées au registre de configuration sont :

- **config-register {valeur}**
 - o Mode de configuration globale
 - o Permet de modifier la valeur du registre de configuration
 - o Le paramètre **valeur** est exprimé en hexadécimal (préfixe 0x)
 - o Toute modification de la valeur est prise en compte lors du redémarrage
- **show version**
 - o Affiche la valeur du registre de configuration

5.1.4. Mode SETUP

Le mode SETUP, aussi connu sous le nom de dialogue de configuration initiale ou interactive, constitue une des routines de la configuration initiale. L'objectif principal du mode SETUP est de créer rapidement une configuration minimale, à savoir de définir :

- Le nom d'hôte
- Le mot de passe du mode privilégié
- Le mot de passe des lignes VTY
- La configuration basique du protocole SNMP
- L'adresse IP pour une interface

Les options configurables via le mode SETUP peuvent varier en fonction de la version d'IOS utilisée.

Le mode SETUP peut être lancé manuellement grâce à la commande **setup** depuis le mode privilégié ou être lancé automatiquement si la configuration de sauvegarde n'est pas chargée (car elle n'existe pas ou a été ignorée).

Il se présente sous la forme d'un questionnaire interactif en mode texte, dans lequel il suffit de répondre aux questions posées par le système. Pour la plupart des questions, 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. Par contre, si on configure le système pour la première fois, il s'agit des valeurs par défaut définies en usine.

Si on ne souhaite plus continuer avec le mode SETUP, on a la possibilité d'utiliser la combinaison de touches **CTRL+C**. Ceci est utile lorsque l'on ne souhaite pas utiliser le mode SETUP pour la configuration basique, ou si une erreur a été commise sur une des réponses. Dans ce dernier cas, il suffit de relancer le mode SETUP pour reprendre le dialogue de configuration à son point de départ.

Lorsque le questionnaire est terminé, la configuration créée est affichée. Le système nous demande alors si l'on souhaite appliquer cette configuration, et par conséquent la sauvegarder dans la NVRAM.

5.2. Gestion d'IOS

La gestion des images IOS n'est pas compliquée. Il suffit d'avoir quelques informations utiles, comme la convention de noms utilisée pour nommer les fichiers, ainsi que les procédures simples de mise à jour du système quelque soit la situation.

5.2.1. Informations générales

Avec l'arrivée des IOS 12.x, une interface unique est maintenant utilisée pour les différents systèmes de fichiers, et se nomme IFS (IOS File System). IFS permet d'accéder aux différents systèmes de fichiers avec une syntaxe uniformisée.

Cette syntaxe se présente sous la forme suivante : **{préfixe}:{répertoire(s)}/{nom du fichier}**

Les préfixes peuvent être :

Préfixe	Signification
flash	Mémoire Flash du routeur
nvrn	Mémoire NVRAM du routeur
system	Mémoire RAM du routeur
ftp	Serveur réseau utilisant le protocole FTP
tftp	Serveur réseau utilisant le protocole TFTP

Il existe beaucoup d'autres préfixes qui ne seront pas étudiés dans ce cours.

Il n'existe pas une seule, mais une multitude de versions d'IOS. C'est pourquoi une convention de noms est définie afin de fournir toutes les informations sur l'image concernée.

Cette convention de noms est la suivante : **{Plateforme}-{Feature Set}-{Format}.{Version}.bin**

- **Plateforme** est le matériel sur lequel l'image est prévue pour fonctionner (exemple : **c2600** pour un routeur de la gamme Cisco 2600).
- **Feature Set** correspond à l'ensemble des fonctionnalités incluses dans l'image (exemple : **js** pour une image de type "Entreprise Plus" et **k9** pour un niveau d'encryption).
- **Format** permet de connaître le format de conditionnement de l'image (exemple : **mz** pour une image compressée).
- **Version** est le numéro de version de l'image IOS (exemple : 123-9 pour un IOS version "12.3(9)").

5.2.2. Gestion des systèmes de fichiers

La gestion des systèmes de fichiers, et plus particulièrement les images IOS ainsi que les fichiers de configuration, passe par l'utilisation de la commande **copy {source} {destination}**.

La source et la destination peuvent être simplement des mots clés (tftp, running-config, startup-config) ou peuvent être exprimées en utilisant la syntaxe uniformisée d'IFS. On peut regrouper l'utilisation de la commande copy en 2 catégories :

- **Import**
 - Source externe (FTP, TFTP) vers une destination interne au routeur (Flash, NVRAM, RAM)
 - Utilisé pour la mise à jour du système
- **Export**
 - Source interne vers une destination externe
 - Utilisé pour la sauvegarde des données

Au travers de cette commande **copy**, on peut donc effectuer une opération importante, à savoir la mise à jour de l'image IOS.

Pour cette opération, il faut donc prendre quelques précautions préliminaires, et procéder comme suit :

- Etape n°1 : Vérifier si la quantité de mémoire Flash disponible est suffisante pour une nouvelle image IOS
- Etape n°2 : Sauvegarder l'image IOS actuelle sur un serveur TFTP ou FTP
- Etape n°3 : Lancer la mise à jour à l'aide de la commande **copy**
- Etape n°4 : Vérification de la validité de l'image IOS par le système (checksum)

5.2.3. Mode RXBoot

Le mode RXBoot, aussi connu sous le nom de ROMmon, peut être utilisé pour l'une des deux raisons suivantes :

- Procédure de récupération des mots de passe
- Récupération du système après un problème d'image IOS

Pour accéder au mode RXBoot, il faut utiliser la combinaison de touches **CTRL+Pause** pendant les 60 secondes suivant le redémarrage du routeur. Le mode RXBoot est reconnaissable de part l'invite de commande affichée (exemple : **Rommon 1>** sur un routeur Cisco 2600).

Les commandes utilisées dans le mode RXBoot sont les suivantes :

- **confreg [valeur]**
 - Sans paramètre, cela permet d'afficher/modifier les paramètres de la ligne console (vitesse, etc.).
 - Avec paramètre, **valeur** correspond à la valeur hexadécimale du registre de configuration à attribuer. Ceci est utile lors de la récupération des mots de passe
- **xmodem -c {nom fichier}**
 - Lance la demande de chargement d'une image IOS au travers de la ligne console
- **dir {système de fichier}**
 - Affiche le contenu d'un système de fichiers
- **boot [{préfixe}:{fichier}]**
 - Démarre le routeur en utilisant une image IOS précise (syntaxe uniformisée d'IFS)
- **set**
 - Permet de visualiser les valeurs des variables d'environnement
- **IP_ADDRESS={IP}**
 - Variable d'environnement spécifiant l'adresse IP du routeur

- **IP_SUBNET_MASK={SM}**
 - o Variable d'environnement spécifiant le masque de sous-réseau du routeur
- **DEFAULT_GATEWAY={IP}**
 - o Variable d'environnement spécifiant l'adresse IP de la passerelle par défaut pour le routeur
- **TFTP_SERVER={IP}**
 - o Variable d'environnement spécifiant l'adresse IP du serveur TFTP à utiliser
- **TFTP_FILE={{répertoire}/{nom fichier}}**
 - o Variable d'environnement spécifiant l'emplacement de l'image IOS sur ce serveur TFTP
- **tftpdnld**
 - o Lance le téléchargement de l'image IOS en utilisant les valeurs des variables d'environnement
- **reset**
 - o Redémarre le routeur
- **i**
 - o Quitte le mode RXBoot et continue la séquence d'amorçage du routeur

Les commandes ci-dessus peuvent varier en fonction de la plateforme utilisée. Elles correspondent au mode RXBoot des dernières plateformes Cisco et ne fonctionnent pas sur les anciennes (tel que les routeurs Cisco 2500).

Il est plausible qu'un problème survienne avec IOS. Ceci peut aller de l'utilisation d'une image non prévue pour la plateforme à l'utilisation d'une image n'ayant pas assez de mémoire RAM pour se lancer, en passant par de bien nombreuses autres possibilités.

Dans ce genre de situations, le seul recours est le mode RXBoot. Pour la récupération d'une image IOS, on peut procéder de 2 manières différentes :

- Méthode **Xmodem**
- Méthode **tftpdnld**

La première méthode (Xmodem) est utilisée lorsque le routeur n'est branché qu'à un ordinateur via son port console (il faut posséder l'image IOS sur l'ordinateur relié au routeur par le câble console) :

- **Etape n°1** : Modifier les paramètres de la ligne console avec la commande **confreg** (vitesse par défaut de 56000 bauds à changer en 115200 bauds)
- **Etape n°2** : Redémarrer le routeur avec la commande **reset** puis entrer de nouveau dans le mode RXBoot
- **Etape n°3** : Lancer la demande de téléchargement avec la commande **xmodem -c {nom fichier}**
- **Etape n°4** : Lancer le téléchargement de l'image grâce au protocole Xmodem depuis le logiciel d'émulation de terminaux
- **Etape n°5** : Une fois le téléchargement terminé, effectuer un redémarrage du routeur

La deuxième méthode (tftpdnld) est utilisée lorsqu'un serveur TFTP est disponible sur le réseau :

- Etape n°1 : Configurer toutes les variables d'environnement
- Etape n°2 : Vérifier les variables d'environnement avec la commande **set**
- Etape n°3 : Lancer le téléchargement de l'image IOS avec la commande **tftpdnld**
- Etape n°4 : relancer la séquence d'amorçage du routeur avec la commande **i** ou **reset**

6. Routage

6.1. Principes fondamentaux

6.1.1. Fonctions de routage et de commutation

La couche réseau fournit un acheminement de bout en bout et au mieux des paquets à travers les réseaux interconnectés. Ceci est effectué par 2 fonctions distinctes :

- **Fonction de routage**
- **Fonction de commutation**

La fonction de routage utilise la table de routage du protocole routé utilisé par le paquet à faire transiter pour déterminer le meilleur chemin à emprunter pour atteindre le réseau de destination. La métrique est utilisée pour offrir une mesure de qualité des différents chemins.

La fonction de commutation permet à un routeur d'accepter un paquet d'une file d'attente d'entrée et de le transmettre à une file d'attente de sortie.

Le but de ces deux fonctions est donc complètement différent et entièrement complémentaire.

Il existe plusieurs méthodes permettant d'optimiser la relation entre les fonctions de routage et de commutation. Ces méthodes permettent l'accélération de la transmission des paquets au travers d'un routeur en mettant en mémoire cache, les décisions de routage déjà prises. Il existe les méthodes suivantes :

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

Par défaut, un routeur Cisco utilise le Fast Switching, qui permet de mettre en mémoire cache les décisions de routage pour chaque destination. Pour cela, la première décision est effectuée normalement, en passant successivement par les fonctions de routage et de commutation. A ce moment là, on place en mémoire cache la décision de routage (l'interface de sortie) ainsi que l'en-tête de trame qui fut généré pour la trame de sortie.

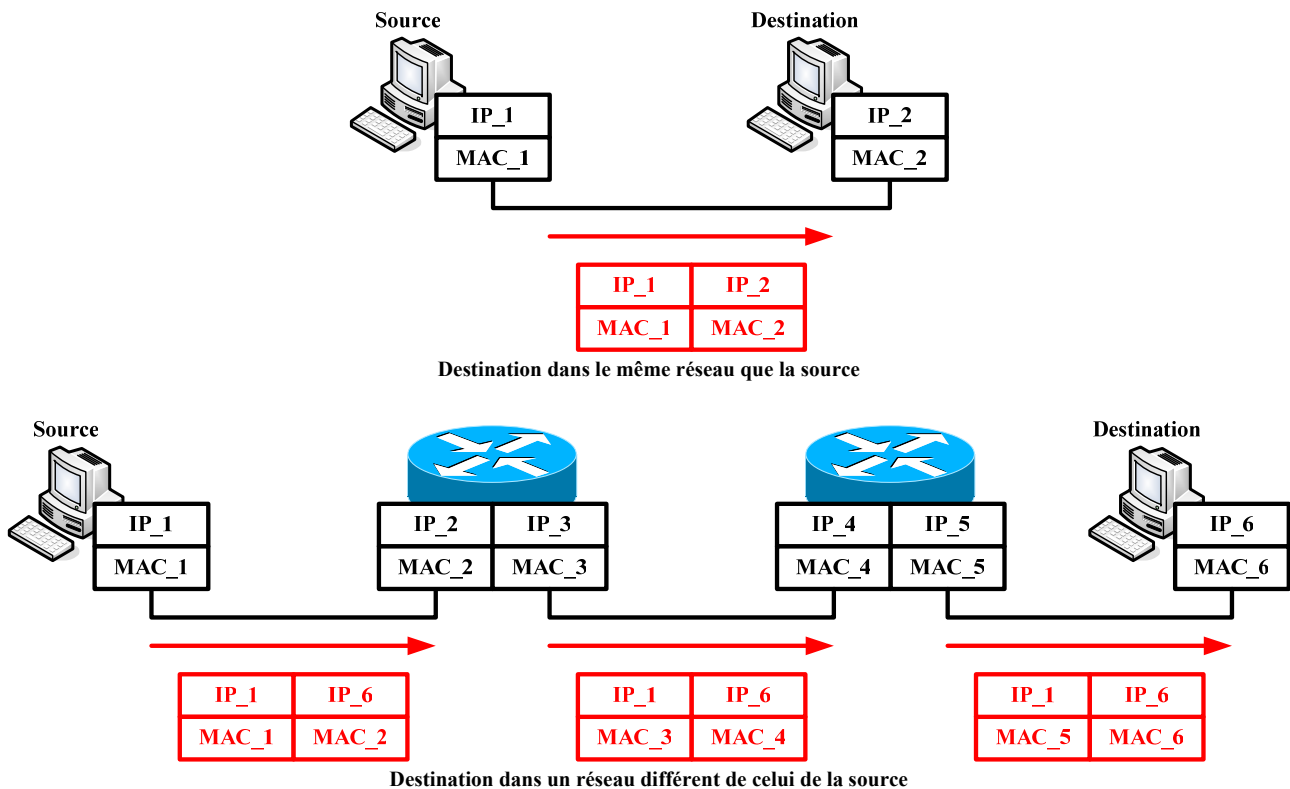
Les paquets suivants pour cette même destination se verront alors automatiquement traités de la même manière que le premier, en utilisant la même interface de sortie ainsi que le même en-tête de trame. Cela permet donc d'économiser le temps de parcours de la table de routage ainsi que le temps de création de l'en-tête pour la nouvelle trame.

Sauf exceptions, ces méthodes ont un inconvénient majeur, à savoir que seule la première décision de routage est mise en mémoire cache. Cela signifie que le partage de charge entre plusieurs liens pour une même destination devient impossible. Il faut donc choisir entre rapidité de transmission par le routeur et répartition de charge.

La commande suivante peut être utilisée :

- **[no] ip route-cache**
 - o Mode de configuration d'interface
 - o Active/désactive le Fast Switching sur l'interface courante
 - o Actif par défaut

6.1.2. Processus de transmission



Le processus de transmission des paquets se déroule comme suit :

- L'hôte source détermine si la destination est locale (même réseau ou sous-réseau) ou distante grâce au couple IP/masque de sous-réseau. Elle calcule ainsi l'adresse IP de sous-réseau de la destination ainsi que la sienne.
- Si les adresses 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 adresses 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. Ce routeur va 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 l'analyse d'une 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.

6.1.3. Table(s) de routage

La table de routage est l'élément central d'un routeur. C'est cette table qui est utilisée par la fonction de routage pour déterminer le meilleur chemin pour chaque destination connue du routeur.

Il existe une seule table de routage par protocole routé, sachant que cette table de routage peut être complétée manuellement (routage statique) ou dynamiquement (protocoles de routage).

Une table de routage possède les champs suivants :

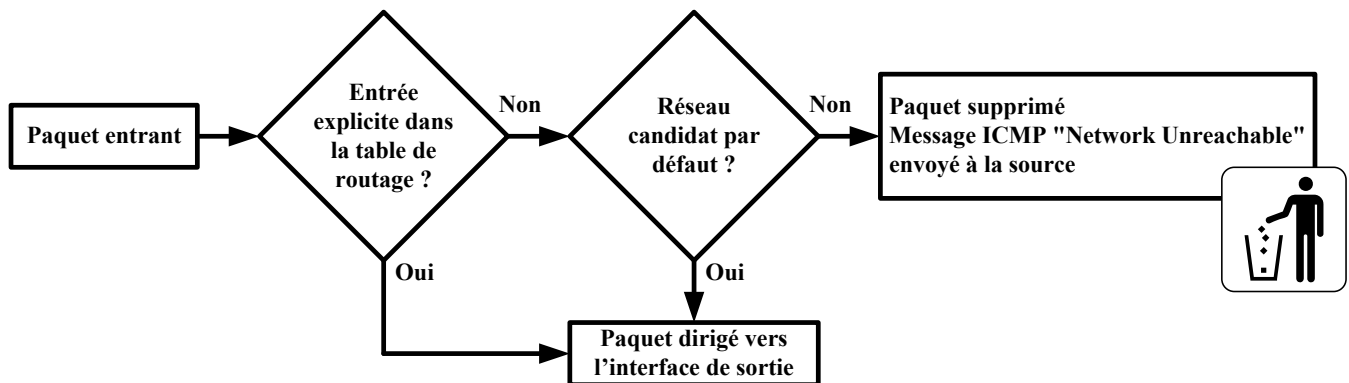
- **Destination**
 - o Jusqu'à 6 ou 16 (IOS >= 12.3(2)T) routes différentes pour une même destination peuvent exister dans la table de routage. Ceci permet la répartition de charge sur plusieurs liens (Round Robin).
 - o Ces entrées doivent obligatoirement avoir un prochain saut différent.
 - o Il ne peut exister qu'une seule entrée dans la table de routage pour une même destination passant par un même prochain saut.
- **Interface de sortie**
 - o Interface locale du routeur vers laquelle le paquet sortira.
- **Prochain saut**
 - o Adresse de couche 3 du prochain routeur sur le chemin pour atteindre le réseau de destination.
- **Métrique**
 - o Il s'agit d'une valeur numérique, utilisée par les protocoles de routage, qui permet la sélection du meilleur chemin et qui est basée sur des critères propres à chaque protocole de routage.
 - o Plus la métrique est petite, meilleure est la route.
- **Distance administrative**
 - o Cette valeur numérique permet d'indiquer un ordre de préférence entre les différents protocoles lorsque plusieurs d'entre eux concourent pour une même entrée dans la table de routage. En effet, il est presque impossible de comparer objectivement les informations fournies par différents protocoles de routage en utilisant leurs métriques calculées avec des critères différents.
 - o Plus la distance administrative est petite, plus le protocole est considéré comme prioritaire.
 - o Les différentes valeurs à connaître sont :

Protocole	Distance administrative
Directement connecté	0
Statique	1
RIP	120
IGRP	100

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

Code	Protocole
C	Directement connecté
S	Statique
R	RIP
I	IGRP
*	Candidat par défaut

Un réseau candidat par défaut (aussi appelé route par défaut) est une entrée de table de routage qui dirige les paquets vers un saut suivant défini, lorsqu'il n'y a pas d'entrée explicite pour le réseau de destination. Ce type de route est utilisé par exemple pour rediriger les paquets d'un réseau LAN vers Internet.



Routage des paquets en fonction des entrées dans la table de routage

Tout paquet qu'un routeur reçoit n'ayant pas d'entrée explicite ou implicite (réseau candidat par défaut) dans la table de routage est détruit. Le message ICMP "Network Unreachable" est alors envoyé par le routeur à la station source du paquet.

La décision prise par la fonction de routage est basée sur le principe de la correspondance la plus longue. Ceci signifie que si plusieurs entrées existent dans la table de routage, la plus précise correspondant à la destination sera choisie.

6.2. Routage statique et dynamique

6.2.1. Caractéristiques et comparatif

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 le nécessite.
- **Dynamique** : Une fois qu'un administrateur réseau a entré les commandes de configuration pour lancer le routage dynamique, les informations relatives aux routes sont mises à jour automatiquement, par un processus de routage.

Le routage statique offre plusieurs applications utiles :

- Le routage dynamique a tendance à révéler toutes les informations connues d'un réseau, 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.

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

6.2.2. Caractéristiques des protocoles de routage

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

Les protocoles de routage peuvent être classés selon l'algorithme qu'ils utilisent :

- Vecteur de distance
- Etat de liens
- Hybride symétrique

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.

Un protocole de routage peut calculer les métriques en fonction de critères tels que :

- **Bande passante** : Le débit d'une liaison, mesuré en bits par seconde.
- **Délai** : Le temps requis pour acheminer un paquet, 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** : Le nombre de routeurs par lesquels un paquet doit passer avant d'arriver à destination.
- **Tics** : L'intervalle de temps entre 2 trames pour une liaison de donnée précise (environ 55 millisecondes).
- **Coût** : Généralement basée sur une dépense monétaire attribuée à un lien par un administrateur réseau.

6.3. Convergence, boucles de routage et solutions

6.3.1. Convergence

La convergence est le fait que tous les dispositifs réseau ont la même vue de la topologie du réseau. Le temps de convergence est donc le temps pendant lequel les dispositifs réseaux n'ont pas la même vue de celui-ci.

Lorsque tous les routeurs d'un réseau utilisent les mêmes informations, le réseau 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.

6.3.2. Boucles de routage

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.

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

- Métrique de mesure infinie (Finite State Metric)
- Split Horizon
- Route Poisoning
- Mises à jour déclanchées (Triggered Updates)
- Compteurs de retenue (Holddown Timers)

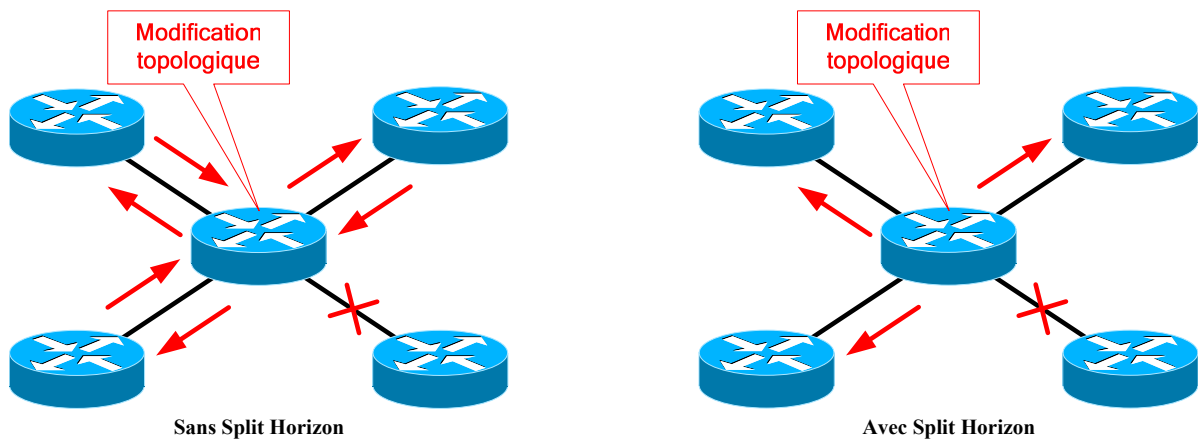
Ces cinq méthodes sont uniquement utilisées par les protocoles de routage à vecteur de distance, afin d'essayer de contrer les plausibles boucles de routage.

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

6.3.3. Métrique de mesure infinie

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.

6.3.4. Split Horizon



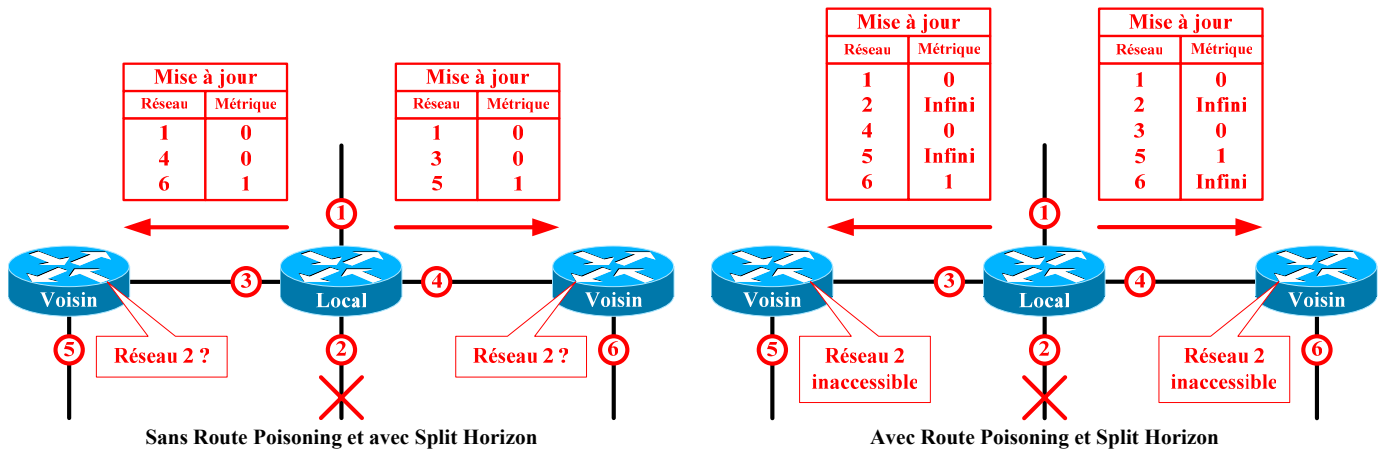
Le principe du Split Horizon est simple. Aucune information de mise à jour ne sera renvoyée par le chemin par lequel on a appris la modification de topologie. Ceci permet d'éviter de renvoyer à la source des informations erronées.

Ceci implique donc que l'information se propage toujours du plus près du réseau de destination au plus éloigné, sans jamais revenir en arrière.

6.3.5. Route Poisoning

Le Route Poisoning, aussi appelé Poison Reverse, est utilisé lorsqu'un réseau devient inaccessible. Au lieu de n'avertir que les routes existantes dans la table de routage aux voisins, le Route Poisoning inclut aussi les routes devenues inaccessibles en leur octroyant une métrique infinie.

Ceci permet d'informer directement les voisins qu'un réseau est devenu inaccessible au lieu d'attendre l'expiration de leur compteur d'invalidité (Invalid Timer).



Combiné au Split Horizon, le Route Poisoning n'exclut pas les routes concernées par la règle du Split Horizon mais leur attribue une métrique infinie.

6.3.6. Mises à jour déclenchées

Les mises à jour déclenchées servent à informer les voisins d'une modification topologique au moment où elle survient. Cela permet de réduire le temps de convergence en n'attendant pas l'expiration de l'intervalle de temps de transmission des mises à jour périodiques.

6.3.7. Compteurs de retenue

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 a eu lieu 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.

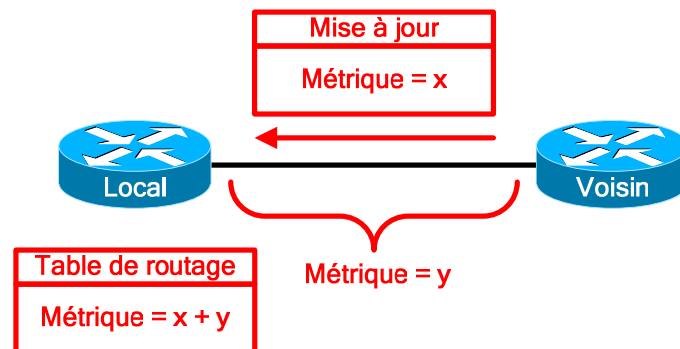
Pour calculer le temps à utiliser pour la configuration des compteurs de retenue, il faut multiplier le plus petit nombre de sauts à effectuer pour atteindre le routeur le plus éloigné par l'intervalle de temps entre les mises à jour.

6.4. Routage à vecteur de distance

L'algorithme de routage à vecteur de distance possède une vision de la topologie du réseau qui est basée sur celle de ses voisins. En effet, les mises à jour de routage envoyées par les protocoles de routage à vecteur de distance contiennent directement la table de routage du routeur émetteur.

Le récepteur n'a donc pour seul travail que de récupérer ces informations, de garder les entrées pertinentes et de modifier les métriques.

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



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

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

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

Les protocoles de routage à vecteur de distance les plus connus sont :

- RIPv1
- RIPv2
- Cisco IGRP
- Cisco EIGRP (vecteur de distance évolué, ou hybride symétrique)

6.5. Routage à état de liens

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

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

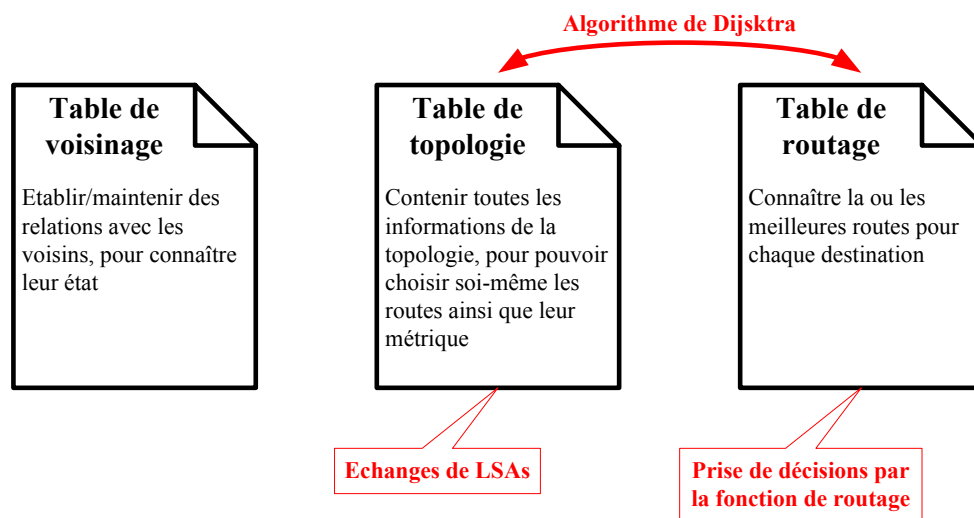
Les mises à jour de routage des protocoles à état des liens possèdent de grandes différences comparées à celles des protocoles à vecteur de distance :

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

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

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

Le routage à état de liens se base donc sur l'utilisation de trois tables distinctes (au contraire des protocoles à vecteur de distance qui ne gèrent que la table de routage) :



Tables utilisées par un protocole de routage à état de liens

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

- **Ressource calculatoire** : Un protocole de routage à état de liens requière une puissance CPU importante pour l'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 de liens car il faut stocker les tables de voisinage ainsi que de topologie en plus de la classique table de routage.

Les protocoles de routage à état de liens les plus connus sont :

- OSPF
- IS-IS

6.6. Systèmes autonomes, protocoles de routage intérieurs et extérieurs

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

Ces AS sont identifiés par un numéro qui est chiffré sur 16 bits. Ce numéro est unique dans le monde et permet d'identifier une organisation aux yeux du reste du monde informatique. Il est attribué par le NIC (Network Information Center).

Pour les protocoles de routage imposant l'indication d'un numéro d'AS et se trouvant dans un réseau privé, ce numéro de système autonome peut être choisi arbitrairement dans la plage de valeurs allant de 64512 à 65535.

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

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

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

Les protocoles de routage sont donc classifiés ainsi :

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

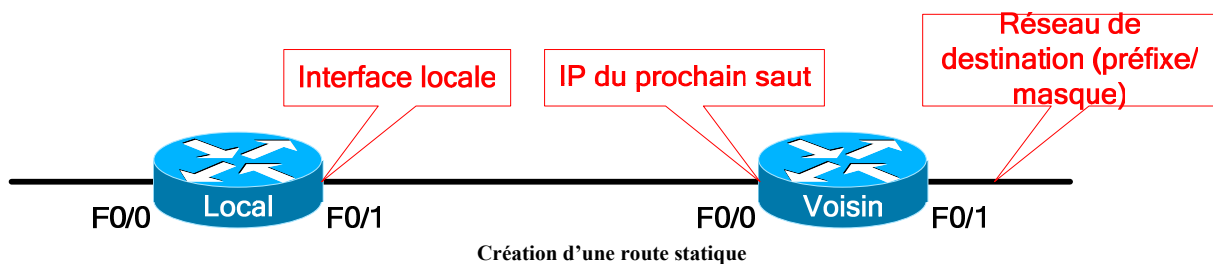
Typiquement, la convergence d'un réseau est restreinte au système autonome. Le temps de convergence dépend donc du protocole utilisé dans le système autonome.

6.7. Configuration par défaut, routage statique et visualisation d'état

Par défaut, seul le routage pour le protocole IP est activé sur un routeur Cisco. Le routage s'effectue automatiquement entre les réseaux directement connectés au routeur, sans avoir à utiliser des routes statiques ou un protocole de routage quelconque.

Les commandes permettant de configurer le routage de base sont les suivantes :

- **{protocole} routing**
 - o Mode de configuration globale
 - o Permet d'activer/désactiver le routage pour un protocole routé particulier
 - o Le paramètre **protocole** correspond au mot clé du protocole voulu (IP, IPX, IPv6, etc.)
- **ip classless**
 - o Mode de configuration globale
 - o Active le routage Classless sur le routeur
 - o Actif par défaut
 - o Permet l'utilisation d'information de routage Classless, telles que les routes par défaut
- **ip route {préfixe} {masque} {prochain saut | interface} [distance administrative]**
 - o Mode de configuration globale
 - o Crée une route statique sur le routeur
 - o La distance administrative permet la création d'une route statique flottante (valeur par défaut = 1)

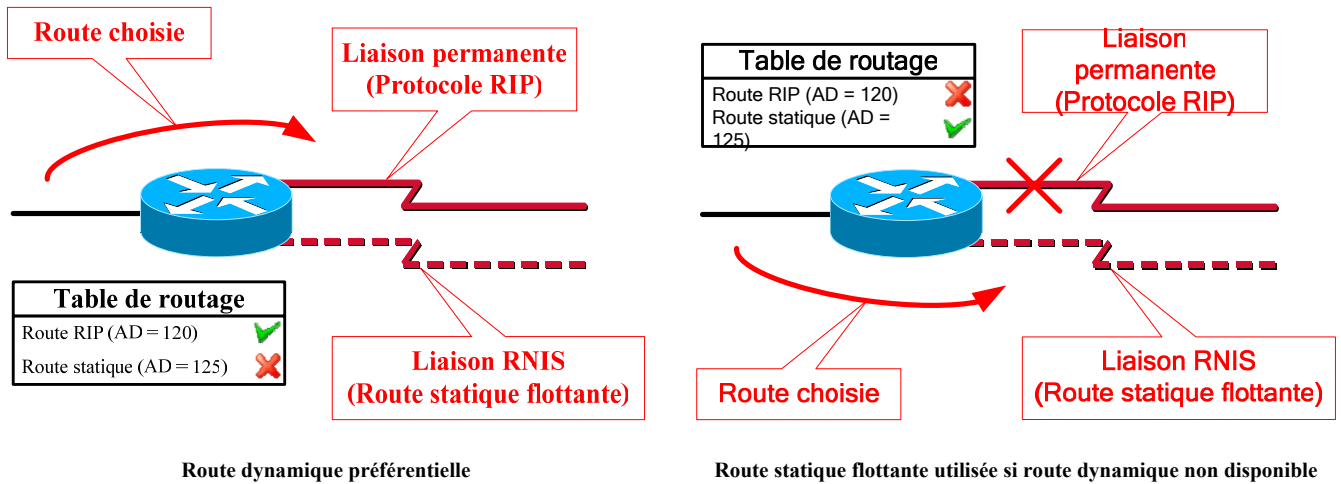


Il est possible de créer une route statique par défaut. Pour cela, il suffit d'utiliser le pseudo réseau ayant pour préfixe 0.0.0.0 et pour masque de sous-réseau 0.0.0.0. Cette route statique sera considérée par le routeur comme un réseau candidat par défaut dans la table de routage.

Les routes statiques sont prioritaires à n'importe quel protocole de routage, à cause de la distance administrative par défaut (égale à 1). Cette distance peut être modifiée afin de rendre une route statique moins préférable à une entrée fournie par un protocole de routage.

Pour cela, il faut expliciter pour la route statique une distance administrative plus grande que celle du protocole de routage.

On crée ainsi une route statique flottante, qui est une route alternative à une autre en cas de défaillance de la première. Une route statique flottante doit être pour la même destination qu'une entrée fournie par un protocole de routage.



Cette route statique flottante n'apparaît dans la table de routage que lorsque l'entrée fournie par le protocole de routage n'est plus valide.

Les commandes utilisées pour la visualisation d'état sont :

- **show ip protocols** : Affiche la liste des protocoles de routage configurés sur le routeur ainsi que les informations générales les concernant (interfaces participant à chaque processus de routage, réseaux avertis, compteurs, etc.).
- **show ip route** : Affiche la table de routage IP.
- **clear ip route [{préfixe} | *]** : Supprime une ou plusieurs routes de la table de routage.

7. Protocole RIP

7.1. 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.

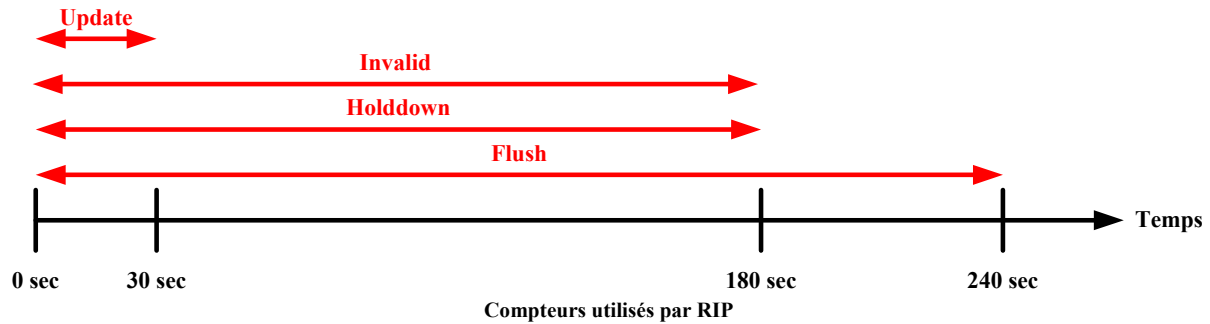
Avantages	Inconvénients
Processus léger	Temps de convergence lent
Implémenté sur tous les systèmes d'exploitation	Nombre de sauts pour calculer les métriques
	Nombre de sauts limité à 15

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).



7.2. Configuration

7.2.1. Commandes

Les commandes liées à la configuration du protocole RIP sont :

- **router rip**
 - Mode de configuration globale
 - Active le protocole RIP
 - Passe dans le mode de configuration du routeur
- **network {préfixe}**
 - Mode de configuration du routeur
 - Spécifie le réseau qui sera inclut dans les mises à jour de routage
 - Détermine les interfaces appartenant à ce réseau qui participent au processus de routage
 - Le **préfixe** doit être un réseau directement connecté au routeur
- **neighbor {IP}**
 - Mode de configuration du routeur
 - Définit l'adresse IP d'un voisin avec lequel RIP échangera des mises à jour de routage
 - Par défaut, aucun voisin n'est défini
- **passive-interface {type} {numéro}**
 - Mode de configuration du routeur
 - Empêche l'interface indiquée d'envoyer des mises à jour
- **[no] ip split-horizon**
 - Mode de configuration d'interface
 - Active/désactive Split Horizon sur l'interface courante
- **timers basic {update} {invalid} {holddown} {flush}**
 - Mode de configuration du routeur
 - Définit les intervalles de temps, en secondes, utilisés par RIP
- **version {1 | 2}**
 - Mode de configuration du routeur
 - Indique la version de RIP utilisée par le routeur
 - Ceci modifie automatiquement le type (RIPv1 ou RIPv2) de mises à jour envoyées et reçues
 - Par défaut, les mises à jour sont de type 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
- **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
- **rip equal-cost {nombre}**
 - o Mode de configuration globale
 - o Indique le nombre d'entrées ayant la même métrique pouvant être insérées dans la table de routage
 - o De 1 à 15 et par défaut à 1

7.2.2. Procédure de configuration

Pour configurer un routeur en utilisant le protocole de routage RIP, il faut procéder comme suit :

- **Etape n°1** : Activer le protocole RIP (commande **router rip**)
- **Etape n°2** : Spécifier les réseaux directement connectés devant participer au processus de routage (commande **network**)
- **Etape n°3 (optionnelle)** : Désactiver l'émission de mises à jour de routage vers les réseaux n'ayant pas de routeur(s) RIP autre(s) que le routeur local (commande **passive-interface**)
- **Etape n°4 (optionnelle)** : Ajuster les différents compteurs de temps (commande **timers basic**)
- **Etape n°5 (optionnelle)** : Choisir la version de RIP à utiliser (commande **version**)
- **Etape n°6 (optionnelle)** : Propager la route par défaut existante sur le routeur local aux autres routeurs RIP du système autonome (commande **default-information originate**)
- **Etape n°7 (optionnelles)** : Activer la répartition de charge entre plusieurs liens de même métrique (commande **maximum-paths**)

Il ne peut y avoir qu'une seule et unique instance du protocole RIP par routeur.

7.3. Vérification

IOS fournit une panoplie de commandes permettant de visualiser l'état du protocole RIP ainsi que d'effectuer du débogage. Ces commandes sont les suivantes :

- **show ip protocols** : Affiche les compteurs RIP, les interfaces participant au processus de routage, les réseaux avertis ainsi que la version pour les mises à jour envoyées et reçues.
- **show ip rip database** : Affiche la FIB (Forward Information Base) de RIP.
- **debug ip rip [events]** : Affiche en temps réel les mises à jour RIP envoyées et reçues.

8. Protocole IGRP

8.1. 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 :

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

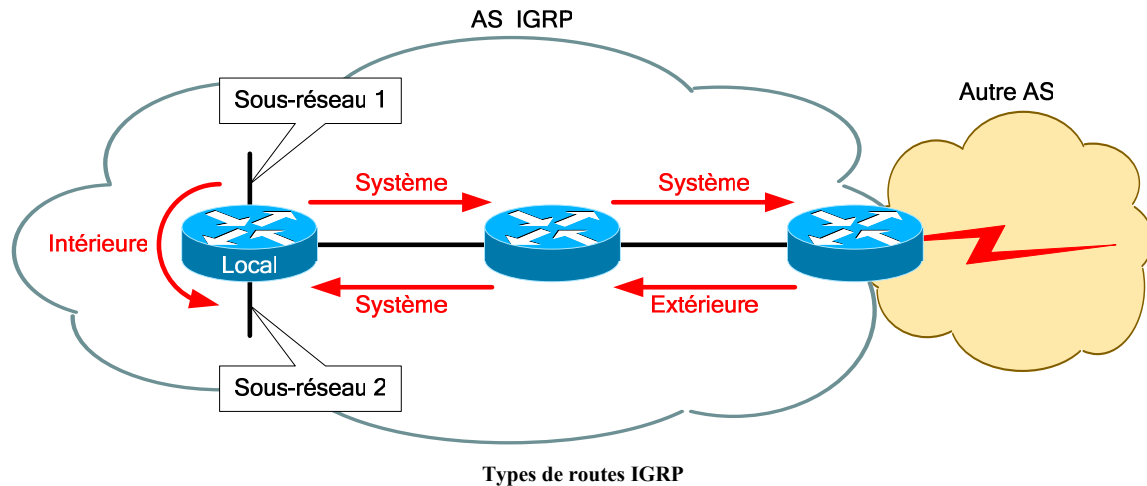
Les différents paramètres de cette formule sont les suivants :

- **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 :

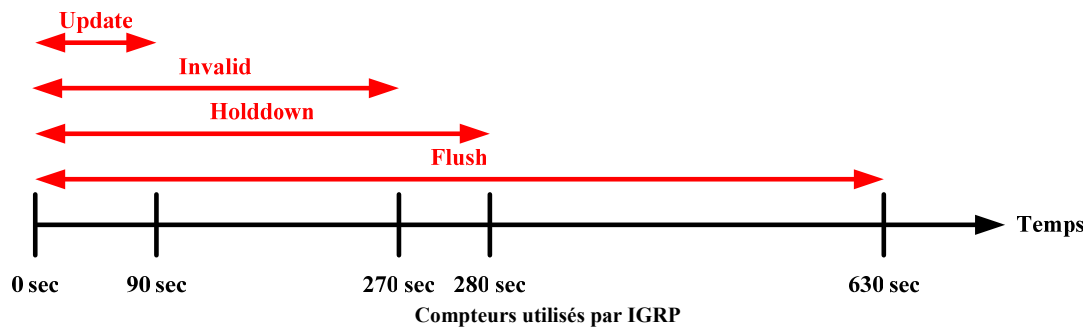
$$\text{Métrique} = \text{Bandwidth} + \text{Delay}$$

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



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.

8.2. Configuration

8.2.1. 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 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
- **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

8.2.2. Procédure de configuration

Pour configurer un routeur en utilisant le protocole de routage IGRP, il faut procéder comme suit :

- **Etape n°1** : Activer le protocole de routage IGRP (commande **router igrp**)
- **Etape n°2** : Spécifier les réseaux directement connectés devant participer au processus de routage (commande **network**)
- **Etape n°3 (optionnelle)** : Désactiver l'émission de mises à jour de routage vers les réseaux n'ayant pas de routeur(s) IGRP autre(s) que le routeur local (commande **passive-interface**)
- **Etape n°4 (optionnelle)** : Ajuster les différents compteurs de temps (commande **timers basic**)
- **Etape n°5 (optionnelle)** : Propager la route par défaut existante sur le routeur local aux autres routeurs IGRP du système autonome (commande **ip default-network**)
- **Etape n°6 (optionnelle)** : Activer la répartition de charge entre plusieurs liens de même métrique (commandes **maximum-paths** et **variance**)

Il ne peut y avoir qu'une seule instance d'IGRP par numéro de système autonome. Il peut donc y avoir plusieurs instances d'IGRP sur un même routeur.

8.3. 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.

9. Protocole ICMP

9.1. Théorie

ICMP (Internet Control Message Protocol) est un protocole faisant partie de la pile de protocoles TCP/IP et fonctionne au niveau de la couche 3 du modèle OSI.

Les messages du protocole ICMP sont classifiés en deux catégories :

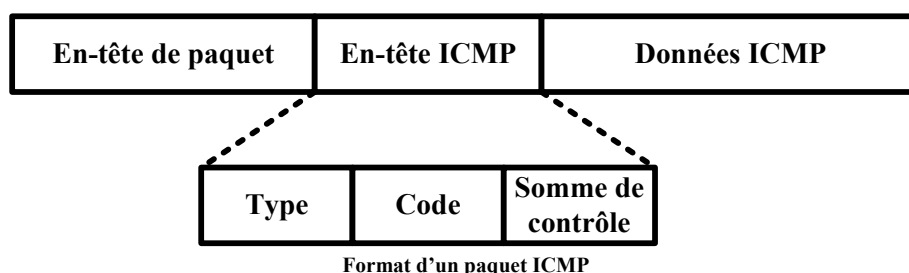
- Messages d'erreurs
- Messages de contrôle

Les messages d'erreurs sont présents pour informer les pairs communiquant d'une erreur de transmission, permettant ainsi de contrer la limitation du protocole IP.

Ces messages d'erreurs ICMP sont eux-mêmes des paquets IP et sont donc aussi sujets aux erreurs de transmission. Afin d'éviter une boucle de messages d'erreurs, les erreurs survenant à des messages ICMP ne génèrent pas de messages d'erreur ICMP.

Les messages de contrôle servent à informer sur l'état du réseau (dispositif congestionné, meilleure passerelle par défaut existante, etc.).

Les messages ICMP sont encapsulés comme toute autre donnée dans un paquet :



9.2. Messages ICMP

9.2.1. Types de messages

Il existe plusieurs types de messages ICMP associés à un numéro de code précis :

Code	Message
0	Echo Reply
3	Destination Unreachable
4	Source Quench
5	Redirect/Change Request
8	Echo Request
9	Router Discovery
10	Router Solicitation
11	Time Exceeded
12	Parameter Problem
13	Timestamp Request
14	Timestamp Reply
15	Information Request
16	Information Reply
17	Address Mask Request
18	Address Mask Reply

9.2.2. Echo Request/Reply

Les messages d'Echo permettent de déterminer si un hôte est joignable ou non. Ceci s'effectue en utilisant la commande ping qui envoie des messages ICMP Echo Request. L'hôte de destination recevant ces paquets renvoie à son tour des messages ICMP Echo Reply.

L'en-tête ICMP utilisé pour ces messages est le suivant :

8 bits	8 bits	16 bits	16 bits	16 bits	Variable
Type (0 ou 8)	Code (0)	Somme de contrôle	Identificateur	Numéro de séquence	Remplissage

En-tête ICMP Echo Request/Reply

Le numéro de séquence est utilisé pour distinguer les différentes requêtes effectuées.

9.2.3. Destination Unreachable

Le message ICMP Destination Unreachable est envoyé par un routeur lorsque ce dernier ne possède pas les informations suffisantes pour transmettre un paquet (typiquement une table de routage n'ayant pas d'entrée pour le réseau de destination).

L'en-tête ICMP utilisé pour les messages ICMP Destination Unreachable est :

8 bits	8 bits	16 bits	16 bits	
Type (3)	Code (de 0 à 12)	Somme de contrôle	Inutilisé (valeur = 0)	En-tête du paquet source + les 64 premiers bits

En-tête ICMP Destination Unreachable

Les différentes valeurs possibles pour le code permettent d'identifier la cause du problème :

Code	Signification
0	Réseau inaccessible
1	Hôte inaccessible
2	Protocole inaccessible
3	Port inaccessible
4	Fragmentation nécessaire mais refusée
5	Echec de la route source
6	Réseau de destination inconnu
7	Hôte de destination inconnu
8	Hôte source isolé
9	Communication avec le réseau de destination administrativement refusée
10	Communication avec l'hôte de destination administrativement refusée
11	Réseau inaccessible pour le ToS utilisé
12	Hôte inaccessible pour le ToS utilisé

9.2.4. Parameter Problem

Un message ICMP Parameter Problem est envoyé lorsqu'un paquet n'a pas pu être transmis à cause d'une erreur d'en-tête IP.

L'en-tête d'un message ICMP Parameter Problem est ainsi :

8 bits	8 bits	16 bits	8 bits	8 bits	
Type (12)	Code (de 0 à 2)	Somme de contrôle	Pointeur	Inutilisé (valeur = 0)	En-tête du paquet source + les 64 premiers bits

En-tête ICMP Parameter Problem

Le champ "Pointeur" permet d'indiquer l'octet de l'en-tête IP posant problème.

9.2.5. Source Quench

Le message ICMP Source Quench est envoyé par un dispositif réseau subissant une congestion réseau.

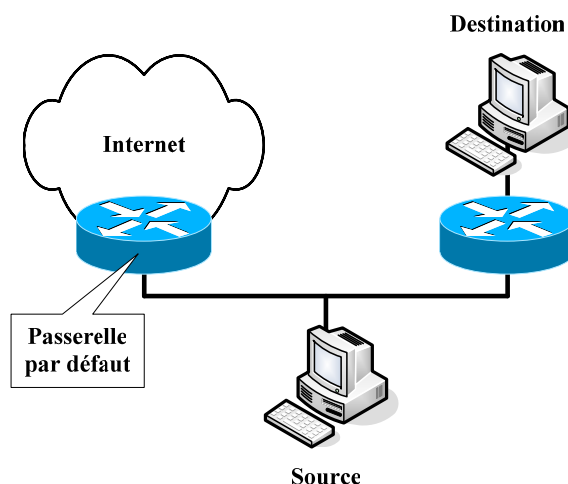
Ne pouvant pas traiter tous les paquets entrant en cas de congestion, il est obligé d'en supprimer. Les sources des paquets supprimés sont averties par ce message ICMP.

9.2.6. Redirect/Change Request

Ce message ICMP permet la notification à la source qu'une meilleure route existe pour une destination précise.

Ce message est envoyé par une passerelle par défaut uniquement si les conditions suivantes sont remplies :

- Interface d'entrée = interface de sortie
- Réseau de la source = réseau du prochain saut
- Route dans la table de routage $\neq$ route par défaut
- Paquet reçu n'est pas un ICMP Redirect
- Routeur configuré pour envoyer des messages ICMP Redirect



L'en-tête du message ICMP envoyé par la passerelle par défaut est le suivant :

8 bits	8 bits	16 bits	32 bits	
Type (5)	Code (de 0 à 3)	Somme de contrôle	IP d'un routeur	En-tête du paquet source + les 64 premiers bits

En-tête ICMP Redirect

Le champ "IP d'un routeur" fournit à la source l'adresse IP du prochain saut à utiliser pour la destination qu'elle a cherchée à atteindre.

Le code peut avoir ces valeurs :

Code	Signification
0	Redirection pour le réseau de destination
1	Redirection pour l'hôte de destination
2	Redirection pour le ToS pour le réseau de destination
3	Redirection pour le Tos pour l'hôte de destination

Pour configurer le ICMP Redirect sur un routeur Cisco, il faut utiliser cette commande :

- **[no] ip redirects**
 - o Mode de configuration d'interface
 - o Active/désactive les messages ICMP Redirect
 - o Actif par défaut

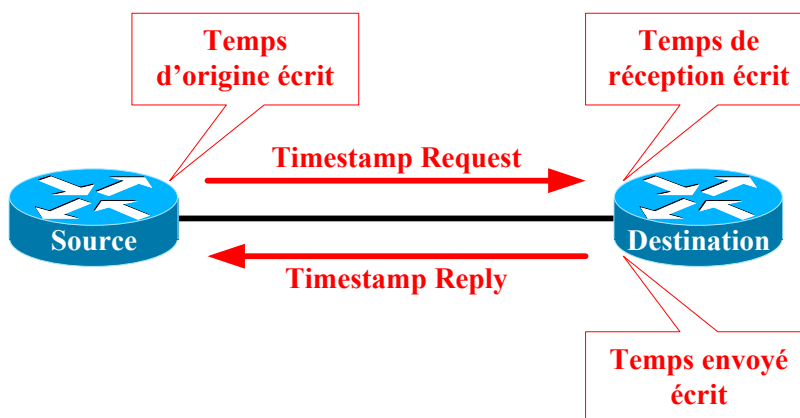
9.2.7. Timestamp Request/Reply

Les deux messages ICMP ont été créés afin d'aider à la synchronisation du temps entre des dispositifs. L'en-tête utilisé est ainsi :

8 bits	8 bits	16 bits	16 bits	16 bits	32 bits	32 bits	32 bits
Type (13 ou 14)	Code (0)	Somme de contrôle	Identifica teur	Numéro de séquence	Temps d'origine	Temps reçu	Temps envoyé

En-tête ICMP Timestamp Request/Reply

Les trois temps envoyés, en millisecondes depuis minuit du temps universel (UT), permettent aux deux pairs de vérifier l'heure de l'autre, et sont inscrits suivant cette séquence :



Processus d'échange de messages ICMP Timestamp

Ces messages sont très utiles pour synchroniser les dispositifs entre eux ainsi que pour déterminer le temps de transmission sur la liaison les reliant. De nos jours, le protocole NTP (Network Time Protocol) est utilisé à la place de ces messages ICMP.

9.2.8. Information Request/Reply

Ces messages étaient utilisés par un hôte pour déterminer son adresse réseau.
Ces messages ICMP sont remplacés par les protocoles BOOTP, RARP et DHCP.

9.2.9. Address Mask Request/Reply

Ces messages permettent à un hôte de demander à sa passerelle par défaut le masque de sous-réseau à utiliser.

9.2.10. Router Discovery/Solicitation

Ces deux messages sont utilisés pour indiquer aux hôtes d'un réseau de l'adresse de leur passerelle par défaut lorsqu'ils ne la connaissent pas.

Le message ICMP Router Solicitation est envoyé par un hôte n'ayant pas de passerelle par défaut.
Le message ICMP Router Discovery est envoyé par le routeur en réponse à un message ICMP Router Solicitation.

10. Résolution de problèmes

10.1. Commandes de vérification

Trois commandes vous permettent de vérifier la configuration des adresses dans votre réseau :

- **telnet {IP ou nom d'hôte} [tcp-port-number]** : Mécanisme de test le plus complet car permet de vérifier toutes les couches du modèle OSI.
- **ping [IP ou nom d'hôte]** : Mécanisme de test de base pour la couche 3 permettant de vérifier la connexion matérielle et l'adresse de couche réseau du modèle OSI pour une destination précise.
- **trace {IP ou nom d'hôte}** : Génération de messages à partir de chaque routeur situé tout au long du chemin jusqu'à la destination.

La commande **telnet** permet, 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'expliciter 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 :

!	Réception réussie d'une réponse d'écho
.	Délai d'attente dépassé pour la réponse à la requête
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

Utilisée sans aucun paramètre depuis le mode privilégié, la commande ping devient ce qui est appelé la commande ping étendue, permettant de modifier les paramètres pour les requêtes.

Lorsqu'on utilise la commande **tracert**, 3 analyseurs sont lancés sur chaque routeur rencontré sur le chemin menant à la destination, afin d'obtenir les temps de réponse pour chacun d'entre eux. Ceci est très utile pour déterminer l'emplacement d'un plausible problème ou d'un goulet d'étranglement.

S'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

10.2. Erreurs courantes et modèle OSI

L'une des méthodes pouvant être utilisée pour la résolution des problèmes est la vérification des différentes couches du modèle OSI en commençant par la plus basse.

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

10.3. Débogage

IOS met à notre disposition toute une panoplie de commandes nous permettant de vérifier en temps réel les interactions et communications. 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 de débogage avec parcimonie car elles exigent un temps processeur important.

Elles sont disponibles depuis le mode privilégié.

En plus des commandes de débogage déjà étudiées, les commandes suivantes sont disponibles :

- **no debug all** : Permet de stopper tous les débogages en cours.
- **undebug all** : Permet de stopper tous les débogages en cours.
- **debug all** : Affiche l'intégralité des informations de débogage disponibles.

10.4. 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.

Pour cette procédure, il faut avoir impérativement un accès physique au routeur, par le biais du port console.

Cette procédure peut varier en fonction de la plateforme utilisée, et est effectuée en 2 redémarrages :

- **Redémarrage n°1** : Modification du registre de configuration depuis le mode RXBoot.
- **Redémarrage n°2** : Modification de la configuration du routeur sous IOS.

Pour le redémarrage n°1, il faut faire ainsi :

- Redémarrer le routeur (interrupteur ou avec la commande **reload** sous IOS).
- Utiliser la combinaison de touches **CTRL+Pause** avant expiration des 60 secondes suivant le redémarrage.
- On se trouve alors dans le mode RXBoot. Il faut maintenant changer la valeur du registre de configuration afin de forcer le routeur à ignorer le fichier de configuration de sauvegarde lors du démarrage :
 - o Commande **o/r 0x2142** (routeurs 2500).
 - o Commande **confreg 0x2142** (routeurs 1600, 1700, 2600, 3600, etc.).
- Sortir du mode RXBoot et relancer le routeur :
 - o Commande **i** (routeurs 2500).
 - o Commande **reset** (routeurs 1600, 1700, 2600, 3600, etc.).

Au redémarrage n°2, nous sommes de nouveaux sous IOS. Il ne nous reste plus qu'à effectuer ces étapes :

- Le mode SETUP nous demande si l'on souhaite effectuer la configuration basique du routeur. Il suffit de refuser en répondant **N** ou en utilisant la combinaison de touche **CTRL+C**.
- On peut ensuite accéder au mode privilégié sans aucun mot de passe. A ce niveau, le routeur a repris sa configuration d'usine.
- Il faut restaurer la valeur initiale du registre de configuration en utilisant la commande **config-register 0x2102** depuis le mode de configuration globale.

Nous avons maintenant la possibilité de restaurer la configuration d'avant, tout en modifiant les mots de passe, ou de laisser le routeur dans sa configuration d'usine.

Pour restaurer la configuration précédente et changer les mots de passe, il faut faire ainsi :

- Importer la configuration précédente (commande **copy start run**).
- Changer les mots de passe des différentes lignes ainsi que pour le mode privilégié.
- Sauvegarder l'ancienne configuration avec les nouveaux mots de passe (commande **copy run start**).
- Redémarrer le routeur (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.

11. ACL

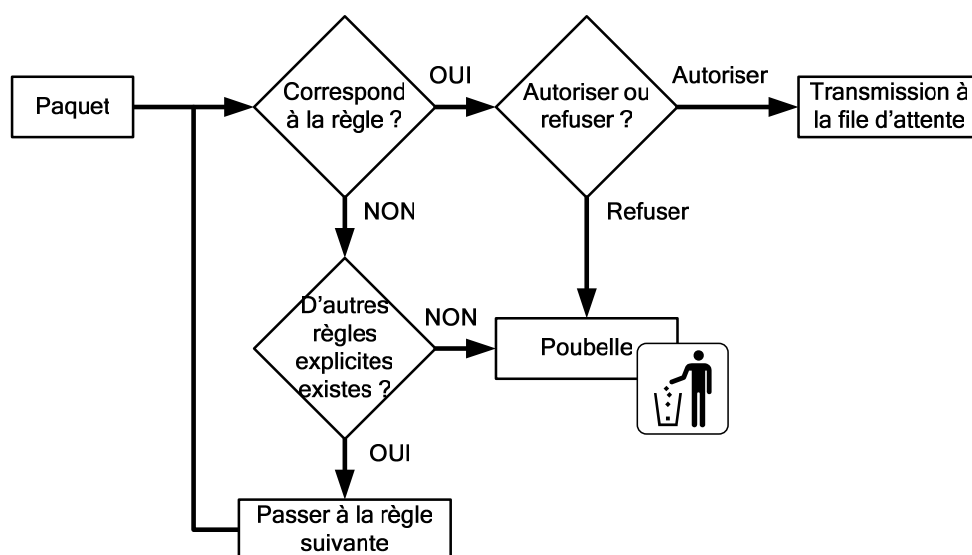
11.1. Théorie

11.1.1. 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.

11.1.2. 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**

11.2. 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.

11.3. 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.

11.4. 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

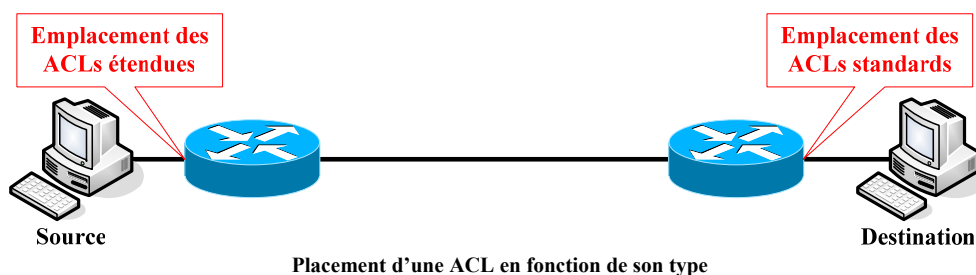
11.5. Mise en place et vérification des ACLs

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).

Table des matières

1. Routage Classless	4
1.1. Introduction au routage Classless	4
1.2. CIDR	5
1.3. VLSM	6
1.4. Procédure de réalisation	6
1.4.1. VLSM Symétrique	6
1.4.2. VLSM Asymétrique	8
1.5. Configuration	9
2. Protocole RIPv2	10
2.1. Rappels sur RIPv1	10
2.2. Spécifications de RIPv2	10
2.3. Configuration	11
2.3.1. Commandes générales	11
2.3.2. Authentification	12
3. Protocole OSPF	13
3.1. Caractéristiques	13
3.2. Définitions	14
3.3. Fonctionnement dans un réseau ne comportant qu'une aire	15
3.3.1. Découverte des routeurs voisins	15
3.3.2. Etablissement des bases de données topologiques	15
3.3.2.1. Dans un réseau point-à-point	15
3.3.2.2. Dans un réseau multi-accès	16
3.4. Opérations OSPF	17
3.4.1. Election du DR / BDR	17
3.4.2. Détermination du Router-ID	17
3.5. Construction de la table de routage	18
3.6. Commandes	19
3.6.1. Commandes générales	19
3.6.2. Authentification	19
3.6.3. Timers	20
3.6.4. Commandes show associées	20
4. Protocole EIGRP	21
4.1. Caractéristiques	21
4.2. Termes et définition	22
4.3. Métriques	23
4.4. Protocole Hello	25
4.4.1. Neighbor Table	26
4.4.2. Topology Table	26
4.5. DUAL	27
4.6. Commandes	28
4.7. Configuration	30

5.	Design de LAN	31
5.1.	Présentation.....	31
5.2.	Méthodologie de conception.....	31
5.3.	Fonction et emplacements des serveurs.....	32
5.4.	Conception de couche 1.....	33
5.5.	Conception de couche 2.....	34
5.6.	Conception de couche 3.....	35
6.	Commutation.....	36
6.1.	Concepts et fonctionnement.....	36
6.2.	Commutateurs.....	38
6.2.1.	Présentation	38
6.2.2.	Démarrage	38
6.2.3.	Configuration de base.....	38
6.2.4.	Voyants d'un commutateur	39
6.2.5.	Commandes.....	40
6.2.6.	Procédure de récupération des mots de passe.....	40
6.3.	Protocole Spanning-Tree	41
6.3.1.	Théorie concernant Spanning-Tree	41
6.3.2.	Théorie concernant Rapid Spanning-Tree.....	42
6.3.3.	Commandes et configuration de Spanning-Tree	43
6.4.	VLAN	44
6.4.1.	Concepts.....	44
6.4.2.	Commandes générales.....	45
6.4.3.	Commandes show associées.....	45
6.4.4.	Configuration	46
6.5.	Trunking	46
6.5.1.	Protocole ISL.....	47
6.5.2.	Protocole 802.1q.....	47
6.5.3.	Comparaison entre ISL et IEEE 802.1q	48
6.5.4.	Commandes associées	48
6.6.	VTP.....	49
6.6.1.	Théorie sur le protocole VTP	49
6.6.2.	Commandes associées	50

1. Routage Classless

1.1. 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.

1.2. 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.

1.3. 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.

1.4. 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.

1.4.1. VLSM Symétrique

Le VLSM symétrique est un plan d'adressage qui fait un découpage récursif de la topologie du réseau de l'entreprise sachant que les différents découpages sont similaires.

Exemple : si l'entreprise a deux bâtiments par ville, on devra avoir deux bâtiments dans chaque ville.

Dans cette procédure, on parle de sous réseau uniquement pour les parties les plus proches des utilisateurs. Tous les autres niveaux de la hiérarchie seront considérés comme une adresse agrégée.

Procédure :

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

Recenser les différents niveaux hiérarchiques de l'entreprise et dessiner la topologie.

- **Etape 2 : Au niveau utilisateur :**

Connaître la taille du sous-réseau.

- **Etape 3 : Recensement :**

Déterminer le nombre de bits nécessaires pour recenser chaque instance du niveau hiérarchique.

- **Etape 4 : Classe d'adresse utilisée :**

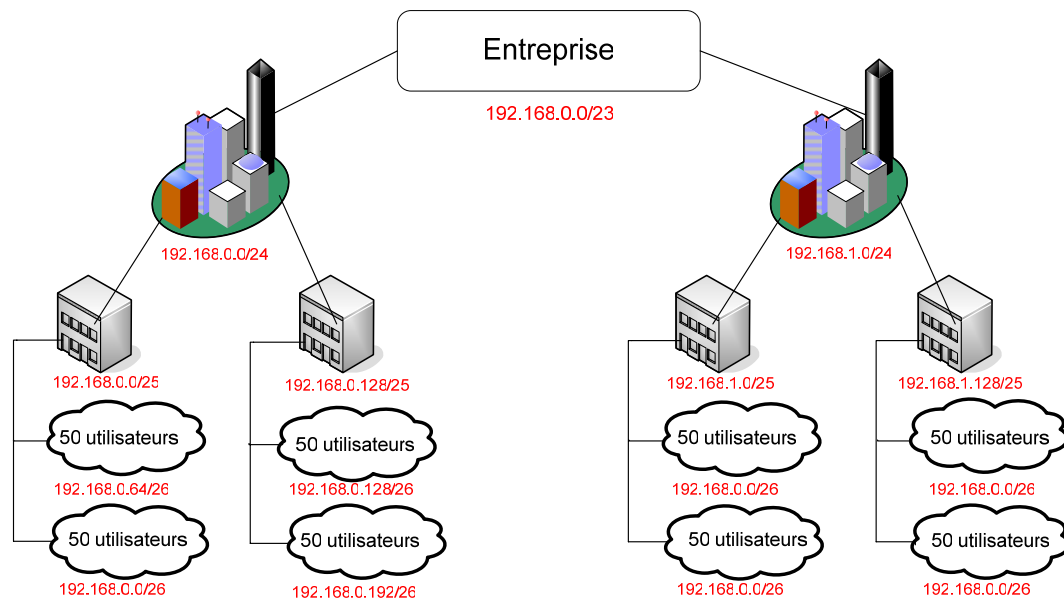
Déterminer la classe d'adresse ou l'agrégat d'adresses (le choix dépendant du contexte), en additionnant tous les bits nécessaires pour identifier chaque niveau hiérarchique de l'entreprise.

- **Etape 5 :**

On procède ensuite au découpage de la classe d'adresse de l'entreprise et de l'attribution à chaque instance du niveau hiérarchique.

Cette procédure est valable quelque soit la méthode d'adressage utilisée (RFC 950 ou 1878) à une différence près, si on applique la règle du $2^n - 1$ ou $2^n - 2$, il faudra l'appliquer une seule fois sur toute la topologie au niveau hiérarchique limitant la perte (induit par le nombre de bits de ce niveau hiérarchique).

Exemple :



Etape 1 : Une entreprise dans deux villes. Deux bâtiments par ville. Deux étages par bâtiment. 50 utilisateurs par étage.

Etape 2 : 50 utilisateurs / sous-réseau +1 adresse pour le broadcast +1 adresse pour le réseau +1 adresse pour la passerelle = 53 adresses IP.

Etape 3 : $2^x \geq 53$ $x=6$ Il faut donc 6 bits par sous-réseau soit un /26 (255.255.255.192)

Etape 4 : Dans ce contexte, on peut découper une classe B (beaucoup de gaspillage) ou agréger plusieurs classe C. On choisira une classe C

Etape 5 : Chaque instance du niveau hiérarchique se voit attribuer un préfixe et un masque. (en rouge sur le dessin)

1.4.2. 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).

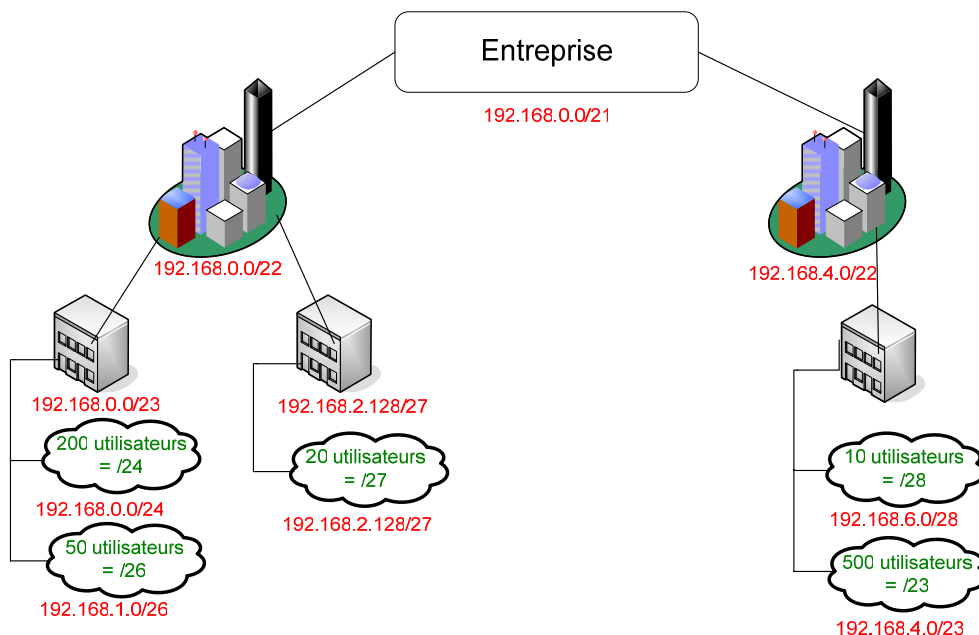
Si le nombre d'utilisateur n'est pas connu à chaque niveau de la hiérarchie, on peut suivre un processus descendant ('top down') : repartir équitablement le nombre d'utilisateur pour un niveau hiérarchique supérieur vers le niveau directement inférieur.

- **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)

1.5. 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**
 - Mode de configuration globale
 - 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**
 - Mode de configuration globale
 - 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}**
 - Mode de configuration globale
 - Permet de passer dans le mode de configuration d'interface
- **ip address {IP} {masque}**
 - Mode de configuration d'interface
 - Permet d'attribuer une adresse IP à cette interface

2. Protocole RIPv2

2.1. Rappels sur RIPv1

RIPv1 est un protocole de routage intérieur classful, à vecteur de distance qui base ses décisions d'acheminement sur une métrique qui emploie essentiellement le nombre de saut. Le nombre maximum de saut est de 15.

- Il transmet des mises à jour de routage complètes toutes les 30 secondes. D'autre part, il lui faut entre 3 et 5 minutes pour converger. Le tableau suivant récapitule les principales caractéristiques de RIPv1 :
- RIPv1 est un protocole de routage intérieur (IGP).
- C'est un protocole de routage à vecteur de distance
- Il utilise une métrique basée sur le nombre de saut.
- Toutes les 30 secondes, il diffuse sa mise à jour de routage par broadcast.
- RIPv1 a une convergence lente.
- Il utilise une métrique de mesure infini (maximum hop count), le split horizon ainsi que des compteurs de retenue (hold down timers) mais aussi le route poisoning pour limiter les effets des boucles de routage.
- RIPv1 est un protocole de routage classful et par conséquent ne supporte pas VLSM et CIDR.

2.2. 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.
- Leur configuration est aisée.

RIPv2 apporte également des fonctionnalités supplémentaires tels que :

- Le support du routage classless.
- La diffusion du masque réseau dans les mises à jour de routage.
- Le support de VLSM.
- 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.

2.3. Configuration

2.3.1. Commandes générales

- **router rip**
 - Mode de configuration globale
 - Active le protocole RIP.
- **version 2**
 - Mode de configuration du protocole de routage
 - Permet d'utiliser RIPv2 à la place de RIPv1
- **network {adresse réseau}**
 - Mode de configuration du protocole de routage
 - Permet d'indiquer les réseaux directement connectés au routeur.
- **ip default-network {adresse réseau}**
 - Mode de configuration du protocole de routage
 - Permet de spécifier une route par défaut.
- **default-information originate**
 - Mode de configuration du protocole de routage
 - Permet de propager la route par défaut dans les mises à jour de routage.
- **no auto-summary**
 - Mode de configuration du protocole de routage
 - Désactive l'auto-agrégation.

2.3.2. Authentication

- **key-chain {nom}**
 - Mode de configuration globale
 - Permet d'identifier un groupe de clef d'authentification.
- **key {id}**
 - Mode de configuration de clé
 - 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}**
 - Mode de configuration de clé
 - Permet de définir un mot de passe pour une clef.
- **ip rip authentication key-chain {nom}**
 - Mode de configuration d'interface
 - Active l'authentification RIP sur une interface
- **ip authentication mode {text | md5}**
 - Mode de configuration d'interface
 - Permet de spécifier le type d'authentification en clair ou crypté.

3. Protocole OSPF

3.1. 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.

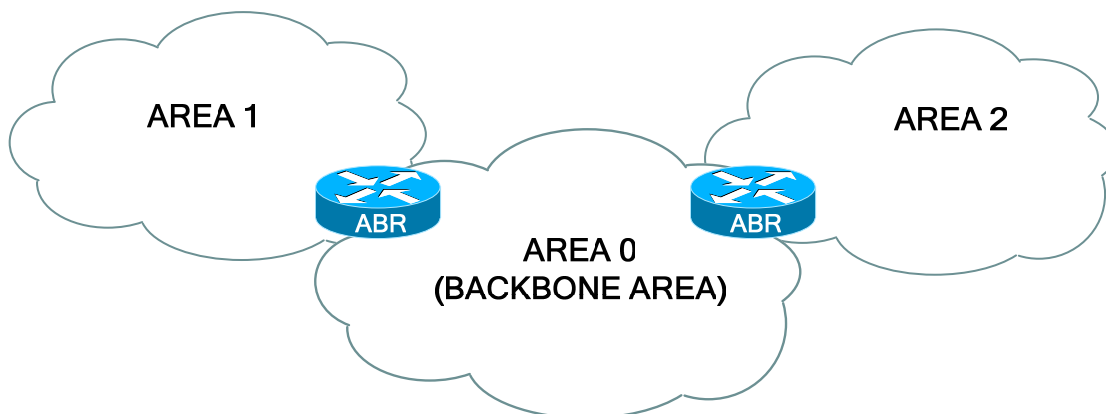
Principales caractéristiques d'OSPF :

- 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**.



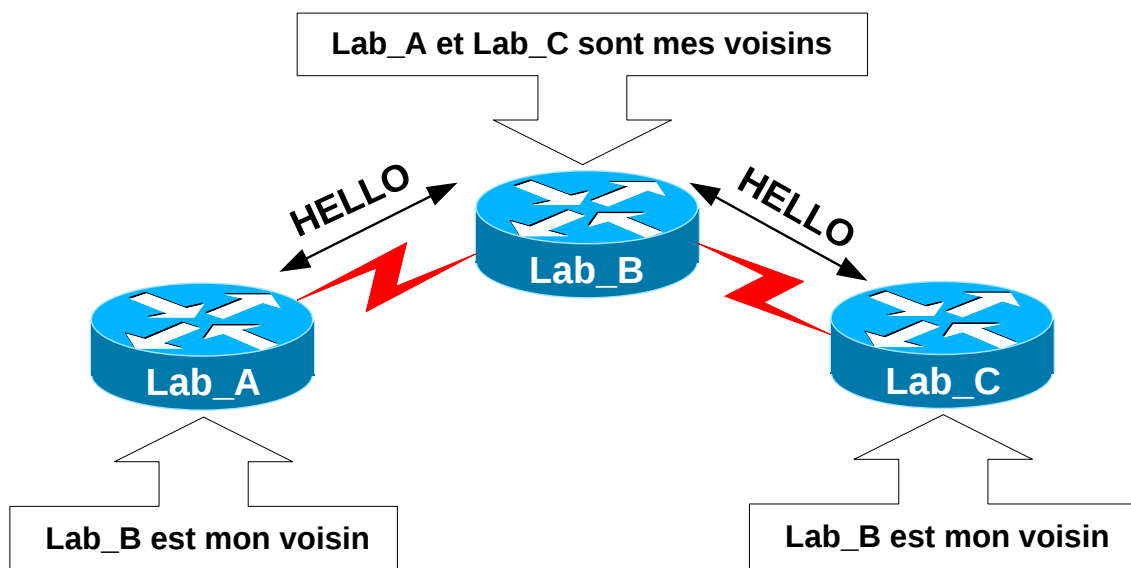
3.2. Définitions

- **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.

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

3.3.1. 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**.

3.3.2. Etablissement des bases de données topologiques

3.3.2.1. Dans un réseau point-à-point

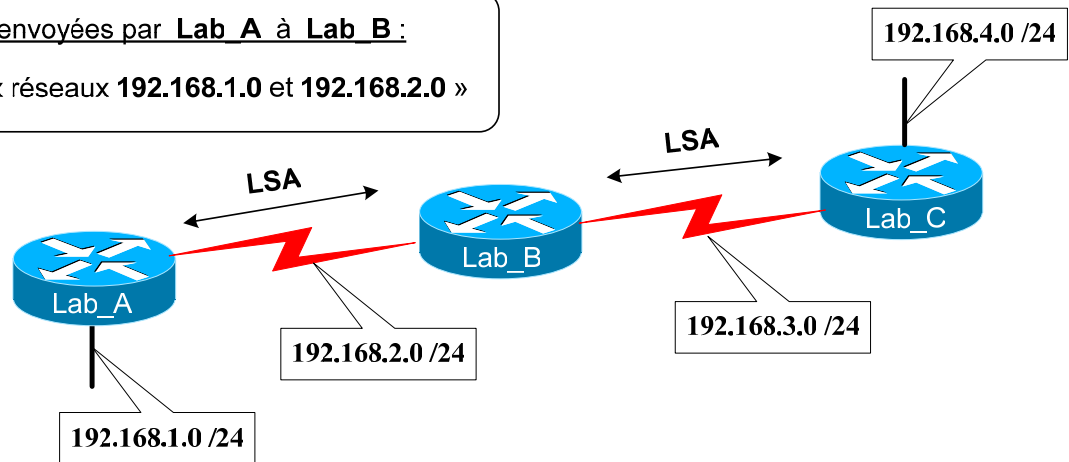
Une fois que chaque routeur a appris l'existence de ses voisins, il va leur envoyer les informations concernant tous les réseaux directement connectés à celui-ci.

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** »



3.3.2.2. Dans un réseau multi-accès

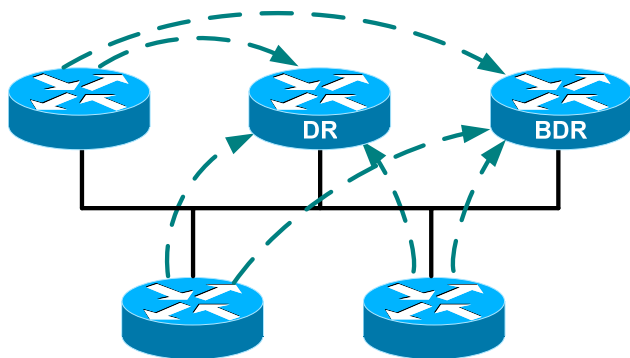
Les réseaux multi-accès fonctionnent suivant le même principe que les réseaux point-à-point à la différence que dans les réseaux multi-accès tous les routeurs sont voisins.

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.

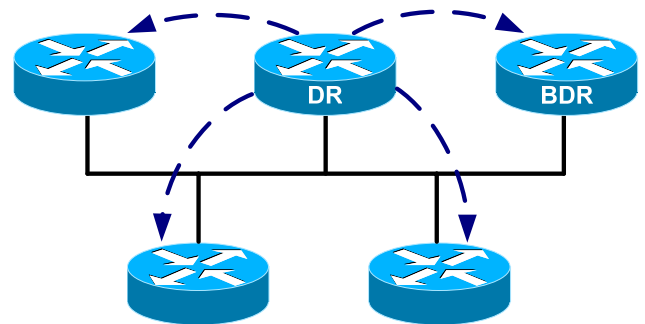
Pour palier à ce problème, le protocole HELLO va élire un **DR** (Designated Router) qui sera chargé de centraliser toutes les informations de modifications topologiques et de les retransmettre par la suite à tous les autres routeurs.

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



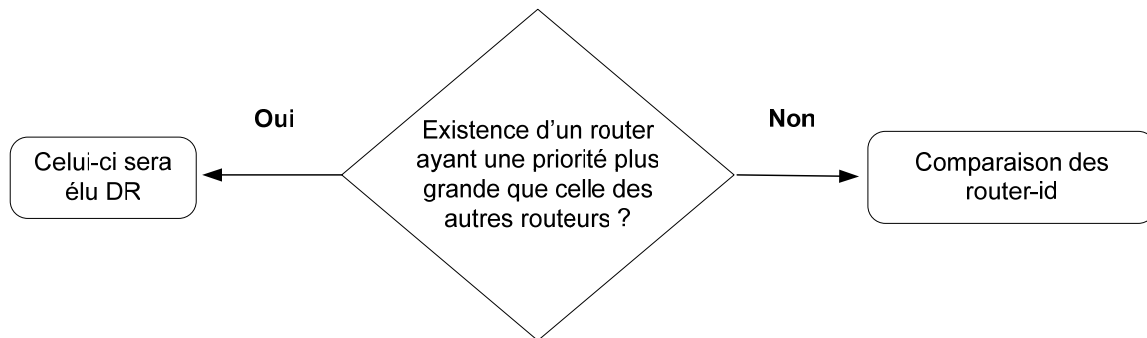
Envoi de paquets LSA à tous les routeurs en utilisant l'adresse multicast 224.0.0.5

3.4. Opérations OSPF

3.4.1. 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.



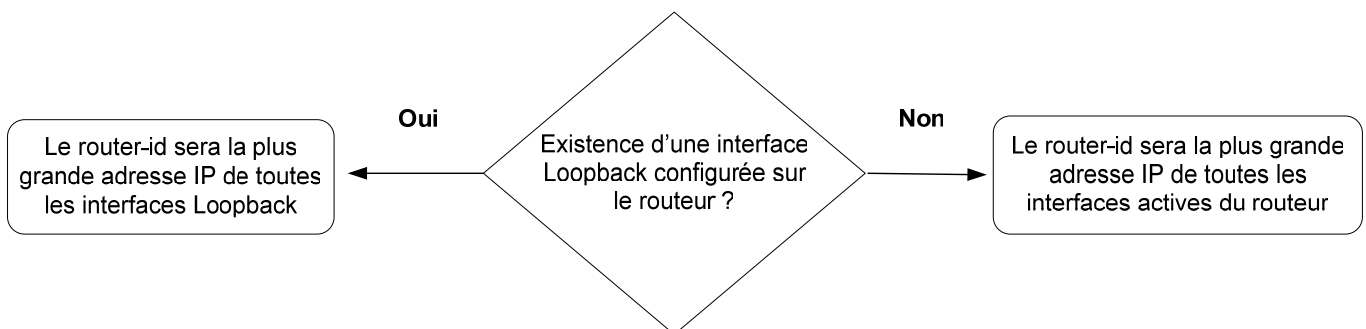
Une fois le DR désigné, le processus d'élection devra ensuite déterminer le BDR, correspondant au routeur ayant la deuxième plus haute priorité ou le deuxième plus grand router-id sur le réseau multi-accès.

3.4.2. 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.

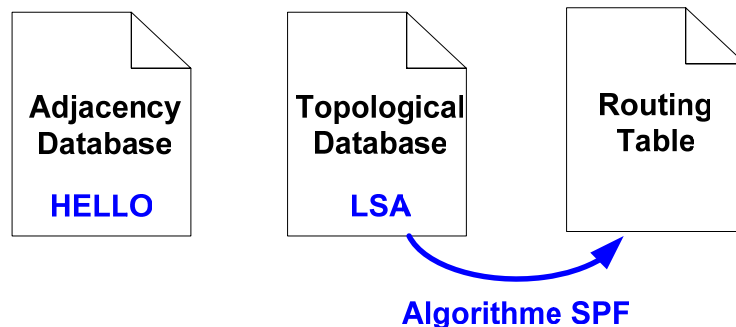
3.5. Construction de la table de routage

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.

3.6. Commandes

3.6.1. Commandes générales

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

3.6.2. Authentification

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

3.6.3. Timers

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

3.6.4. Commandes show associées

- **show ip ospf interface**
 - Mode privilégié
 - Permet d'afficher la priorité de l'interface.
- **show ip protocols**
 - Mode privilégié
 - Affiche les informations sur les protocoles de routage configurés sur le routeur.
- **show ip route**
 - Mode privilégié
 - Affiche la table de routage du routeur.
- **show ip ospf**
 - Mode privilégié
 - 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**
 - Mode privilégié
 - Affiche une liste détaillée des voisins, leur priorité et leur statut.
- **show ip ospf database**
 - Mode privilégié
 - Affiche le contenu de la base de données topologique (router-Id, process-Id).

4. Protocole EIGRP

4.1. Caractéristiques

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

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

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

Les caractéristiques principales d'EIGRP sont :

- Protocole de routage Classless, avec support du VLSM
- Algorithme DUAL
- Mises à jour incrémentales, avec adressage multicast, et de façon fiable (via RTP)
- Utilisation de la bande passante réduite par rapport à IGRP
- Utilisation d'une métrique composite
- Découverte de voisins
- Principe de successeur, avec de multiples FS
- Agrégation de routes manuelle
- Etat des routes (Active et Passive)
- Partage de charge entre chemins n'ayant pas les mêmes métriques
- Compatibilité avec IGRP
- Distance administrative de 90

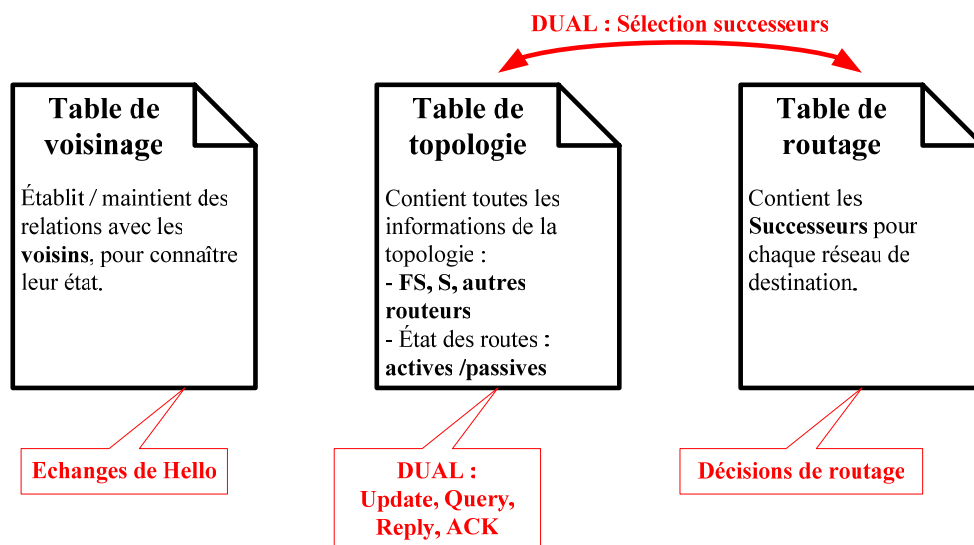
Pour chaque protocole routé utilisé, EIGRP maintient 3 tables distinctes :

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

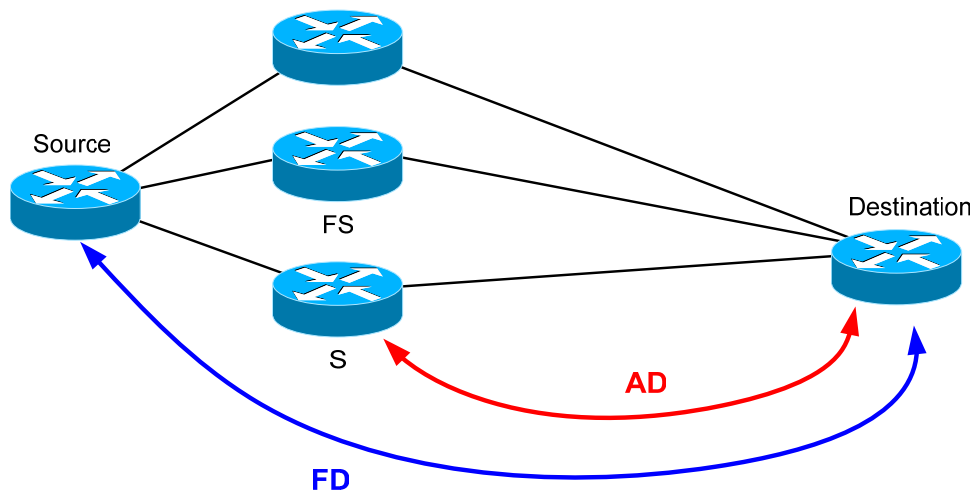
4.2. Termes et définition

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)**
 - Métrique diffusée par un voisin dans sa mise à jour de routage. Elle correspond à la métrique depuis ce voisin, connu localement comme le prochain saut.
- **Reported Distance (RD)**
 - Autre nom pour l'AD.
- **Feasible Distance (FD)**
 - Plus petite métrique pour une destination donnée. C'est la métrique pour la route actuellement dans la table de routage.
- **Feasible Condition (FC)**
 - Condition vérifiée quand un voisin informe une AD plus petite que la FD du routeur local pour une même destination.
- **Feasible Successor (FS)**
 - Voisin vérifiant la FC. Il est potentiellement éligible en tant que successeur.
- **Successor**
 - Voisin utilisé comme prochain saut pour une destination donnée. C'est le FS ayant la plus petite métrique.
- **Stuck In Active (SIA) (aussi appelé Query Scoping)**
 - Etat d'un routeur lorsqu'une route reste active après dépassement d'un certain temps.



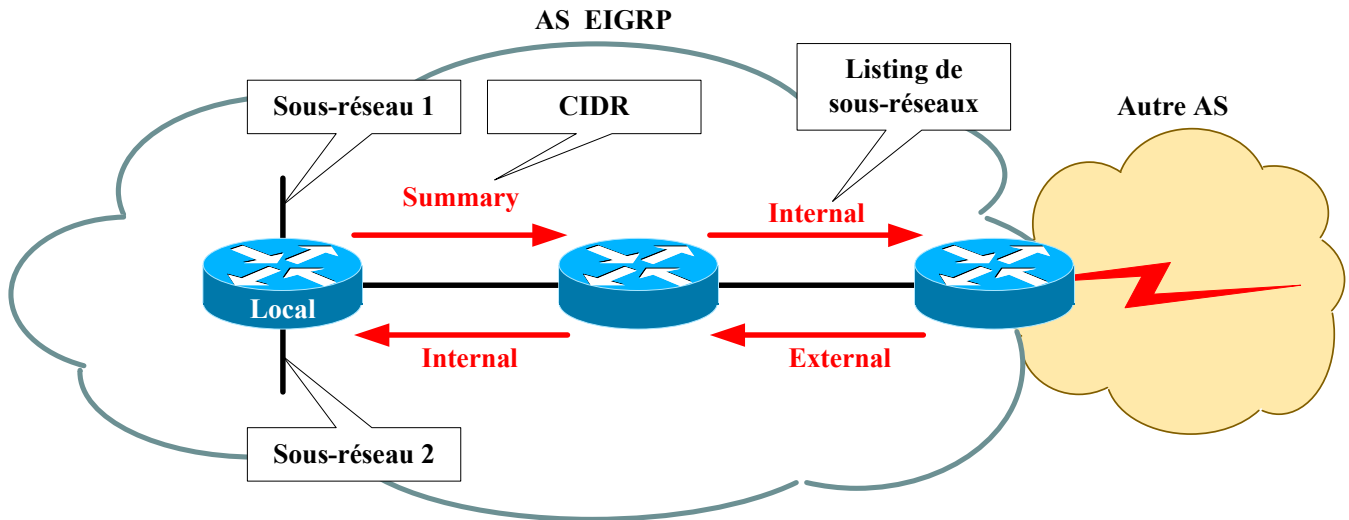
Représentation schématique de quelques définitions

4.3. Métriques

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

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

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



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.

4.4. Protocole Hello

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

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

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

- **Hello**
 - 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.

4.4.1. Neighbor Table

Un routeur est considéré comme voisin si :

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

4.4.2. Topology Table

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

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

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

Elle contient les informations suivantes :

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

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

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

- **Un nouveau réseau est disponible :**
 - o Un paquet Update avertit de l'existence d'un nouveau réseau.
 - o 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 :
 - o Un paquet Reply ou Query est reçu, modifiant ainsi une ou plusieurs entrées dans la table de topologie.
 - o Il y a modification du coût pour une interface locale via configuration.
- **Un réseau devient inaccessible :**
 - o Un paquet Update, Query ou Reply reçu informe la table de topologie qu'un réseau est inaccessible.
 - o Aucun paquet Hello n'est reçu d'un voisin menant à ce réseau avant expiration du Holdtime.
 - o Le réseau est directement connecté et l'interface du routeur perd le signal de portuseuse.

4.5. DUAL

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

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

- **Passive :** Il y a une recherche dans la table de topologie d'une route acceptable pour remplacer l'ancienne présente dans la table de routage :
 - o 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).
 - o Après examen, il existe au moins un FS.
 - o 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 :
 - o 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).
 - o Après examen, il n'existe aucun FS. Le routeur passe en mode actif et envoie des paquets Query à ses voisins.
 - o 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.
 - o Le FS proposant la plus petite AD sera alors choisi comme successeur à l'entrée non valide de l'ancienne table de routage.

4.6. Commandes

Les commandes de configuration d'EIGRP sont les suivantes :

- **router eigrp {n° AS}**
 - o Mode de configuration globale
 - o Active l'algorithme du protocole de routage pour IP.
 - o Permet de passer en mode de configuration de ce protocole de routage.
- **network {réseau} [masque générique]**
 - o Mode de configuration du protocole de routage
 - o 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.
 - o Inclut les informations concernant ces réseaux dans les mises à jour de routage transmises.
 - o 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**
 - o Mode de configuration du protocole de routage
 - o Permet d'activer (par défaut) ou de désactiver l'agrégation de routes automatique aux frontières Classful.
- **ip summary-address eigrp {n° AS} {réseau} {masque}**
 - o Mode de configuration d'interface
 - o Permet de configurer manuellement un agrégat de routes à une frontière Classless.
 - o 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}**
 - o Mode de configuration du protocole de routage
 - o 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.
 - o Le multiplicateur est un entier pouvant aller de 1 (valeur par défaut) à 128.
- **maximum-paths {nombre}**
 - o Mode de configuration du protocole de routage
 - o 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}**
 - o Mode de configuration d'interface
 - o Informe les protocoles de routage utilisant la bande passante pour le calcul des métriques de la véritable bande passante de la liaison.
 - o 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.
 - o Le paramètre **BP** est exprimé en Kbps.

- **passive-interface {type} {numéro}**
 - Mode de configuration du protocole de routage
 - Empêche l'émission et la réception de mises à jour de routage en empêchant la formation d'une relation de voisinage sur l'interface spécifiée.
- **metric weights {TOS} {K1} {K2} {K3} {K4} {K5}**
 - Mode de configuration du protocole de routage
 - Modifie des coefficients entrants en jeu dans le calcul des métriques d'EIGRP.
 - La valeur de **TOS** doit toujours être de 0.

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

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

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

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

4.7. Configuration

La procédure de configuration du protocole EIGRP est la suivante :

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

5. Design de LAN

5.1. Présentation

La conception d'un réseau est un des facteurs les plus importants pour en assurer la stabilité. Les objectifs de cette conception incluent des facteurs tels que :

- **La fonctionnalité**
 - Un réseau doit apporter aux utilisateurs les fonctionnalités suffisantes et nécessaires à leurs besoins
- **L'évolutivité**
 - Un réseau doit pouvoir prendre en charge de nouvelles fonctionnalités sans pour autant devoir reconsidérer la structure initiale
- **L'adaptabilité**
 - Un réseau doit pouvoir s'adapter sans nécessiter de trop complexes configurations
- **La facilité de gestion**
 - Un réseau doit être relativement simple à administrer

Au cours de ce chapitre, nous allons analyser les différents points à observer lors de la conception d'un réseau local. L'analyse portera sur les points suivants :

- Fonctions et emplacements des serveurs
- Détection des collisions (couche 2)
- Segmentation (couche 2 et 3)
- Domaines de broadcast (couche 3)

5.2. Méthodologie de conception

Pour qu'un réseau local soit efficace et réponde aux besoins des utilisateurs, il doit être mis en œuvre selon une suite d'étapes systématiquement planifiées, comprenant notamment les étapes suivantes :

- Le regroupement des besoins et des attentes des utilisateurs
- L'analyse des besoins
- La conception de la structure LAN des couches 1 à 3
- La création de documents sur la mise en œuvre logique et physique du réseau

La première étape de conception d'un réseau consiste à recueillir des données sur la structure de l'organisation. Ces informations comprennent :

- L'historique et l'état en cours de l'organisation
- La croissance prévue
- Les politiques d'exploitation et les procédures de gestion
- Les procédures et les systèmes administratifs ainsi que les points de vue des futurs utilisateurs du réseau local.

Un réseau local est un outil qui sera utilisé par les différents membres de l'entreprise. Le niveau de compétence de ces derniers ainsi que l'utilisation qu'ils comptent faire du réseau sont des éléments déterminants dans la conception.

Ces informations contribuent à identifier et à clarifier les problèmes. Vous devez également déterminer s'il existe des documents sur les politiques déjà en place. Le bon sens et une étude approfondie des besoins des utilisateurs sont les clefs d'un réseau efficace.

Il est également vital de prévoir le rôle des personnes qui vont participer à l'administration du réseau (adressage, maintenance, etc.). Par exemple, la présence d'une tierce entreprise utilisée pour la maintenance est un élément important.

Les ressources d'une organisation pouvant affecter la mise en œuvre d'un nouveau réseau local sont classées en deux catégories : les ressources matérielles/logicielles et les ressources humaines.

Le matériel informatique et les logiciels existants de l'organisation doivent être répertoriés par écrit, et les besoins futurs dans ce domaine doivent être définis. Un rapport écrit sur ces besoins permet d'évaluer les coûts et d'établir un budget pour la mise en place du réseau local. Un schéma présentant la topologie logique du réseau est également un élément important qui permet de bien visualiser le réseau dans son intégralité.

Un schéma logique représente le modèle de la topologie du réseau sans les détails relatifs au parcours d'installation précis des câbles. Il s'agit du plan de base du réseau local. La topologie logique comprend les éléments suivants :

- L'emplacement exact des locaux techniques du répartiteur principal MDF et des répartiteurs intermédiaires IDF.
- Le type et le nombre de câbles utilisés pour interconnecter le répartiteur principal MDF et les répartiteurs intermédiaires IDF ainsi que le nombre de câbles de réserve disponibles pour accroître la bande passante entre les locaux techniques.
- Un document décrivant en détail tous les parcours de câbles, les numéros d'identification et le port de l'interconnexion horizontale ou verticale auquel aboutissent les câbles.

5.3. Fonction et emplacements des serveurs

On distingue 2 types de serveurs :

- Les serveurs d'entreprise :
 - Serveurs dédiés à une application
 - Prennent en charge tous les utilisateurs du réseau (Exemple : DNS, messagerie)
 - Doivent être installés dans le répartiteur principal (MDF)
- Les serveurs de groupes de travail :
 - Offrent des services tels que l'impression ou encore le partage de fichiers
 - Prennent en charge un ensemble spécifique d'utilisateurs
 - Doivent être installés dans les répartiteurs intermédiaires (IDF)

Dans le répartiteur principal MDF et les répartiteurs intermédiaires IDF, les commutateurs LAN de couche 2 liés à ces serveurs doivent avoir un débit minimal de 100 Mbits/s.

5.4. Conception de couche 1

Le câblage physique est l'un des éléments les plus importants à prendre en considération lors de la conception d'un réseau. Les questions relatives à la conception comprennent le type de câble à utiliser (généralement, des câbles de cuivre ou à fibre optique) ainsi que la structure globale du câblage.

Les médias de câblage de couche 1 comprennent le câble à paires torsadées blindées (ou non) de catégorie 5 et le câble à fibre optique, avec la norme TIA/EIA-568-A pour la disposition et la connexion des méthodes de câblage.

En plus des limites de distance, vous devez évaluer avec soin les points forts et les points faibles des diverses topologies, car l'efficacité d'un réseau est directement liée au câblage sous-jacent. Si vous prévoyez d'apporter des modifications importantes à un réseau, il est essentiel d'effectuer une vérification complète des câbles pour identifier les zones qui nécessitent une mise à niveau ou une réinstallation.

Qu'il s'agisse de la conception d'un nouveau réseau ou de la réinstallation du câblage d'un réseau existant, vous devez utiliser des câbles à fibre optique dans le réseau de backbone et le câblage vertical, avec des câbles à paires torsadées blindées (ou non) de catégorie 5 pour le câblage horizontal.

La mise à niveau des câbles doit être prioritaire sur toutes les autres modifications à apporter. En outre, il est impératif de s'assurer, sans exception, que ces systèmes sont conformes aux normes en vigueur.

Dans une topologie en étoile simple comportant un seul local technique, le répartiteur principal MDF comprend un ou plusieurs tableaux d'interconnexions horizontales. Les câbles d'interconnexion horizontale servent à relier le câblage horizontal de la couche 1 aux ports du commutateur LAN de la couche 2.

Le port uplink du commutateur LAN qui, selon le modèle, diffère des autres ports parce qu'il n'est pas interconnecté, est connecté au port Ethernet du routeur de la couche 3 via un câble de raccordement. À ce stade, l'hôte d'extrémité est doté d'une connexion physique complète au port du routeur.

Lorsque des hôtes de grands réseaux dépassent la limite des 100 mètres fixée pour le câble à paires torsadées non blindées de catégorie 5, il n'est pas rare d'installer plusieurs locaux techniques.

La création de plusieurs locaux techniques entraîne la création de plusieurs zones d'interconnexion de réseaux (IDF).

Les normes TIA/EIA568-A précisent que les répartiteurs intermédiaires IDF doivent être connectés au répartiteur principal MDF par le biais d'un câblage vertical appelé câblage de backbone. Une interconnexion verticale permet d'interconnecter les divers répartiteurs intermédiaires IDF au répartiteur principal (MDF).

Comme les câbles verticaux sont en général plus longs que la limite des 100 mètres imposée pour les câbles à paires torsadées non blindées de catégorie 5, le câble à fibre optique est habituellement utilisée.

5.5. Conception de couche 2

L'objectif des équipements de couche 2 est d'assurer la commutation ainsi que la détection des erreurs et la réduction des congestions du réseau. Les deux équipements de couche 2 les plus courants (autres que la carte réseau dont chaque hôte du réseau doit être doté) sont les ponts et les commutateurs LAN. Les équipements de cette couche déterminent la taille des domaines de collision et de broadcast.

Les collisions et la taille du domaine de collision sont deux facteurs qui nuisent aux performances d'un réseau. La commutation LAN permet de micro segmenter le réseau afin d'éliminer les collisions et de réduire la taille des domaines de collision.

Grâce à une autre caractéristique importante, un commutateur LAN peut attribuer la bande passante par port, ce qui laisse davantage de bande passante aux câbles verticaux, aux liaisons montantes (uplinks) et aux serveurs.

Si vous installez un commutateur LAN au répartiteur principal MDF et aux répartiteurs intermédiaires IDF ainsi qu'un câble vertical entre le répartiteur principal et les répartiteurs intermédiaires, le câble vertical acheminera tout le trafic de données entre le répartiteur principal et les répartiteurs intermédiaires.

La capacité de ce parcours doit être supérieure à celle des parcours reliant les répartiteurs intermédiaires IDF et les stations de travail. Les câbles horizontaux utilisent des paires torsadées non blindées de catégorie 5 et aucun branchement de câble ne doit dépasser 100 mètres de longueur de façon à obtenir des liaisons à des débits de 10 Mbits/s ou de 100 Mbits/s. Dans un environnement normal, un débit de 10 Mbits/s convient pour le câble de branchement horizontal.

Comme les commutateurs LAN asymétriques permettent de combiner des ports à 10 Mbits/s et à 100 Mbits/s sur un même commutateur, l'étape suivante consiste à déterminer le nombre de ports à 10 Mbits/s et à 100 Mbits/s nécessaires pour le répartiteur principal MDF et pour chacun des répartiteurs intermédiaires IDF.

Vous pouvez déterminer ce nombre en consultant les besoins des utilisateurs spécifiant le nombre de câbles de branchement horizontaux par salle dans chaque zone d'interconnexion de réseaux ainsi que le nombre de câbles verticaux.

L'autre méthode permettant de mettre en œuvre une commutation LAN consiste à installer des concentrateurs LAN partagés sur les ports du commutateur et de connecter plusieurs hôtes à un seul port du commutateur. Tous les hôtes connectés au concentrateur LAN partagé partagent le même domaine de collision et la même bande passante.

Les concentrateurs à média partagé sont généralement utilisés dans un environnement de commutateurs LAN pour créer davantage de points de connexion à l'extrémité des câbles horizontaux.

Cette solution est acceptable, mais vous devez vous assurer que la taille des domaines de collision n'augmente pas et que les besoins de l'hôte en matière de bande passante respectent les spécifications définies à l'étape des besoins du processus de conception du réseau.

5.6. Conception de couche 3

Les équipements de couche 3, tels que les routeurs, peuvent être utilisés pour créer des segments LAN uniques et permettre la communication entre les segments sur la base de l'adressage de couche 3, tel que l'adressage IP. La mise en œuvre des équipements de couche 3, tels que les routeurs, permettent de segmenter le réseau local en réseaux physiques et logiques uniques.

Les routeurs fournissent également la connectivité aux réseaux WAN tels qu'Internet. Le routage de couche 3 détermine également le flux du trafic entre les segments physiques uniques du réseau en fonction de l'adressage de couche 3 (par exemple, un réseau IP ou un sous-réseau).

Le nombre total de broadcasts, tels que les requêtes ARP, est une question importante dans un réseau. Grâce aux VLAN, vous pouvez limiter le trafic de broadcast au sein de chaque VLAN et, par conséquent, créer des domaines de broadcast plus petits.

Les VLAN permettent également de sécuriser le réseau en créant des groupes de VLAN selon leur fonction. Une association à un port physique est utilisée pour mettre en œuvre l'attribution de VLAN statiques. Comme le routeur détermine si le réseau VLAN 1 peut communiquer avec le réseau VLAN 2, vous pouvez créer un système de sécurité fondé sur l'attribution des VLAN.

Les routeurs fournissent une évolutivité au réseau parce qu'ils servent de pare-feu vis-à-vis des broadcasts. De plus, comme les adresses de couche 3 ont généralement une structure, ils accroissent l'évolutivité en divisant les réseaux et les sous-réseaux, ce qui renforce la structure de ces adresses.

Une fois les réseaux divisés en sous-réseaux, l'étape finale consiste à développer et à expliquer par écrit le système d'adressage IP à utiliser. La technologie de routage filtre les broadcasts et les multicasts de liaison de données. En ajoutant des ports de routeur ainsi que des adresses réseau ou de sous-réseau, vous pouvez, si nécessaire, segmenter l'inter réseau.

Les routeurs permettent de créer des sous-réseaux IP pour renforcer la structure des adresses. Avec des ponts et des commutateurs, toutes les adresses inconnues encombrant chaque port doivent être évacuées.

Avec des routeurs, les hôtes utilisant des protocoles d'adressage de couche réseau peuvent résoudre la recherche d'hôtes sans provoquer d'encombrement réseau :

- Si l'adresse de destination est locale, l'hôte émetteur peut encapsuler le paquet dans un en-tête de liaison de données et transmettre une trame d'unicast directement à la station. Le routeur ne voit pas la trame et, bien sûr, n'a pas besoin de la traiter. L'hôte émetteur peut utiliser une requête ARP. Dans ce cas, un broadcast est généré. Cependant, comme il s'agit d'un broadcast local, le routeur ne le transmet pas.
- Si la destination n'est pas locale, la station émettrice transmet le paquet au routeur. Le routeur envoie la trame à destination ou au saut suivant en fonction de sa table de routage.

En raison de cette fonctionnalité de routage, il est évident que les grands réseaux locaux évolutifs doivent comporter quelques routeurs.

6. Commutation

6.1. Concepts et fonctionnement

Au début des LAN, les équipements de réseau utilisaient un seul bus électrique. En effet, tous les équipements du LAN partageaient la bande passante d'un seul bus. C'est le cas des normes Ethernet 10Base2, 10Base5 et 10Base-T.

Dans une réunion, quand plusieurs personnes prennent la parole en même temps, cela crée une cacophonie et il devient très difficile, voire impossible de comprendre les interlocuteurs. Il faut donc appliquer une convention afin qu'il n'y ait qu'une personne à la fois qui prenne la parole.

La même problématique se retrouve dans les LAN où les équipements du réseau se partagent le même espace de « discussion ». Quand 2 hôtes envoient un signal en même temps, ceux-ci se chevauchent rendant impossible leur interprétation : on parle de collision.

Pour résoudre ce problème, l'algorithme **CSMA/CD** (Carrier Sense Multiple Access with Collision Detection) est appliqué et définit de quelle façon accéder au bus. Malgré l'apport de CSMA/CD l'utilisation du réseau n'était pas optimale. Les LAN étaient confrontés aux problèmes de collisions, congestions, latence et de remise de données de type broadcast.

Les ponts transparents, puis les commutateurs (ou **switch**) dans un second temps permirent de résoudre ces phénomènes.

Les ponts offrent principalement les avantages suivants:

- La réduction de la taille des domaines de collisions par la segmentation.
- L'augmentation de la bande passante (due à la réduction de la taille des domaines de collision).

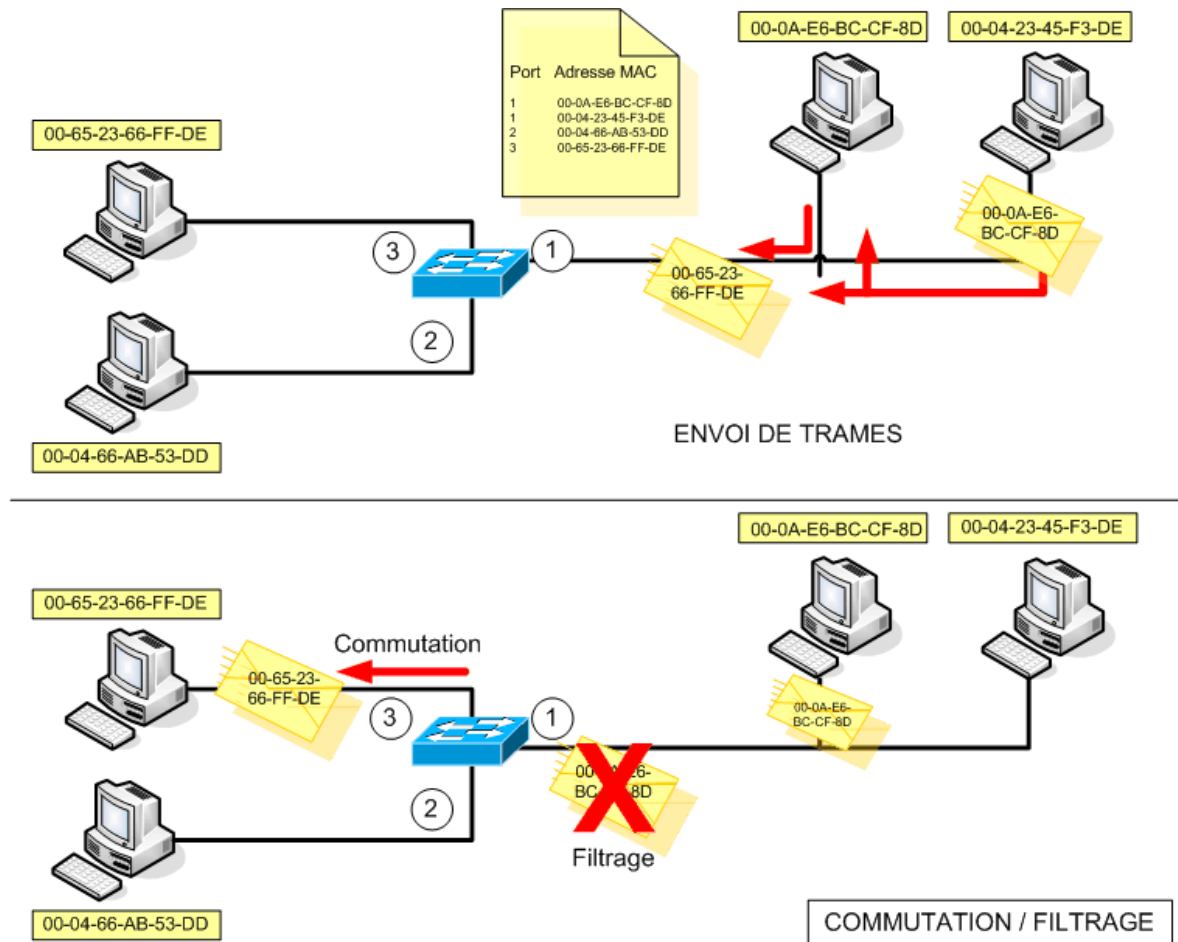
Les ponts et les commutateurs ont la même logique de fonctionnement, d'ailleurs un commutateur est un pont multiport.

Contrairement à un concentrateur qui se contente de régénérer, resynchroniser et retransmettre les bits sur le média, le pont est capable de prendre des décisions d'acheminement. Pour cela, il utilise les adresses **MAC** (Media Access Control). De ce fait, les ponts, comme les commutateurs, sont des équipements de couche 2 du modèle OSI.

Quand un pont reçoit une trame, il examine l'adresse MAC (Media Access Control) de destination et détermine s'il doit filtrer ou transmettre la trame. Les décisions d'acheminements se fondent sur une table de pontage où le pont inscrit toutes les adresses MAC et le port par lequel elles sont arrivées.

Quand une trame arrive à un port, le pont va consulter sa table de pontage pour connaître le port par lequel la trame doit être envoyée à l'adresse MAC de destination. Dans ce cas, si :

- Le port de destination est le même que celui qui a reçu la trame, la trame ne sera pas transmise sur d'autres ports : c'est le **filtrage**.
- Le port de destination est un port différent de celui par lequel la trame a été reçue, le pont transmet cette trame sur le port correspondant : c'est la **commutation**.



Le fonctionnement du commutateur est identique. L'apprentissage des adresses MAC se déroule comme suit :

- Lorsque le commutateur reçoit une trame par un de ses ports, il inscrit dans sa table de pontage la correspondance entre l'adresse MAC source et le port associé.
- Supposons que la table de pontage soit déjà créée et complète. Le commutateur examine l'adresse physique de destination de la trame reçue et cherche dans sa table l'entrée associée à l'adresse.
- Une fois le port de destination identifié, le commutateur commute la trame sur le port correspondant.

S'il n'y a pas d'entrée dans la table de pontage, le commutateur crée une entrée correspondante et transmet les données par tous ses ports excepté le port source. Quand le destinataire répondra à l'émetteur, le commutateur pourra inscrire l'entrée correspondante.

L'IEEE a défini trois catégories d'adresses MAC :

- Adresse **unicast** : adresse physique identifiant une seule carte réseau.
- Adresse de **broadcast** : avec cette adresse tous les noeuds du LAN doivent traiter la trame. L'adresse de broadcast a pour valeur FFFF.FFFF.FFFF
- Adresse **multicast**. Permet à un ensemble de noeuds de communiquer entre eux. L'adresse multicast a pour valeur 0100.5Exx.xxxx où x peut prendre n'importe quelle valeur.

En transmettant les trames reçues à un autre port, le commutateur crée un bus unique entre la source et la destination (micro segmentation). L'utilisation de la bande passante est optimale, 100% de la bande passante est utilisée.

L'algorithme CSMA/CD n'est plus employé car il n'y a pas de collision. On peut alors utiliser le mode de fonctionnement full-duplex, c'est-à-dire que la source et la destination peuvent émettre et recevoir en même temps.

6.2. Commutateurs

6.2.1. Présentation

Un commutateur est un équipement réseau de couche 2. Il en existe une grande variété avec des caractéristiques différentes :

- Nombre de ports
- Type de port (10/100 Mbits, gigabit)
- Type de commutation (Store and Forward, Cut Through)
- Facilité d'installation en armoire etc...

Les différents types de commutation :

- **Store and forward:** Le commutateur attend d'avoir reçu toute la trame avant de la transmettre. Cette méthode offre une grande vérification d'erreur car le commutateur a le temps de vérifier la valeur FCS. 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:
 - **Fragment Free:** Filtrage des fragments de collision (inférieur à 64 octets). Le commutateur attend d'avoir reçu les 64 premiers octets avant de commencer à transmettre la trame. La détection des collisions subies doit être détectée au niveau des 64 premiers octets.
 - **Fast Forward:** Pas de vérification d'erreurs. La trame est transmise dès que l'adresse de destination est identifiée.

6.2.2. Démarrage

Avant le démarrage du système d'exploitation une procédure POST (Power On Self Test) est lancée pour tester le bon état du matériel.

Le voyant indique l'échec ou la réussite du POST : une lumière ambre indique l'échec, alors qu'une couleur verte indique que la procédure s'est terminée avec succès.

6.2.3. Configuration de base

Pour configurer un commutateur il convient de se connecter via le port console à l'aide d'un câble du même nom. Une fois la connexion lancée, on se retrouve sur une interface de ligne de commande : la CLI (Command-Line Interface).

A l'instar de l'IOS des routeurs, il existe différents modes de configuration : le mode utilisateur, le mode privilégié et le mode de configuration globale. Les mêmes commandes sont utilisées pour accéder à ces différents modes.

6.2.4. Voyants d'un commutateur

Voyant	Etat et signification	
Système	Voyant éteint : le système est hors tension.	
	Voyant vert : le système est sous-tension.	
	Voyant ambre : problème suite au POST.	
RPS (Remote Power Supply)	Ce voyant indique si l'alimentation de sécurité est utilisée.	
Port	Chaque port a son voyant qui donne des indications sur l'état du port selon le mode choisi.	
Bouton mode	Permet de choisir entre les 4 modes: Stat, Util, Duplex et Speed.	
Bouton mode	Stat	Donne des informations sur l'état des ports. Une lumière verte indique que le port est opérationnel. Quand elle clignote elle témoigne d'une activité. Si la lumière est éteinte le port est non opérationnel.
	Util	Ce mode utilise l'ensemble des voyants de ports pour donner des informations sur l'utilisation générale du commutateur.
	Duplex	Quand le voyant est allumé le port fonctionne en mode full duplex. Eteint, c'est le mode half duplex qui est employé.
	Speed	Un voyant allumé indique un débit de 100 Mbits, un voyant éteint un débit de 10Mbits.



Face avant et arrière d'un commutateur Cisco Catalyst 2950

6.2.5. Commandes

- **enable**
 - Depuis le mode utilisateur permet d'accéder au mode privilégié.
- **configure terminal**
 - Depuis le mode privilégié permet d'accéder au mode de configuration globale.
- **show version**
 - Permet de vérifier la version de l'IOS et la valeur du registre de configuration.
- **show running-config**
 - Permet d'afficher le fichier de configuration actif.
- **show interface FastEthernet [numéro de l'interface]**
 - Affiche le statut de l'interface, le débit, l'auto négociation et les statistiques de l'interface.
- **show flash ou dir:flash**
 - Affiche la version de l'image de l'IOS contenue dans la mémoire flash, la taille de la mémoire et la mémoire utilisée.
- **show interface status**
 - Affiche le mode opérationnel du port.
- **show controllers ethernet-controller**
 - Affiche les statistiques sur les données reçues et envoyées au niveau matériel.
- **show post**
 - Indique si le routeur a effectué le POST.
- **reload**
 - Redémarre le commutateur.
- **erase startup-config**
 - Efface le fichier de configuration de sauvegarde.
- **delete flash:vlan.dat**
 - Supprime la base de donnée de VLAN. Sur les Catalyst 1900 c'est la commande delete nvram qui est employée.
- **show mac-address-table**
 - Permet d'afficher les adresses MAC apprises par le commutateur.
- **clear mac-address-table**
 - Permet d'effacer les entrées de tables configurées par l'administrateur.
- **mac-address table static [adresse MAC de l'hôte] interface Fast Ethernet [numéro de l'interface] vlan [numéro du vlan]**
 - Permet d'attribuer une adresse MAC statique à une interface.
- **show port security**
 - Permet de vérifier le statut de sécurité appliqué aux ports.
- **interface [type] [numéro/sous numéro]**
 - Permet de passer dans le mode configuration de l'interface.
- **interface range [type] [numéro/premier numéro – dernier numéro]**
 - Permet de passer dans le mode de configuration de plusieurs interfaces.

6.2.6. Procédure de récupération des mots de passe

- Appuyez sur le bouton mode en même temps que la mise sous tension du commutateur.
- Pour initialiser la flash tapez **flash_init**, puis **load_helper** et enfin **dir:flash**.
- Ensuite renommez le fichier de configuration avec la commande **rename flash: config.txt flash: config.old**.

Le fichier de configuration ne sera pas chargé au prochain démarrage du commutateur.

Pour des raisons d'espace, pensez à supprimer le fichier config.old avec la commande : **delete flash: config.old**.

Laboratoire SUPINFO des Technologies Cisco

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

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

6.3. Protocole Spanning-Tree

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.

6.3.1. 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.

Un réseau interconnecté est dit convergent lorsque tous les ports ont pris un état de blocage ou de transmission. Le processus de convergence prend 15 secondes pour le processus d'écoute, plus 15 secondes pour le processus de découverte et 20 secondes pour bloquer les ports ou les mettre dans l'état de transmission.

Lorsqu'une modification topologique est détectée l'arbre est recalculé et le trafic ne reprend totalement qu'après le temps de convergence nécessaire.

6.3.2. 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.

6.3.3. Commandes et configuration de Spanning-Tree

- **spanning-tree {identifiant de vlan} root**
 - Depuis le mode de configuration globale, permet de désigner le commutateur racine. A l'issue de cette commande la priorité du commutateur sera modifiée pour être plus basse que le commutateur qui devrait être racine.
- **spanning-tree {identifiant de vlan} [priority priorité]**
 - Depuis le mode de configuration globale, permet de changer le niveau de priorité.
- **spanning-tree cost {coût}**
 - Depuis le mode de configuration spécifique de l'interface, permet de modifier le coût STP.
- **channel-group {numéro du groupe de canal} mode [auto | desirable | on]**
 - Active l' Etherchannel (agrégation de liens) de l'interface.
- **show spanning-tree**
 - Affiche des informations détaillées sur le protocole STP en cours ainsi que l'état de chaque port.
- **show spanning-tree interface {interface}**
 - Affiche les informations Spanning-tree du port spécifié.
- **show spanning-tree vlan {vlan id}**
 - Affiche les informations Spanning-tree du VLAN spécifié.
- **debug spanning-tree**
 - Affiche les informations de changement topologique STP.
- **show etherchannel {numéro de groupe de canal} [brief | detail | port | port-channel | summary]**
 - Affiche les informations sur le statut des EtherChannels sur le commutateur.

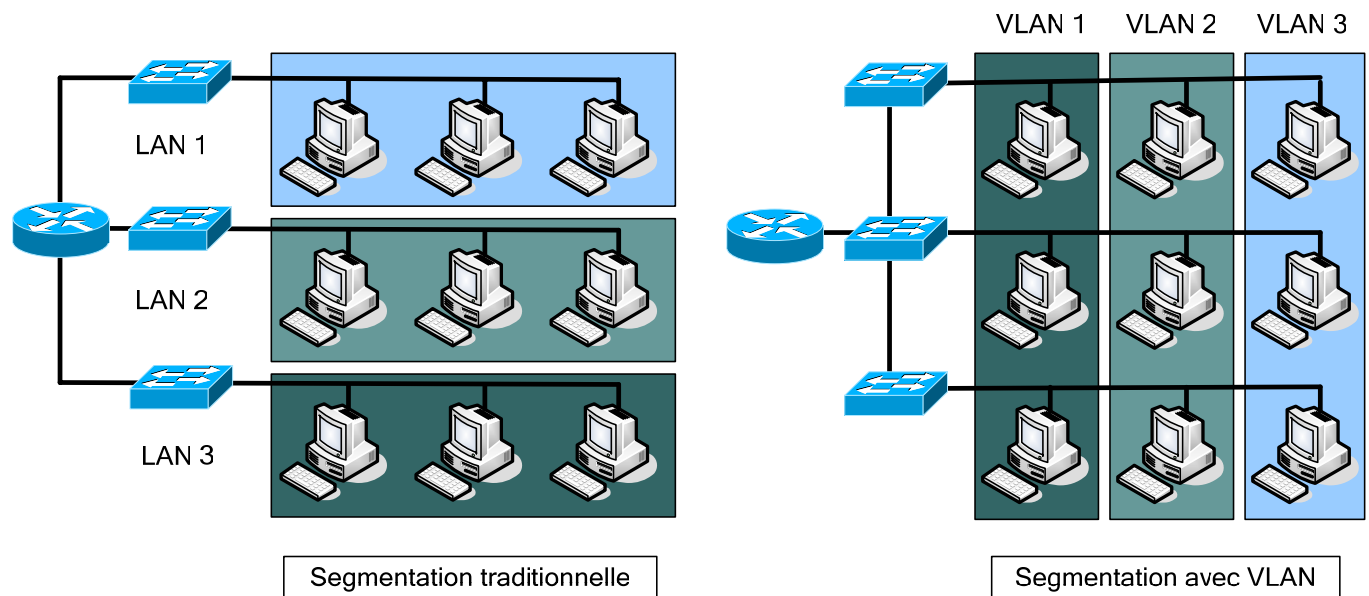
6.4. VLAN

6.4.1. 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 statiques** : 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 dynamiques** : 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 donnée liant MAC et VLAN (remplie par l'administrateur), renverra alors le VLAN d'appartenance au commutateur.

6.4.2. 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.

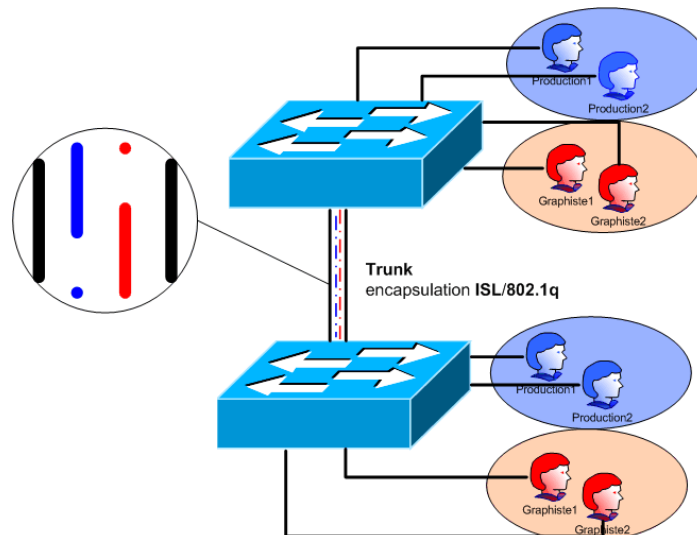
6.4.3. Commandes show associées

- **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.

6.4.4. 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**.

6.5. 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.

6.5.1. 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.

6.5.2. 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
----------	---------

6.5.3. 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.	

6.5.4. Commandes associées

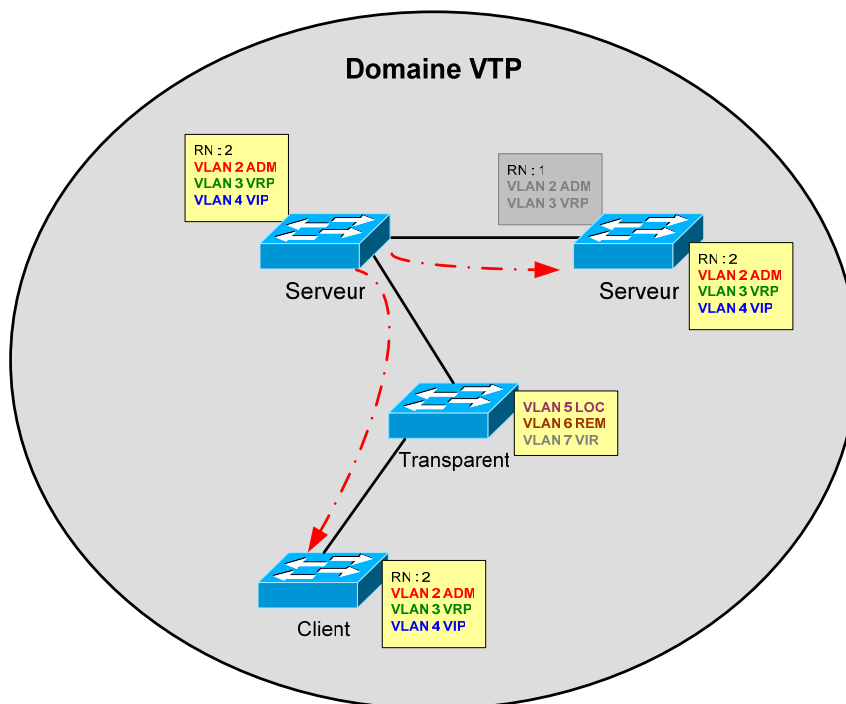
- **switchport mode trunk**
 - Depuis le mode de configuration spécifique du port, active le trunking.
- **switchport trunk [allowed | encapsulation | native | pruning]**
 - Quand l'interface est en mode trunking permet respectivement, de spécifier le type de caractéristiques VLAN, le type d'encapsulation (ISL, 802.1q), les caractéristiques natives et les caractéristiques de pruning des VLANs.
- **show port capabilities [numéro/sous-numéro]**
 - Affiche les fonctionnalités supportées par l'interface.
- **show trunk**
 - Permet de vérifier la configuration du trunking.

6.6. VTP

6.6.1. Théorie sur le protocole VTP

VTP (Virtual 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'ait 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.

6.6.2. 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.

Table des matières

1. NAT et PAT	4
1.1. Adressage privé et public	4
1.2. Translation d'adresses	4
1.2.1. Principe du NAT	4
1.2.2. Principe du PAT	6
1.3. Configuration	6
1.3.1. Commandes	6
1.3.2. Procédure de configuration	7
1.3.3. Vérification	7
2. Protocole DHCP	8
2.1. Introduction	8
2.1.1. Comparatif entre BOOTP et DHCP	8
2.1.2. Opération DHCP	9
2.1.3. Relais DHCP	10
2.2. Configuration	11
2.2.1. Commandes	11
2.2.2. Procédure de configuration	12
2.2.3. Vérification	12
3. Réseaux WAN	13
3.1. Définitions	13
3.2. Equipements et dispositifs	14
3.3. Normes WAN	15
3.4. Classement des différents types de liaison WAN	17
4. Conception WAN	19
4.1. Communication dans un WAN	19
4.2. Premières étapes de la conception WAN	19
4.3. Modèle de réseau hiérarchique	20
4.3.1. Modèle à 3 couche	21
4.3.2. Modèle à 2 couche	22
4.3.3. Modèle à 1 couche	23
5. Protocole PPP	24
5.1. Etude du protocole	24
5.2. Etablissement d'une session	25
5.3. Authentification/Configuration	25
5.3.1. Procédure de configuration du protocole PAP	27
5.3.2. Procédure de configuration du protocole CHAP	27
6. Technologies RNIS	29
6.1. Technologie	29
6.2. Termes et équipements	30
6.3. Normes	31
6.4. Utilisation/Implémentation	32
6.5. Routage à établissement de la connexion à la demande (DDR)	33
6.6. Commandes	34
6.7. Configuration	36

7. Technologies Frame Relay	37
7.1. Technologie	37
7.2. Interface LMI & DLCI	38
7.3. Fonctionnement, table de commutation et processus de transmission.....	39
7.4. Les sous interfaces Frame Relay.....	41
7.5. Commandes	42
7.6. Configuration.....	44
8. Initiation à l'administration réseau	45
8.1. Stations de travail et serveurs	45
8.1.1. Stations de travail	45
8.1.2. Serveurs.....	45
8.2. Systèmes d'exploitation réseau.....	45
8.2.1. Systèmes d'exploitation réseau Microsoft Windows	46
8.2.2. Systèmes d'exploitation réseau UNIX et Linux.....	46
8.2.3. Système d'exploitation réseau Apple	47
8.3. Gestion du réseau.....	47
8.3.1. Introduction à la gestion réseau.....	47
8.3.2. Modèle de gestion réseau et OSI.....	48
8.4. Protocole SNMP	49
8.4.1. Introduction	49
8.4.2. Fonctionnement	49
8.4.3. MIB	52
8.4.4. Configuration	53
8.4.5. RMON.....	54
8.5. Syslog	57
8.5.1. Fonctionnement.....	57
8.5.2. Configuration	58

1. NAT et PAT

1.1. 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.

1.2. 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)

1.2.1. 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.

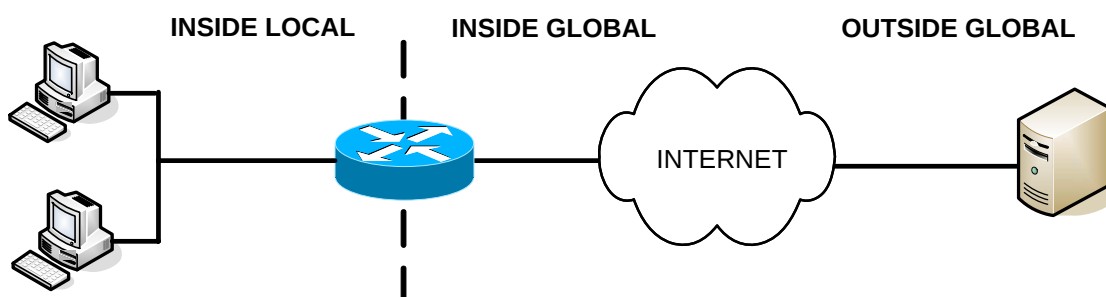
Cette translation d'adresse est effectuée principalement sur les routeurs de bordure d'une entreprise connectée à Internet. Le réseau utilisant les adresses IP privées est ainsi appelé le réseau interne (**inside**), tandis que la partie du réseau utilisant des adresses IP publiques (Internet) est appelé 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**
 - Adresse IP attribuée à un hôte dans le LAN.
- **Inside global address**
 - Adresse(s) IP attribuée(s) par le FAI reconnue(s) par l'Internet pour représenter le LAN.
- **Outside local address**
 - 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**
 - 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.

1.2.2. 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.

1.3. Configuration

1.3.1. Commandes

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

1.3.2. 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.

1.3.3. 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.

2. Protocole DHCP

2.1. 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.

2.1.1. 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).

Laboratoire SUPINFO des Technologies Cisco

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

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

2.1.2. 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 :**

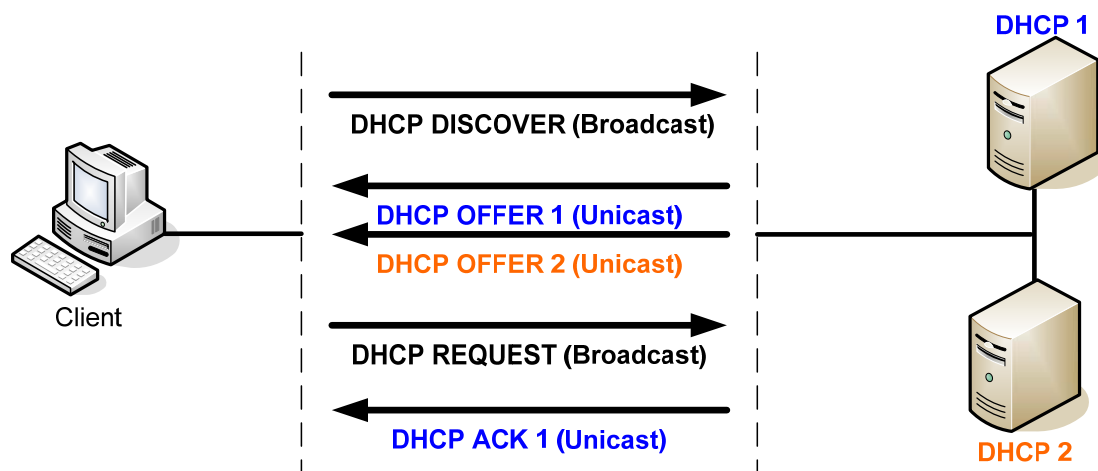
- Les serveurs DHCP recevant le broadcast et pouvant répondre à la demande, envoient une requête en unicast au client. Ce DHCP OFFER contient toutes les informations nécessaires au client (IP, adresse de passerelle, durée du bail, serveur DNS, WINS, etc.).

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.

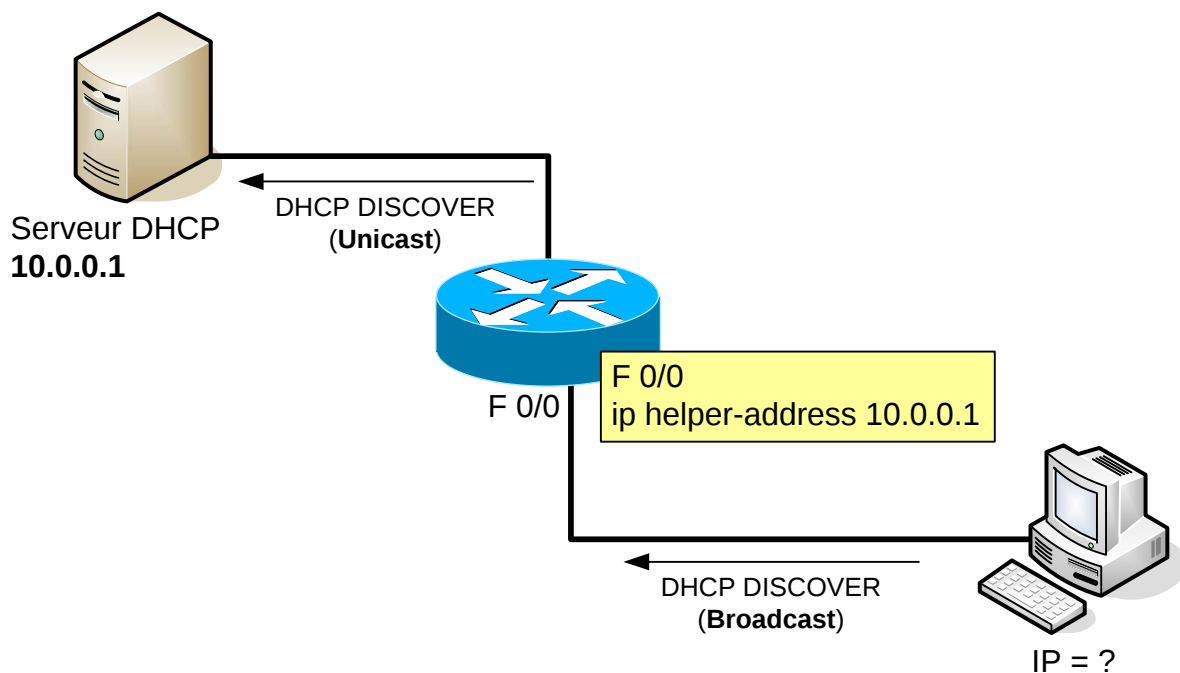
2.1.3. Relais DHCP

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



2.2. Configuration

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

2.2.1. 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.

2.2.2. 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`)

2.2.3. 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

3. Réseaux WAN

3.1. Définitions

Caractéristiques principales des réseaux WAN :

- Fonctionnent sur de vastes étendues géographiques.
- Utilisent les services d'un opérateur Télécom.
- Transportent différents types de trafic (Voix, données, vidéo).
- Axés sur les couches physique et liaison de données du modèle OSI.

La boucle locale est la partie située entre le POP du client et le central téléphonique de l'opérateur.

Un réseau WAN, d'un point de vue général, est un ensemble de liaisons reliées aux différents opérateurs, qui sont interconnectés.

Le rôle des opérateurs Télécom est de fournir une communication bout à bout, en utilisant diverses méthodes de commutation (circuits, paquets, cellules), tout en fournissant des services.

Les trois grands types de services fournis par un opérateur Télécom sont :

Établissement de la communication :

- Aussi appelé signalisation, ce service permet d'établir ou de mettre fin à la communication entre les utilisateurs du système téléphonique.

Transit des données :

- **Multiplexage temporel** : Principe simple qui permet d'allouer l'intégralité de la bande passante disponible d'une liaison par tranche de temps fixe, affectée à chaque utilisateur.
- **Partage de bande passante** : Il existe une bande passante totale disponible sur le backbone, et les clients qui y sont rattachés se la partagent.

Le chemin de réseau WAN reliant les ETDD est appelé :

- Liaison.
- Circuit.
- Canal.
- Ligne.

Le but principal de l'ETCD est de servir d'interface entre l'ETDD et la liaison de communication WAN de l'opérateur :

L'ETDD fournit les données de l'utilisateur (Exemple : routeur).

L'ETCD convertit le format des données de l'utilisateur en un format acceptable par les unités du service réseau WAN (Exemple : modem, unité CSU/DSU, TA, NT1).

Il existe deux types de circuits :

- **Circuit point-à-point** : Circuit physique dédié aux deux extrémités (Exemple : Circuit POTS ou RNIS une fois la commutation de circuits effectuée).
- **Circuit virtuel** : Circuit logique passant au travers d'un nuage (Exemple : Frame Relay, X.25).

Les circuits virtuels se découpent en deux catégories :

- **SVC :**

Établi dynamiquement sur demande et fermé en fin de transmission.

Communication en trois phases : Etablissement du circuit, transfert des données et fermeture du circuit.

Consomme de la bande passante à cause des différentes phases de la communication.

Coûts liés à la disponibilité (Temps) du circuit réduit.

- **PVC :**

Établi en permanence.

Est utilisé pour transmettre des débits de données constantes.

Communication en une phase : Transfert des données.

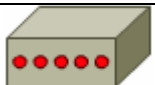
Consommation en bande passante réduite par rapport à un SVC.

Coûts supérieurs en raison de la continuité de service.

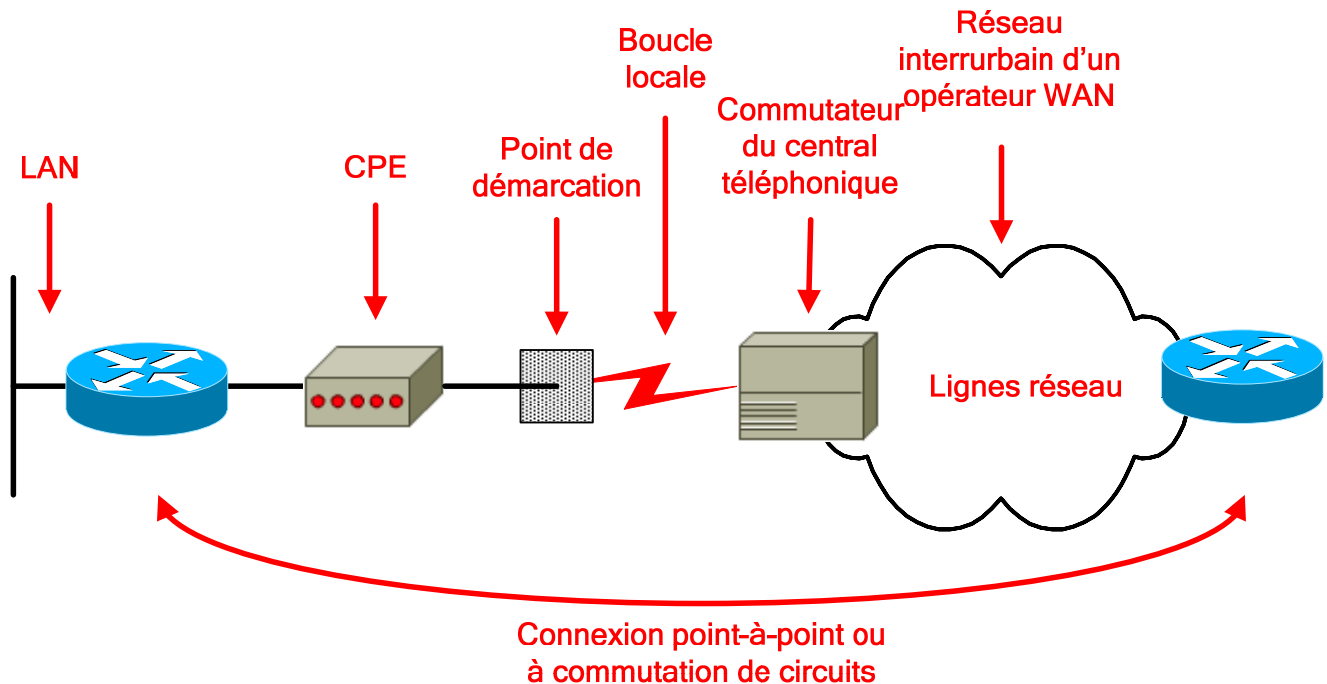
Exemples de lignes WAN et bande passante associée :

Type de ligne	Bande passante
T1	1.544 Mbits/s
E1	2.048 Mbits/s
E3	34.064 Mbits/s
T3	44.736 Mbits/s

3.2. Equipements et dispositifs

	Routeur
	Serveur de communication
	Commutateurs WAN (ATM, RNIS, etc.)
	Modem (Unité CSU/DSU, TA, NT1, etc.)

- **Routeur** : Dispositif de routage, offrant différents services dont des ports d'interface de réseau LAN et WAN.
- **Serveur de communication** : Concentrateur de communications utilisateur entrantes et sortantes.
- **Commutateur WAN** : Unité multiport qui assure les commutations du trafic WAN.
- **Modem** : Equipement de conversion d'un signal numérique en un signal analogique par l'intermédiaire du principe de modulation/démodulation.
- **Unité CSU/DSU** : Interface numérique (ou deux interfaces séparées, si les parties CSU et DSU sont séparées) qui adapte l'interface ETTD à celle d'un ETCD. Cette unité est généralement intégrée au routeur.



- **CPE** : Equipement placé dans les locaux du client, lui appartenant ou étant loué à l'opérateur (Exemple : modem).
- **Point de démarcation de service** : Démarcation entre la partie client et la partie opérateur (boucle locale). C'est à ce point que la responsabilité de chaque partie (Client et opérateur) s'arrête.
- **Boucle locale** : Partie reliant le point de démarcation de service au central téléphonique de l'opérateur.
- **Commutateur du central téléphonique** : Point de commutation le plus proche du client.
- **Réseau interurbain** : Unités et commutateur (appelés lignes réseau) situés dans le nuage de l'opérateur.

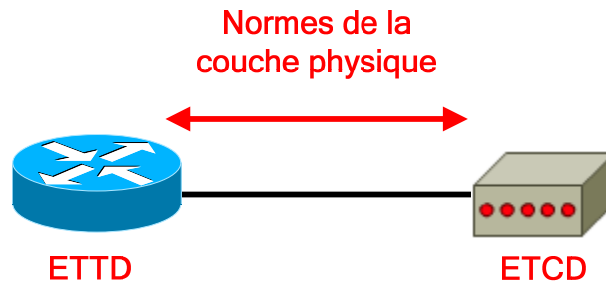
3.3. Normes WAN

Les normes des réseaux WAN décrivent généralement les méthodes d'acheminement de la couche physique ainsi que la configuration exigée pour la couche liaison de donnée, notamment :

- L'adressage.
- Le contrôle de flux.
- L'encapsulation.

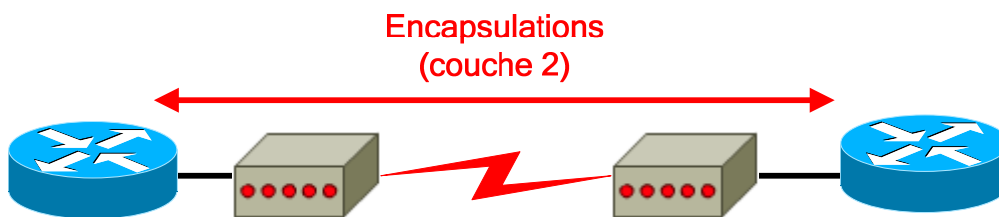
Les principaux organismes définissant et gérant les normes WAN sont :

- **UIT-T** (Union Internationale des Télécommunications - secteur de normalisation des Télécommunications), anciennement appelée CCITT (Comité Consultatif International Télégraphique et Téléphonique).
- **ISO** (International Standards Organization).
- **IETF** (Internet Engineering Task Force).
- **EIA** (Electrical Industries Association).
- **TIA** (Telecommunications Industry Association).



La couche physique d'un réseau WAN décrit principalement l'interface entre l'ETTD (unité connectée) et l'ETCD (fournisseur) :

- **EIA/TIA-232** : Similaire à la norme V.24 et anciennement appelée RS-232. Prévue pour les circuits asymétriques dont la bande passante peut atteindre 64 Kbits/s.
- **EIA/TIA-449** : Version plus rapide que l'EIA/TIA-232 (2 Mbits/s).
- **EIA/TIA-612/613** : Décrit l'interface HSSI (pour T3, E3, SDH STM-0, etc.).
- **V.24**.
- **V.35** : Décrit un protocole synchrone, utilisé pour la communication dans un réseau de paquets.
- **X.21** : Pour les lignes numériques synchrones.
- **G.703** : Connexions utilisant des connecteurs BNC et fonctionnant à des débits E1.
- **EIA-530** : Deux mises en œuvres électriques des normes EIA/TIA-449 :
 - o **RS-422** : Transmissions symétriques.
 - o **RS-423** : Transmissions asymétriques.



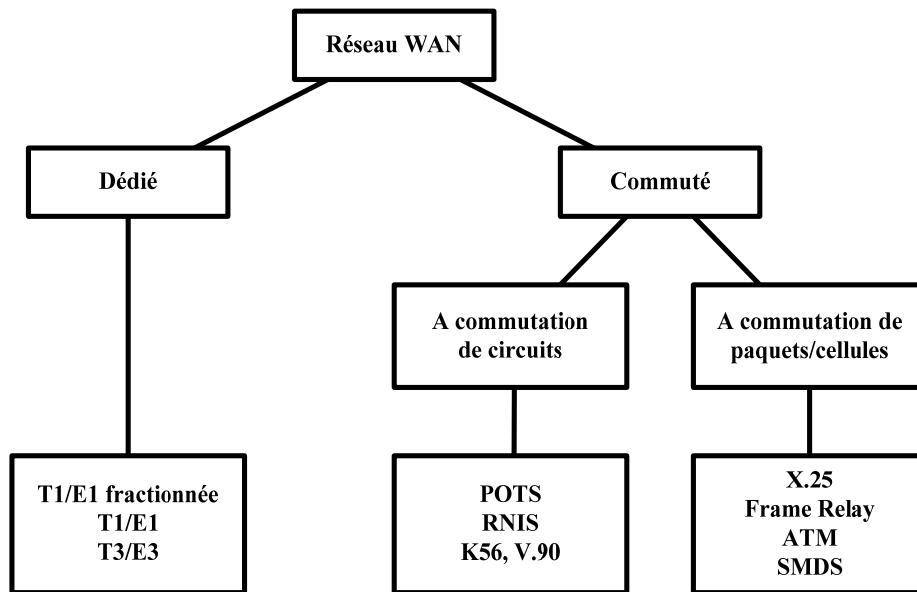
La couche liaison de données définit le mode d'encapsulation des données sur les réseaux WAN :

- **Frame Relay** :
 - o Encapsulation simplifiée.
 - o Dépourvue de mécanismes de correction des erreurs.
 - o Prévu pour des unités numériques haut de gamme.
 - o Transmet les données très rapidement par rapport aux autres encapsulations WAN.
 - o Il existe deux variantes pour cette encapsulation, à savoir Cisco et IETF.
- **PPP** :
 - o Comprend un champ identifiant le protocole de couche réseau.
 - o Vérifie la qualité de la liaison au moment de l'établissement d'une connexion.
 - o Gère l'authentification grâce aux protocoles PAP et CHAP.
- **RNIS** : Ensemble de services numériques pour la voix et les données sur le réseau commuté classique.
- **LAPB** :
 - o Encapsulation des paquets à la couche 2 de la pile X.25 sur des réseaux à commutation de paquets.
 - o Egalement sur des liaisons point-à-point, si elles ne sont pas fiables ou possèdent un délai inhérent (Exemple : liaison par satellite).
 - o Apporte la fiabilité et le contrôle de flux sur une base point-à-point.
- **HDLC** :
 - o Peut être incompatible entre fournisseurs car chacun a sa propre mise en œuvre.
 - o Prend en charge les configurations point-à-point et multipoints.
 - o Dérivé du protocole SDLC.
 - o Protocole par défaut pour les interfaces série d'un routeur Cisco.

- o Extrêmement simplifié : Pas de fonctions de fenêtrage ni de contrôle de flux.
- o Champ d'adresse contenant uniquement des 1, avec un code propriétaire à 2 octets indiquant le type de verrouillage de trame du fournisseur.

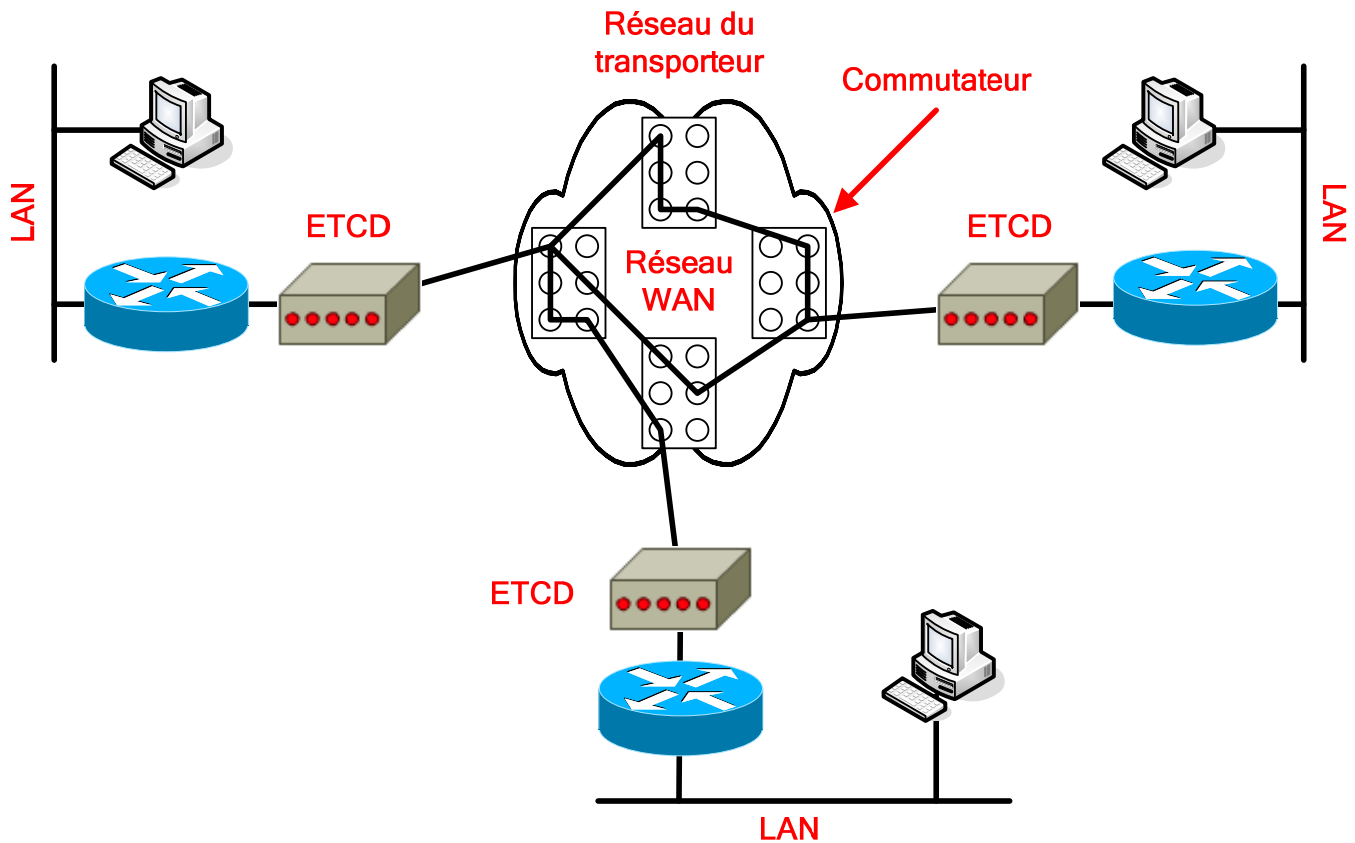
Le protocole HDLC est recommandé sur une liaison reliant deux équipements utilisant IOS. Dans le cas contraire, il est recommandé d'utiliser le protocole PPP.

3.4. Classement des différents types de liaison WAN



Les différents types de liaison WAN habituellement disponibles sont :

- **Liaisons dédiées** (aussi appelées liaisons spécialisées ou lignes louées) :
 - o Fournissent un service continu.
 - o Il s'agit d'un lien physique dédié qui va directement d'un port du routeur client à un port du routeur de l'opérateur, sans passer par un environnement commuté.
 - o Il est nécessaire d'avoir un port par liaison client sur le routeur de l'opérateur.
 - o Fournies par des liaisons série synchrone point-à-point.
 - o Cette liaison point-à-point est utilisée pour :
 - Une liaison physique directe.
 - Des liaisons virtuelles constituées de plusieurs liaisons physiques.
 - o Conviennent aux grands volumes d'information et aux trafics constants.
- **Connexions commutées** :
 - o **A commutation de circuits** :
 - Commutation physique des centraux téléphoniques afin d'obtenir la liaison point-à-point.
 - o **A commutation de paquets/cellules** :
 - Commutation « logique » effectuée au niveau de la couche 2 du modèle OSI.



Les deux grands types de liaison à commutation sont :

- **Commutation de circuits :**
 - o Circuit physique dédié par commutation des centraux téléphoniques.
 - o Établi, maintenu et fermé à chaque session.
 - o Établi à la demande.
 - o Sert aussi de ligne de secours aux circuits haut débit.
 - o Offre une bande passante dédiée.
- **Commutation de paquets/cellules :**
 - o Utilisation d'un PVC similaire à une liaison point-à-point.
 - o Possibilité d'acheminer des trames de taille variable (paquets) ou de taille fixe (cellules).
 - o Les unités du réseau partagent une liaison point-à-point unique.
 - o Plus souple et utilise mieux la bande passante que les services à commutation de circuits.

4. Conception WAN

4.1. Communication dans un WAN

La communication WAN est généralement appelée « service », car elle a un coût par rapport au temps d'utilisation (Facture forfaitaire ou basée sur la consommation) contrairement à la communication LAN (Uniquement les frais d'installation du matériel), et se caractérise habituellement par :

- Un débit relativement faible (par rapport aux réseaux LAN).
- Des délais importants (liés aux distances).
- Un taux d'erreurs généralement élevé (Réseaux WAN plus soumis aux interférences extérieures).

Le choix d'un service WAN dépend principalement des critères suivants :

- Optimisation de la bande passante.
- Réduction des coûts.
- Optimisation de l'efficacité du service.

Les besoins liés aux services WAN sont parmi les facteurs suivants :

- Augmentation de l'utilisation des réseaux (Applications client/serveur, multimédia, etc.).
- Évolution permanente des exigences relatives aux logiciels (Qualité, etc.).
- Nombre de connexions à distance en constante augmentation (Utilisateurs éloignés ou mobiles, sites répartis dans le monde, communication avec les clients et les fournisseurs, etc.).
- Croissance des intranets et extranets d'entreprise (bande passante).
- Utilisation de plus en plus importante des serveurs d'entreprise.

4.2. Premières étapes de la conception WAN

Les deux principaux objectifs de la conception et de la mise en œuvre d'un WAN sont :

- Disponibilité des applications (Accès aux applications = efficacité du réseau).
- Coût (Utilisation rentable des ressources).

Ces deux critères sont fondamentalement contradictoires. Il est donc nécessaire d'observer une pondération entre la relative importance de la disponibilité des ressources et les prix de revient globaux.

La première étape de la conception d'un réseau WAN est de recueillir des informations :

- Données sur la structure et les processus de l'entreprise.
- Déterminer les personnes susceptibles de nous aider à concevoir le réseau.
- Identifier les besoins des utilisateurs (concernant la disponibilité des applications) :
 - o Temps de réponse.
 - o Débit.
 - o Fiabilité.

Les différentes méthodes d'évaluation des besoins des utilisateurs sont :

- **Les profils des utilisateurs** : Définition des besoins des divers groupes d'utilisateurs.
- **Des entretiens, groupes de discussion et sondages** : Etablissement d'une base de référence.
- **Des entretiens aux groupes d'utilisateurs clés** : Méthode de collecte de renseignements par échantillonnage.
- **Tests du facteur humain** : Test en laboratoire avec un groupe représentatif d'utilisateurs. C'est la méthode d'évaluation la plus coûteuse et significative.

Cette analyse des besoins des utilisateurs a pour but de déterminer :

- Le type de trafic passé.
- Le niveau du trafic.
- Le temps de réponse des systèmes hôtes.
- La durée d'exécution des transferts de fichiers.
- L'utilisation de l'équipement réseau existant.

Les besoins ne sont pas statiques, il faut donc prendre en compte :

- L'accès au réseau changeant en fonction du temps (Période de pointe).
- Les différences liées au type de trafic (Sensibilité aux paquets abandonnés, exigence en bande passante).
- La nature aléatoire du trafic réseau (les heures d'utilisation peuvent changer).

Ensuite, il reste à effectuer un test de sensibilité en brisant des liaisons stables et à observer le résultat. On peut utiliser une de ces deux méthodes :

- **Supprimer une interface active** : Observation de la redirection du trafic, d'une probable perte de connectivité.
- **Modifier la charge réseau** : Observation du comportement du réseau lors de la saturation du réseau.

4.3. Modèle de réseau hiérarchique

Il existe deux structures de modèle de réseau :

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

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

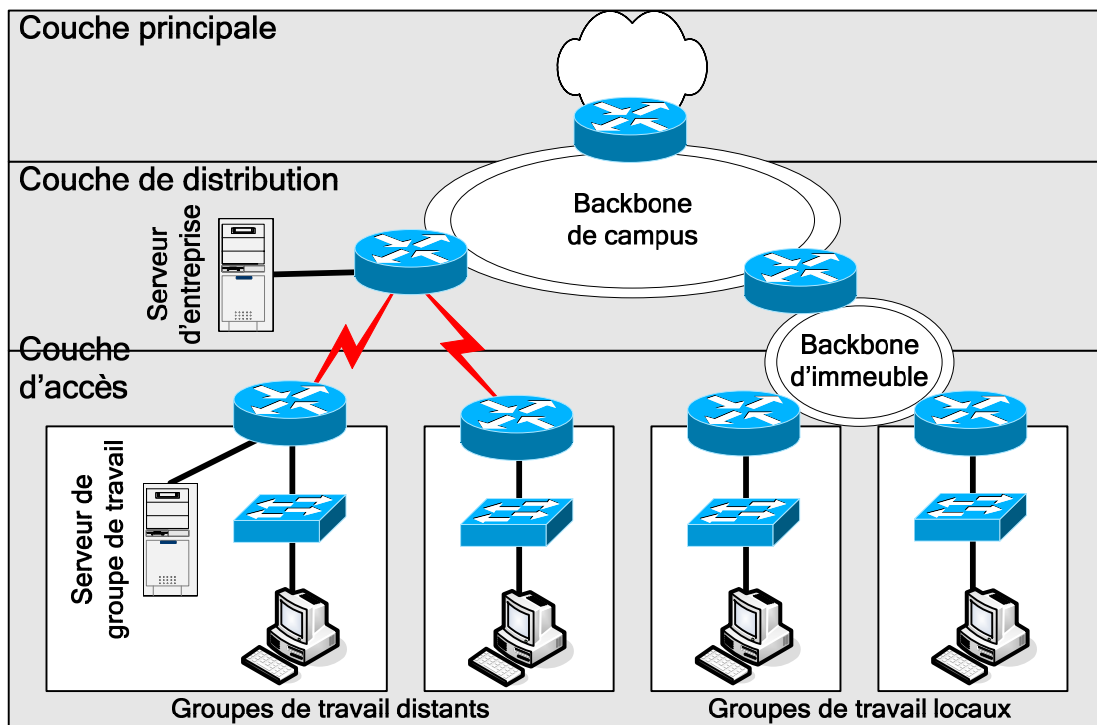
- Faciliter les modifications et la compréhension du réseau (Réseau modulaire).
- Limiter les coûts et la complexité des mises à niveau du réseau (appliquées à un sous-ensemble uniquement).
- Limiter les coûts de construction et d'élaboration du réseau.
- Faciliter l'identification des points de défaillance.

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

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

Les couches, dans un modèle de conception, sont séparées par des dispositifs de couche 3 du modèle OSI, qui séparent le réseau en domaines de broadcast.

4.3.1. Modèle à 3 couche



Les couches de ce modèle sont :

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

La couche principale :

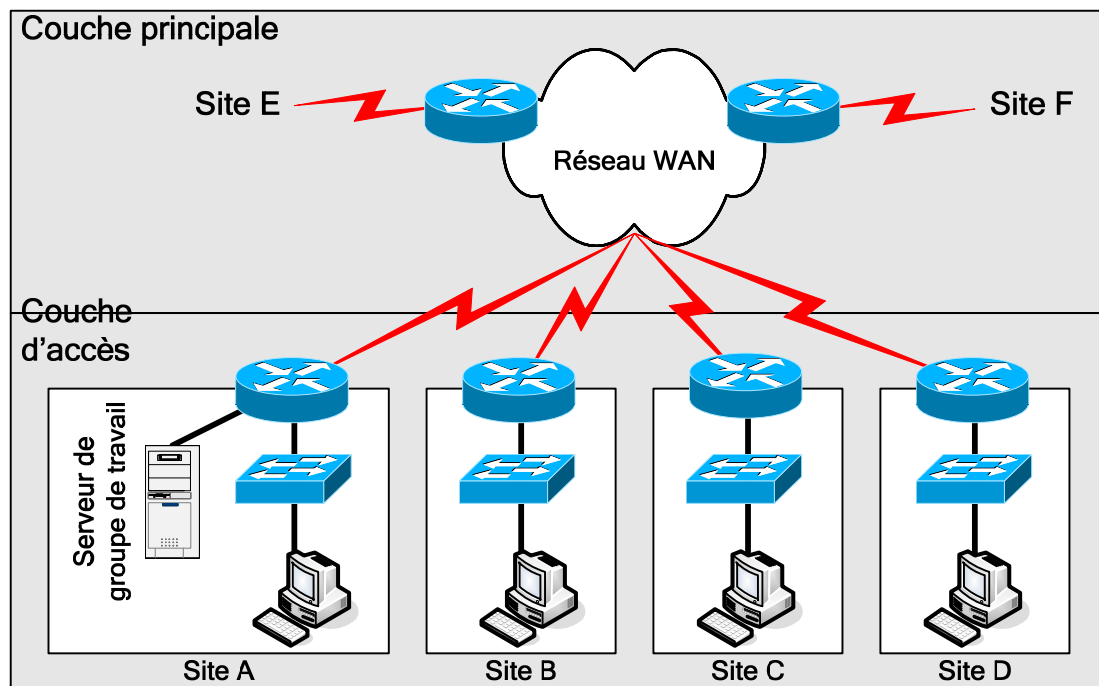
- Assure la communication (la plus rapide possible) entre les sites éloignés.
- Comporte habituellement des liaisons point-à-point.
- Aucun hôte présent, que des unités de communication.
- Services présents (Frame Relay, T1/E1, SMDS) loués auprès d'un fournisseur de services.
- Ne s'occupe pas du filtrage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

La couche distribution :

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

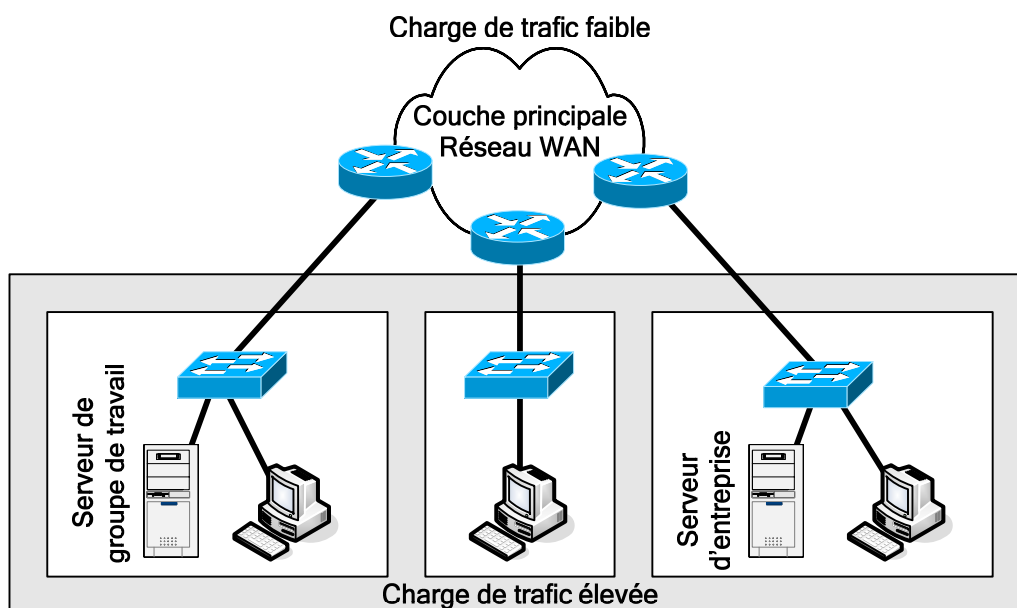
La couche d'accès :

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

4.3.2. Modèle à 2 couche

Dans un modèle à 2 couches, les sites distincts sont interconnectés directement par l'intermédiaire de liaisons WAN, représentant la couche principale. Chaque site peut contenir plusieurs LAN.

4.3.3. Modèle à 1 couche



Un réseau à une couche (Modèle linéaire) est mis en œuvre si l'entreprise n'a pas beaucoup d'emplacements éloignés, et si l'accès aux applications se fait principalement à l'intérieur du LAN.

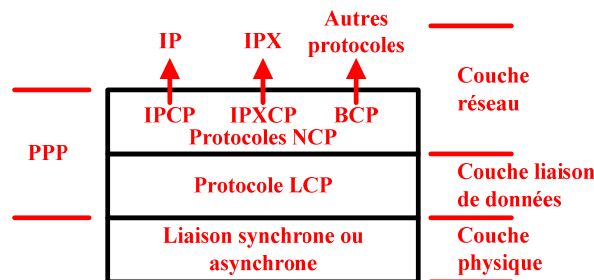
5. Protocole PPP

5.1. Etude 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).
 - o 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.

5.2. Etablissement 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**.

5.3. 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 :

- o Confirmation de l'authentification.
 - o Interruption de la connexion.
- **PAP n'est pas très efficace :**
 - o Mots de passe envoyés en clair.
 - o 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) :
 - o Confirmation.
 - o Réponse.
 - o Acceptation ou refus.
- **Méthode d'authentification plus évoluée :**
 - o Vérification régulière de l'identité du nœud distant (A l'établissement puis à tout moment).
 - o Authentification dans les deux sens.
 - o Impossibilité de tenter une authentification sans avoir reçu une demande de confirmation.
 - o Authentification cryptée via l'algorithme MD5 lors du transit sur la liaison.
- **Efficacité contre le piratage :**
 - o Utilisation d'une valeur de confirmation variable, unique et imprévisible.
 - o Répétition des demandes de confirmation visant à limiter la durée d'exposition aux attaques.
 - o 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 :

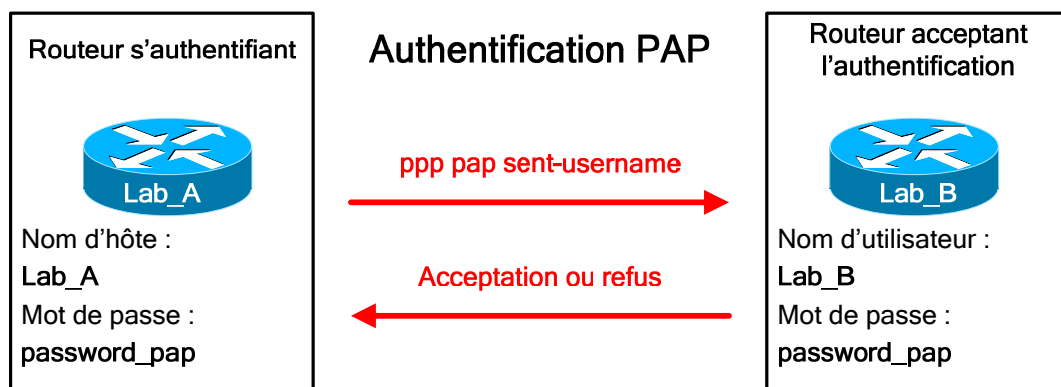
- **username {nom} password {mot_de_passe} :**
 - o Mode de configuration globale.
 - o Paramètre nom : Nom d'hôte qu'on souhaite accepter.
 - o Paramètre mot_de_passe : Mot de passe à utiliser pour l'authentification. Celui-ci doit correspondre au mot de passe du mode privilégié crypté du routeur distant 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

5.3.1. Procédure de configuration du protocole PAP



Nous allons d'abord étudier la configuration qu'il faut utiliser pour une authentification unidirectionnelle.

```
Lab_A (config-if)# encapsulation ppp
Lab_A (config-if)# ppp authentication pap callin
Lab_A (config-if)# ppp pap sent-username Lab_A password password_pap
```

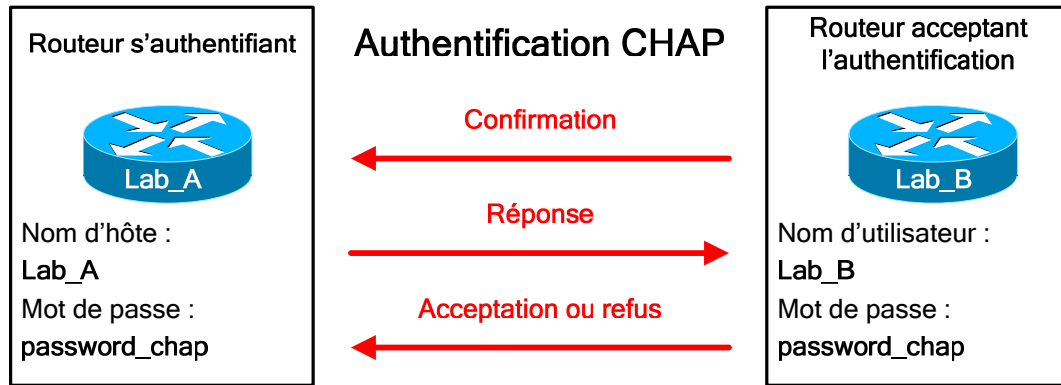
```
Lab_B (config)# username Lab_A password password_pap
Lab_B (config-if)# encapsulation ppp
Lab_B (config-if)# ppp authentication pap
```

Pour une authentification bidirectionnelle, il suffit de procéder comme suit :

```
Lab_A (config)# username Lab_B password password_pap
Lab_A (config-if)# encapsulation ppp
Lab_A (config-if)# ppp authentication pap
Lab_A (config-if)# ppp pap sent-username Lab_A password password_pap
```

```
Lab_B (config)# username Lab_A password password_pap
Lab_B (config-if)# encapsulation ppp
Lab_B (config-if)# ppp authentication pap
Lab_B (config-if)# ppp pap sent-username Lab_B password password_pap
```

5.3.2. Procédure de configuration du protocole CHAP



Le schéma d'authentification ci-dessus représente l'authentification dans un seul sens, il va donc falloir répéter ce schéma dans les deux sens de l'authentification CHAP.

Pour cela, nous allons effectuer les tâches de configuration suivantes sur le routeur Lab_A :

```
Lab_A (config)# username Lab_B password password_chap
Lab_A (config-if)# encapsulation ppp
Lab_A (config-if)# ppp authentication chap
```

Les commandes à utiliser sur le routeur Lab_B sont :

```
Lab_B (config)# username Lab_A password password_chap
Lab_B (config-if)# encapsulation ppp
Lab_B (config-if)# ppp authentication chap
```

6. Technologies RNIS

6.1. 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 multi liaison 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.

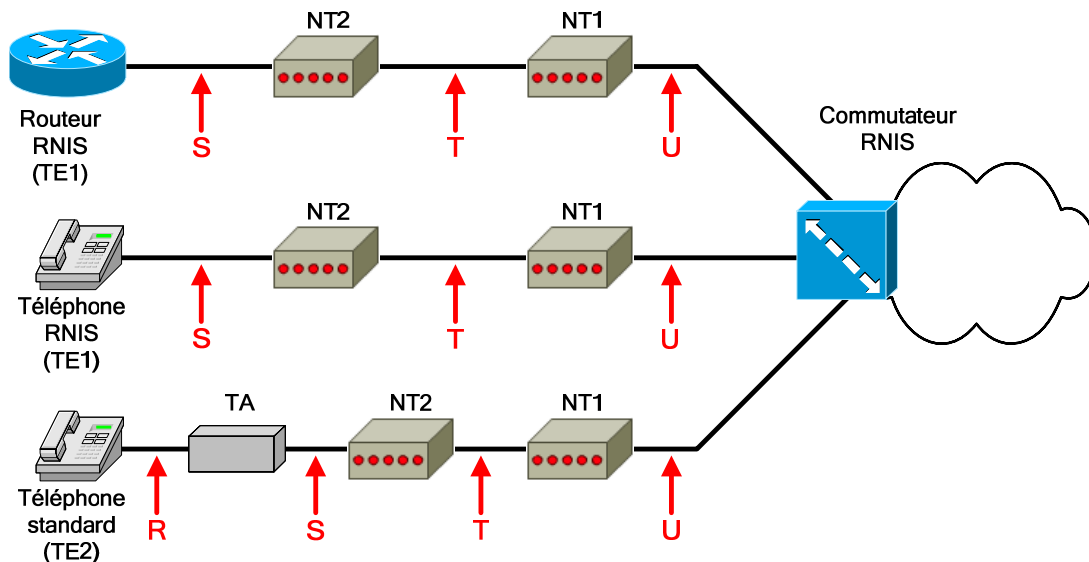
Le protocole LAPD (Couche 2) est utilisé sur le canal D et permet une circulation et une réception adéquate des flux d'information de contrôle et de signalisation. Ce protocole est similaire à HDLC et à LAPB (X.25).

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.

6.2. Termes et 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).

6.3. 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.

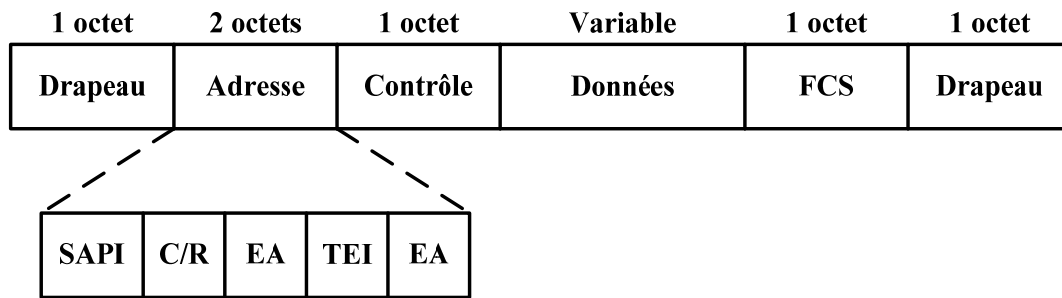
Les différentes normes que nous étudierons en fonction des couches du modèle OSI sont :

- **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) :

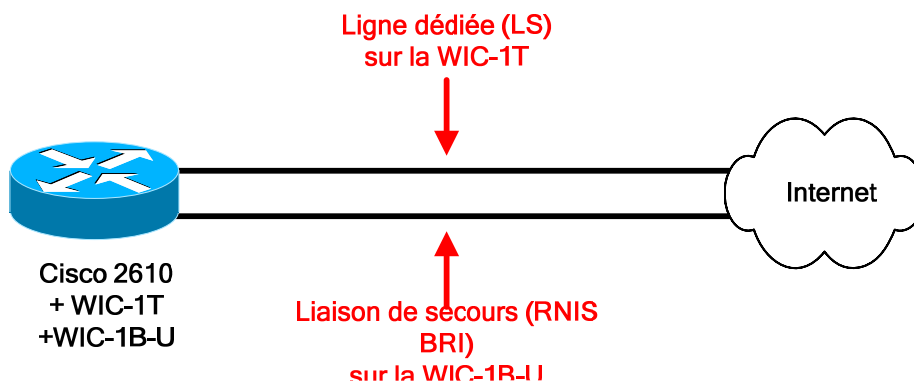


- **Drapeau** : Similaire au champ HDLC.
- **Adresse** : Peut comporter 1 ou 2 octets (Dépend de la valeur des bits EA).
 - o **SAPI** : Bits d'identification du point d'accès (6 bits). Indique le portail où les services LAPD sont fournis à la couche 3.
 - o **C/R** : Bit de commande/réponse.
 - o **EA** : 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).

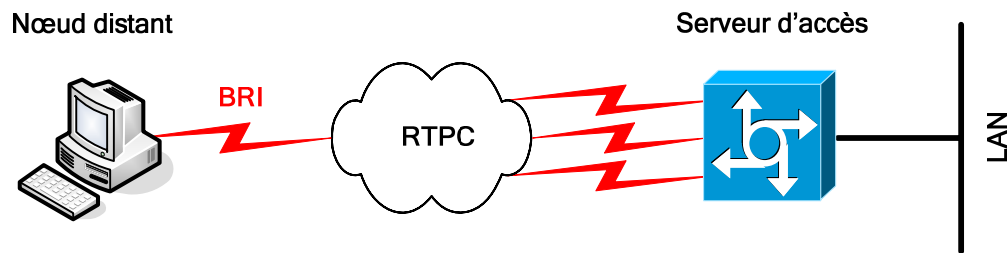
6.4. Utilisation/Implémentation

La technologie RNIS a de nombreuses applications :

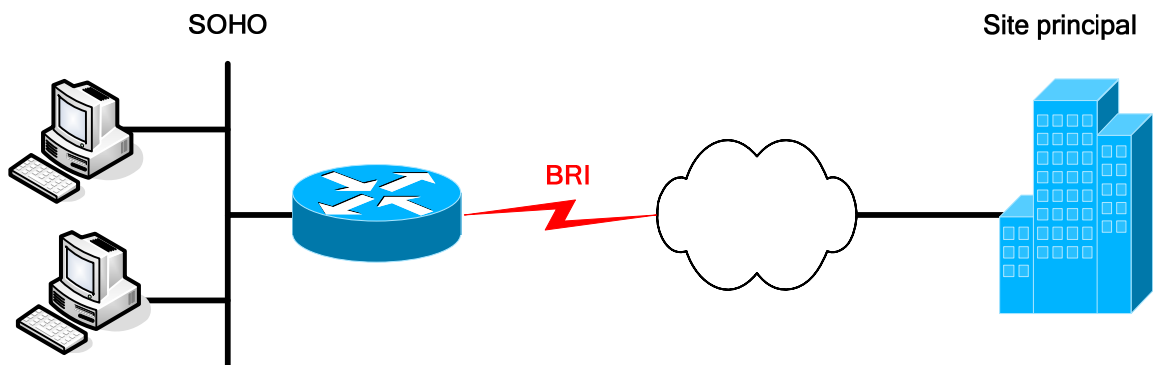
- Solution alternative aux lignes dédiées.
- Accès à distance :
 - o Nœuds distants.
 - o Connectivité des petits bureaux et bureaux à domicile (SOHO – Small Office / Home Office).



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).

6.5. 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..

6.6. 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 rendre 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.

6.7. 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.

7. Technologies Frame Relay

7.1. 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 minime (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.

7.2. 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.

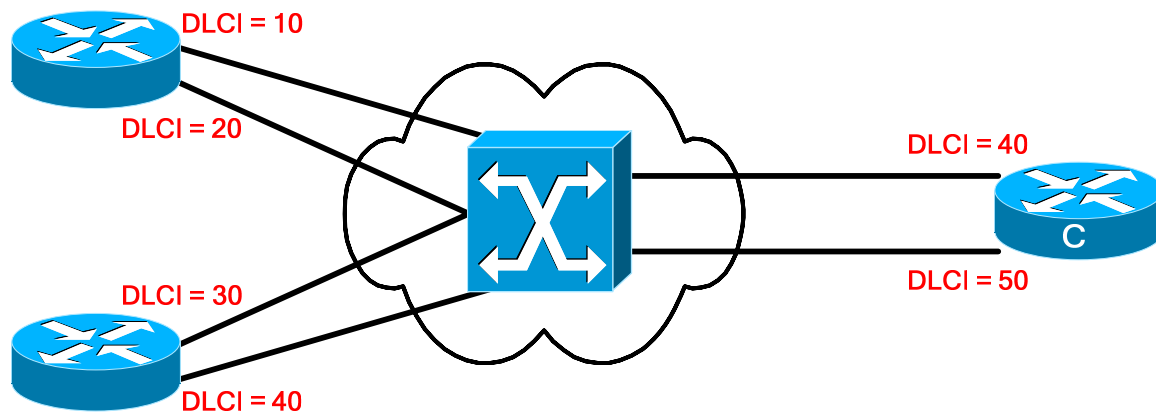
Il existe des extensions LMI, qui sont optionnelles :

- **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**.

7.3. Fonctionnement, table de commutation et 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 initiale 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.

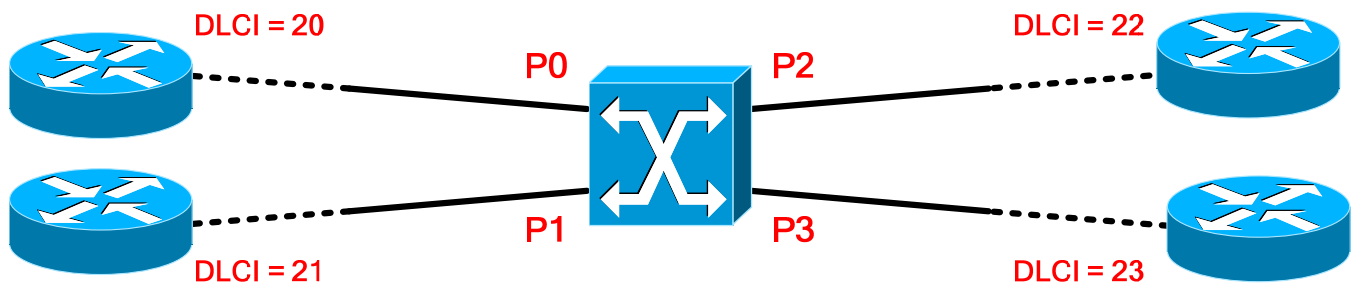


Table de commutation du port P0

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 :

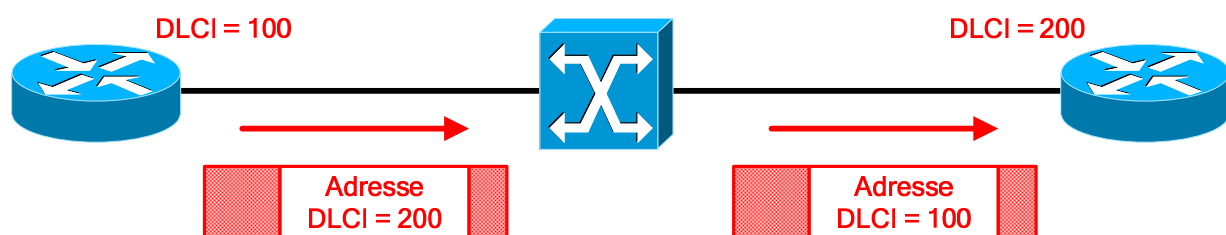
- Port d'entrée.
- DLCI d'entrée.
- Port de sortie.
- DLCI de sortie.

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 :

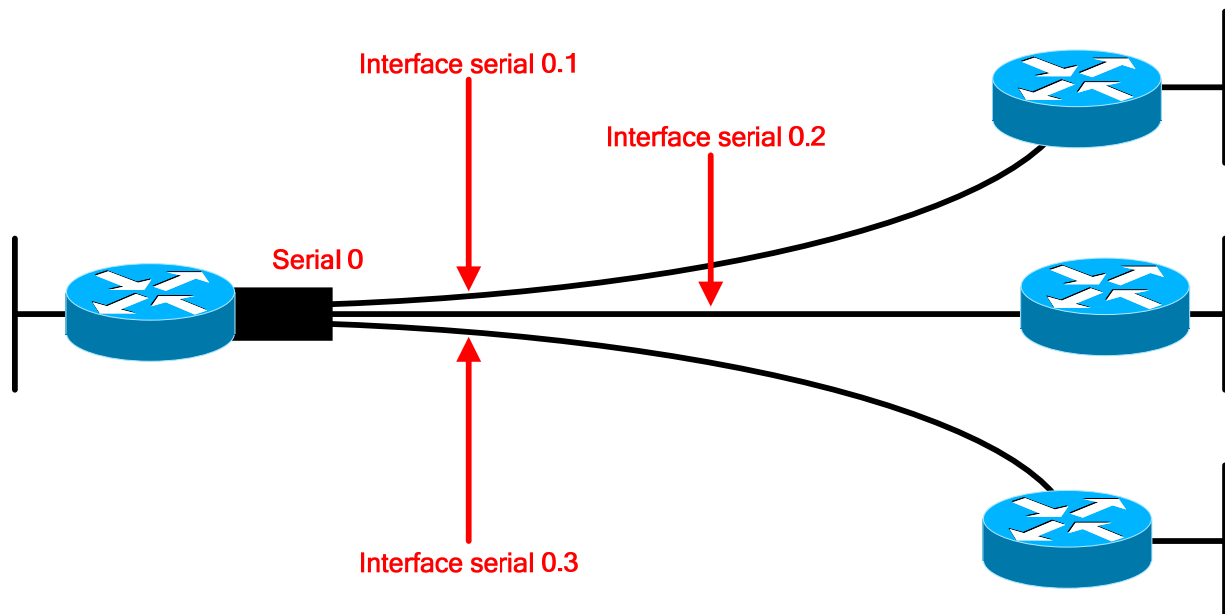
- Émission d'un message de demande d'état au commutateur Frame Relay (donne l'état du routeur local et demande celui des connexions des routeurs distants).
- Le commutateur répond avec un message d'état, contenant les DLCI des routeurs distants qui sont accessibles au routeur local.
- Pour chaque DLCI 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 :

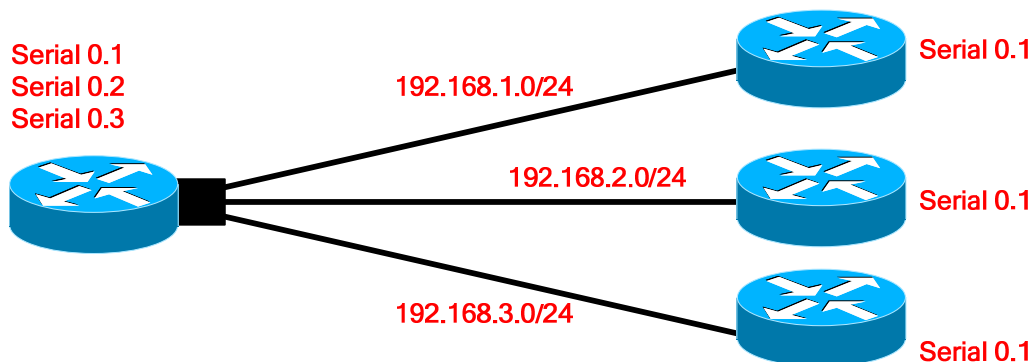
- Le routeur source encapsule les données à transmettre dans une trame Frame Relay, dont la valeur du champ Adresse correspond au DLCI du destinataire, puis l'envoie.
- Le commutateur reçoit cette trame, et utilise la table de commutation du port d'entrée afin de déterminer le port de sortie, et donc 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.

7.4. Les sous interfaces Frame Relay



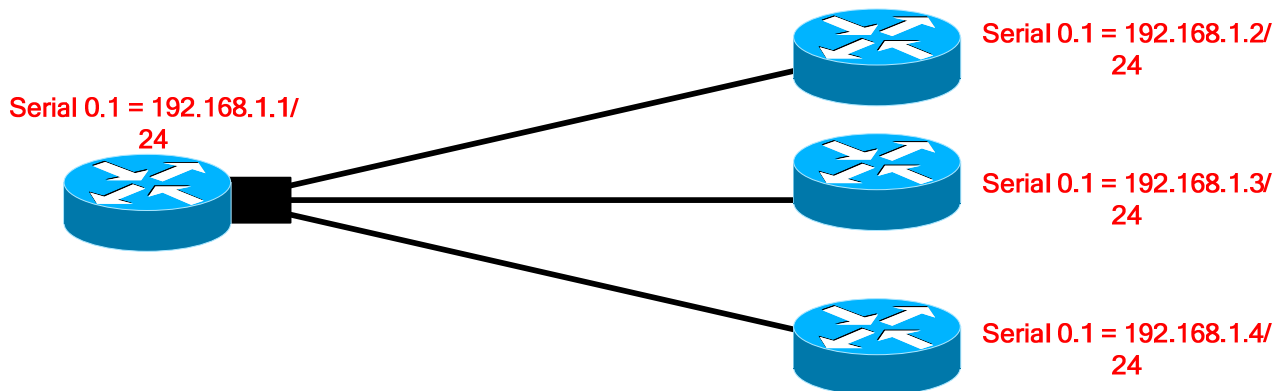
Les sous-interfaces sont des subdivisions logiques d'une interface physique et peuvent être de deux types :

- 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.

7.5. Commandes

Les commandes concernant Frame Relay sont les suivantes :

- **interface serial {numéro} :**
 - o Mode de configuration globale.
 - o Permet de passer dans le mode de configuration de l'interface souhaitée.
- **interface serial {numéro.sous-numéro} {multipoint | point-to-point} :**
 - o Mode de configuration globale.
 - o Permet de passer dans le mode de configuration de la sous-interface souhaitée.
 - o Le paramètre multipoint ou point-to-point définit le type de sous-interface utilisée.
 - o 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] :**
 - o Mode de configuration d'interface.
 - o Précise l'encapsulation des trames pour l'interface courante.
 - o Le paramètre cisco est la valeur par défaut, et est à utiliser si on est raccordée à un autre équipement Cisco.
 - o Le paramètre ietf est utile pour se connecter à un dispositif non Cisco.
- **frame-relay interface-dlci {dlci} :**
 - o Mode de configuration de sous-interface.
 - o 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 ETTD, 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} :**
 - Mode de configuration d'interface.
 - Permet d'explicitier le type d'interface Frame Relay locale.
 - La valeur par défaut est **dte**.
 - **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 :**
 - Mode de configuration globale. Permet d'activer la commutation de PVC sur une unité ETCD (Commutateur Frame Relay).
 - Active l'interface LMI.
- **frame-relay route {dlci_src} interface {type} {numéro} {dlci_dest} :**
 - Mode de configuration d'interface.
 - Permet de créer une entrée dans la table de commutation Frame Relay.
 - Il faut indiquer le DLCI source, l'interface locale de sortie et celui de la destination.
 - 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.

7.6. 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.

8. Initiation à l'administration réseau

8.1. Stations de travail et serveurs

Les premiers ordinateurs personnels (PC) furent conçus pour fonctionner de manière autonome. Le système d'exploitation utilisé sur ces machines autorisait l'accès au fichier et aux ressources du système à un utilisateur à la fois. Peu à peu, les PC ont envahis les espaces de travail, nécessitant de la part des systèmes d'exploitation des fonctions de réseau, permettant le partage de ressource.

Ces systèmes d'exploitation réseaux classifient les ordinateurs en 2 grandes familles :

- Les stations de travail
- Les serveurs

8.1.1. Stations de travail

Une station de travail est un poste utilisateur qui exécute une application et qui est connecté à un serveur à partir duquel il obtient des données partagées. La plupart d'entre elles dispose de connexions réseaux et supporte les accès multi-utilisateurs. Une station de travail peut être de type :

- Ordinateur de bureau.
- Ordinateur portable.
- Ordinateur sans disque dur.

8.1.2. Serveurs

Un serveur est un ordinateur exécutant un système d'exploitation réseau auquel des stations de travail viendront se connecter. De manière générale, les serveurs sont des machines plus puissantes et plus robustes que les stations de travail.

8.2. Systèmes d'exploitation réseau

Un système d'exploitation est un environnement au travers duquel les applications et les services sont exécutés sur une machine. Un système d'exploitation réseau aussi appelé NOS¹, permet la communication entre plusieurs équipements et ressources à travers le réseau. C'est un système multi-tâches et Multi-Utilisateurs capable d'exécuter plusieurs programmes à la fois. Les caractéristiques d'un tel système sont :

- Performance
- Gestion et supervision
- Sécurité
- Evolutivité
- Robustesse/tolérance de panne

Il existe plusieurs familles de système d'exploitation réseau (Windows, Unix, Linux, Apple). Les plus connus sont détaillés ci-dessous.

¹ Network Operating System

8.2.1. Systèmes d'exploitation réseau Microsoft Windows

Microsoft dispose, dans son offre commerciale, de plusieurs NOS :

- Windows NT4 Server

Sortie en Juillet 1996, Windows NT4 peut s'exécuter à la fois en tant que station de travail (NT4 Workstation) ou en tant que serveur (NT4 Server). Windows NT utilise une structure de domaine afin de contrôler les accès utilisateurs et les accès aux ressources. Chaque domaine NT nécessite la présence d'un contrôleur de domaine contenant la base SAM². Lorsqu'un utilisateur se connecte au domaine NT, les informations du compte de l'utilisateur sont envoyées à la base de données SAM. Si le compte est valide, l'utilisateur est authentifié sur le domaine et a accès à la station de travail.

- Windows2000 Server

Sortie en février 2000, Windows 2000 existe en version « Professional » et « Server »

Basé sur le noyau de Windows NT4, Windows2000 Server intègre également la technologie « Plug and Play ». La gestion des utilisateurs et des ressources d'un domaine peut maintenant se faire en tant qu'objets. Ceux-ci peuvent être placés dans des conteneurs, dont la gestion peut être déléguée à un utilisateur ou un groupe. Tout cela est possible via la technologie Active Directory.

- Windows2003 Server

Sortie en Avril 2003, Windows 2003 Server reprend les points forts de Windows2000 Server. Il accroît les fonctionnalités de migration depuis Windows NT4, tout en étant compatible avec un domaine NT4. Divers services réseaux ont été améliorés tel que « IIS Web Server ». De plus, la technologie .NET a été directement intégrée au système.

8.2.2. Systèmes d'exploitation réseau UNIX et Linux

- UNIX

Unix est un nom donné à un groupe de systèmes d'exploitations issu des laboratoires Bell de 1969. C'est un système multi-utilisateurs et multitâches qui prend en compte les protocoles réseau d'Internet. Au fil du temps plusieurs entreprises ont contribué au développement d'Unix, ce qui entraîna dans les années 1980 sa commercialisation sous diverses appellations :

- Hewlett Packard UNIX (HP-UX)
- Santa Cruz Operation (SCO) UNIX
- Sun Solaris
- IBM UNIX (AIX)

Berkeley Software Design, Inc. (BSD UNIX) distribuera également sa version d'Unix qui produira des dérivés tels que :

- FreeBSD
- OpenBSD
- NetBSD

Unix sous ses diverses formes compose et consolide aujourd'hui sa position de système d'exploitation fiable et sécurisé. Cependant Unix est souvent associé à du matériel coûteux et propriétaire, mais la création de Linux est en train de changer cette image.

- Linux

En 1991, frustré par l'état des systèmes d'exploitation de bureau, mais aussi par les coûts et les problèmes de licence, un étudiant Finlandais du nom de Linus Torvald se mit à travailler sur un système d'exploitation destiné

² Security Accounts Management Database

Laboratoire SUPINFO des Technologies Cisco

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

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

aux ordinateurs à base de processeur 80386. Son système était semblable à Unix et, particularité de ce dernier, le code était ouvert et gratuit pour tous les utilisateurs. Son travail mena à une collaboration Internationale entre la communauté des développeurs et dès la fin des années 1990, Linux était devenue une alternative aux serveurs Unix et aux systèmes de bureau Windows.

A l'instar d'Unix il existe plusieurs versions de Linux dont :

- Red Hat Linux – distribué par Red Hat Software
- OpenLinux – distribué par Caldera
- Corel Linux
- Slackware
- Debian GNU/Linux
- SuSE Linux

Linux est doté de composants réseaux intégrés permettant de se connecter à un réseau local, établir une connexion réseau commutée vers l'Internet ou faire du tunnelling. La pile de protocole TCP/IP est d'ailleurs directement intégrée au noyau Linux.

8.2.3. Système d'exploitation réseau Apple

Apple dispose également d'une version Serveur de son fameux système d'exploitation Mac OS X. Ce dernier dénommé Mac OS X Server est capable de gérer des ordinateurs sous divers systèmes d'exploitation Apple et concurrents (Mac OS 9, Microsoft Windows, Unix et Linux etc..). Le noyau de Mac OS X qui a pour nom de code « Darwin » est dérivé de la technologie BSD4.5 et 5.0. Il en résulte une combinaison de la technologie serveur open source la plus populaire, combinée à l'installation et l'utilisation aisée des systèmes Apple. Les applications réseaux classiques sont toutes supportées (NTP, SMTP, DNS, LDAP, etc..) et le partage des données avec des clients Unix et Windows est également supporté (Nfs, Samba).

8.3. Gestion du réseau

8.3.1. Introduction à la gestion réseau

Un réseau évolue. A mesure que ce dernier s'étend, il devient une ressource de plus en plus cruciale pour l'organisation. Sa gestion se complique et conséquence de tout cela, le réseau devient de plus en plus complexe. Dans ce cas de figure, la tâche de l'administrateur devient ardue : la non constatation de la défaillance d'un service peut avoir des conséquences graves en environnement de production.

L'administrateur doit gérer le réseau de manière active, diagnostiquer les problèmes, prévoir les pannes, et les empêcher de survenir. Les mauvaises performances et la perte de ressources réseaux ne sont pas acceptables pour les utilisateurs. Il devient très difficile pour un administrateur, voire impossible, d'assurer toutes ces tâches sans aide logicielle ni outils de gestion automatique du réseau.

Les facteurs qui régissent l'administration réseau sont les suivants :

- **Contrôle des ressources de l'entreprise** – Gestion efficace des ressources réseaux. Le cas échéant, les résultats fournis ne seront pas à la hauteur d'une administration efficace.
- **Contrôle de la complexité** – Contrôler l'évolution du réseau afin d'éviter que trop de complexité n'entraîne la perte de contrôle de ce dernier.
- **Amélioration du service** – S'assurer que l'utilisateur bénéficie d'un meilleur service, sinon égal à l'ancien, à mesure que le réseau évolue.
- **Équilibrage des divers besoins** – Les applications mises à la disposition des utilisateurs doivent l'être avec un niveau donné de support, de disponibilité et de sécurité.
- **Réduction des temps d'arrêt** – Assurer la redondance des services en environnement haute disponibilité.

- **Contrôle des coûts** – Surveiller et contrôler l'utilisation des ressources, de cette manière les utilisateurs peuvent être satisfaits à coût raisonnable.

L'administration réseau implique les tâches ci-dessous:

- La surveillance de la disponibilité du réseau
- L'amélioration de l'automatisation
- La surveillance des temps de réponse
- La mise en place de fonctionnalités de sécurité
- Le réacheminement du trafic
- Le rétablissement de la fonctionnalité
- L'enregistrement d'utilisateurs

8.3.2. Modèle de gestion réseau et OSI

Afin d'avoir un modèle commun à tout les constructeurs, l'ISO s'est occupé de créer un standard pour la gestion du réseau. La tâche de produire un modèle d'administration réseau commun fut assignée à un comité dirigé par le groupe OSI.

Le comité en charge de cette modélisation en est arrivé à un modèle d'administration découpé en quatre parties :

- **Le modèle d'organisation** Il définit les différents composant de l'administration réseau, Administrateur, NMS, agent SNMP etc., ainsi que leurs relations.
- **Le modèle d'information** Il définit la structure de stockage des informations d'administration appelé SMI. Cette structure définit la syntaxe des informations d'administration. Le contenu de la SMI est appelé MIB.
- **Le modèle de communication** Il définit la manière dont les données sont acheminées depuis la NMS jusqu'aux agents SNMP. Il traite du protocole de communication (SNMP).
- **Le modèle fonctionnel** Il traite des applications d'administration réseau qui s'exécutent sur la NMS.

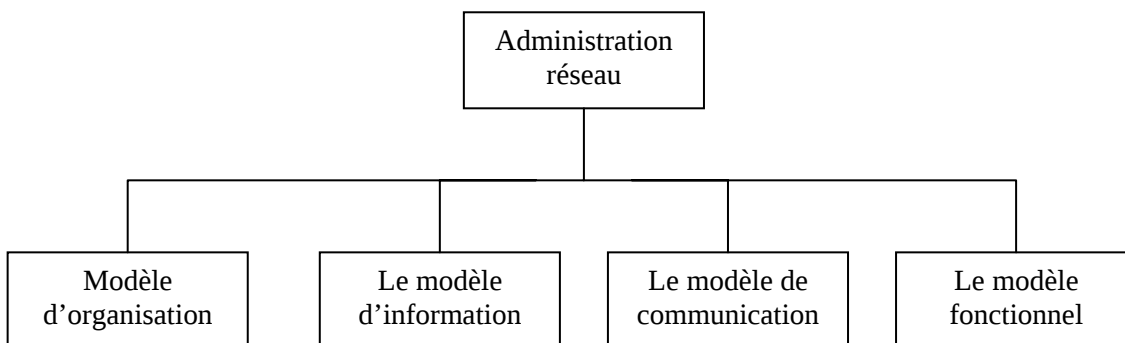


Figure 1- Modèle de gestion réseau de l'ISO

8.4. Protocole SNMP

8.4.1. Introduction

SNMP (Simple Network Management Protocol) a été adopté comme norme pour les réseaux TCP/IP en 1989. Ce protocole désigne un ensemble de normes d'administration, notamment :

- Un protocole de communication
- Une spécification de structure de base de données
- Un ensemble d'objets de données

Très populaire et présent dans la plupart des réseaux d'entreprise, SNMP connu une mise à niveau (SNMPv2c) en 1993, améliorant entre autre la structure des informations d'administration, l'authentification ainsi que le protocole lui-même. SNMP évolue pour en arriver à la version 3 (SNMPv3) qui prend en charge l'authentification et le cryptage des communications tout en restant rétro compatible.

8.4.2. Fonctionnement

SNMP est un protocole de la couche application conçu pour faciliter l'échange d'informations d'administration entre les équipements réseaux. On peut par exemple l'utiliser pour accéder à des données d'informations d'administrations tels que le nombre de paquets en sortie sur l'interface WAN d'un routeur, le nombre de connexions TCP ouvertes ou même la quantité d'erreur détectées sur cette même interface.

La quantité d'informations accessibles et récupérables est très nombreuse et détaillée. SNMP est un protocole simple, mais ses fonctions sont suffisamment efficaces pour gérer les problèmes liés à l'administration des réseaux hétérogènes. Le modèle organisationnel de l'administration réseau SNMP comporte quatre éléments :

- La station de gestion du réseau (NMS : Network Management System)
- Les agents de supervision (Agent SNMP)
- La base d'information de management (MIB : Management Information Base)
- Le protocole de gestion réseau.

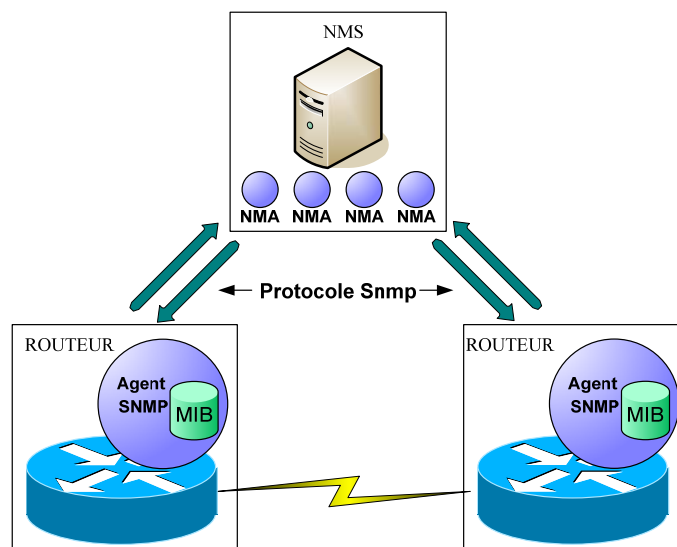


Figure 2-Fonctionnement de SNMP

La NMS est généralement une station de travail autonome. Elle se compose d'un ensemble de logiciels appelé NMA.

Ceux-ci intègrent une interface utilisateur permettant aux administrateurs de superviser le réseau en récupérant des informations sur les agents SNMP. Ceux-ci sont situés sur les différents équipements réseaux (routeur, pont, commutateur, répéteur, serveur d'application).

Un agent SNMP peut répondre à une requête d'exécution d'action de la part de la NMS. Il peut également remonter des informations utiles, non sollicitées par la NMS, telles que la perte de connectivité entre deux routeurs, ou un dysfonctionnement du service de messagerie de l'entreprise.

Un agent SNMP peut effectuer un suivi de ces éléments :

- Le nombre et l'état de ses circuits virtuels.
- Le nombre de certains types de messages d'erreur reçus.
- Le nombre d'octets et de paquets entrant et sortant de l'équipement.
- La longueur maximale de la file d'attente de sortie pour les routeurs et autres équipements inter réseaux.
- Les messages de broadcast envoyés et reçus.
- L'état d'activation des interfaces réseau.

Afin de permettre à une NMS de dialoguer avec un agent SNMP, le protocole définit une chaîne de caractère : « l'identifiant de communauté ». Les échanges ne sont possibles qu'entre agents et NMA d'une même communauté SNMP.

Cette forme très basique de vérification reste une simple identification implémentée dans le protocole SNMP (SNMPv1).

Ceci représentant une faille de sécurité de taille (cet identifiant transitant en clair), la version 2 de SNMP a bénéficié de l'implémentation de mécanismes d'authentification et d'intégrité (chiffrement symétrique à clé privée utilisant l'algorithme HMAC-MD5-96).

Celle-ci posant des problèmes de rétro compatibilité, la version 3 a été conçue pour parer à ces problèmes. SNMPv3 permet donc une sécurité accrue ainsi qu'une rétro compatibilité.

A un identifiant de communauté, peut être affecté des permissions en lecture seulement ou en lecture/écriture sur les objets.

La communauté par défaut pour la lecture seule est « public », et « private » pour l'accès en lecture et écriture.

Version	Authentification	Confidentialité	Cryptage	Fonctionnement
SNMPv1	Non	Non	Non	Identification assurée par l'appartenance à la communauté SNMP
SNMPv2c	Oui	Oui	Oui	Authentification par chiffrement symétrique Problème de rétro compatibilité
SNMPv3	Oui	Oui	Oui	Authentification par chiffrement symétrique Rétro compatible

Tableau 1-Différences SNMPv1SNMPv2c, SNMPv3

SNMP est un protocole de la couche application qui utilise les ports UDP 161 (NMS) et 162 (Agent). Il fonctionne selon un système d'échange de messages.

Ces derniers peuvent être de types :

- **Get** : Récupération de la valeur d'un objet de la MIB à partir de l'agent, nécessite au moins les droits en lecture.
- **Set** : Affecter une valeur à l'un des objets MIB grâce à l'agent, nécessite les droits en lecture et écriture.
- **Trap** : Utilisé par l'agent afin de signaler des informations jugées « importantes » à la NMS.

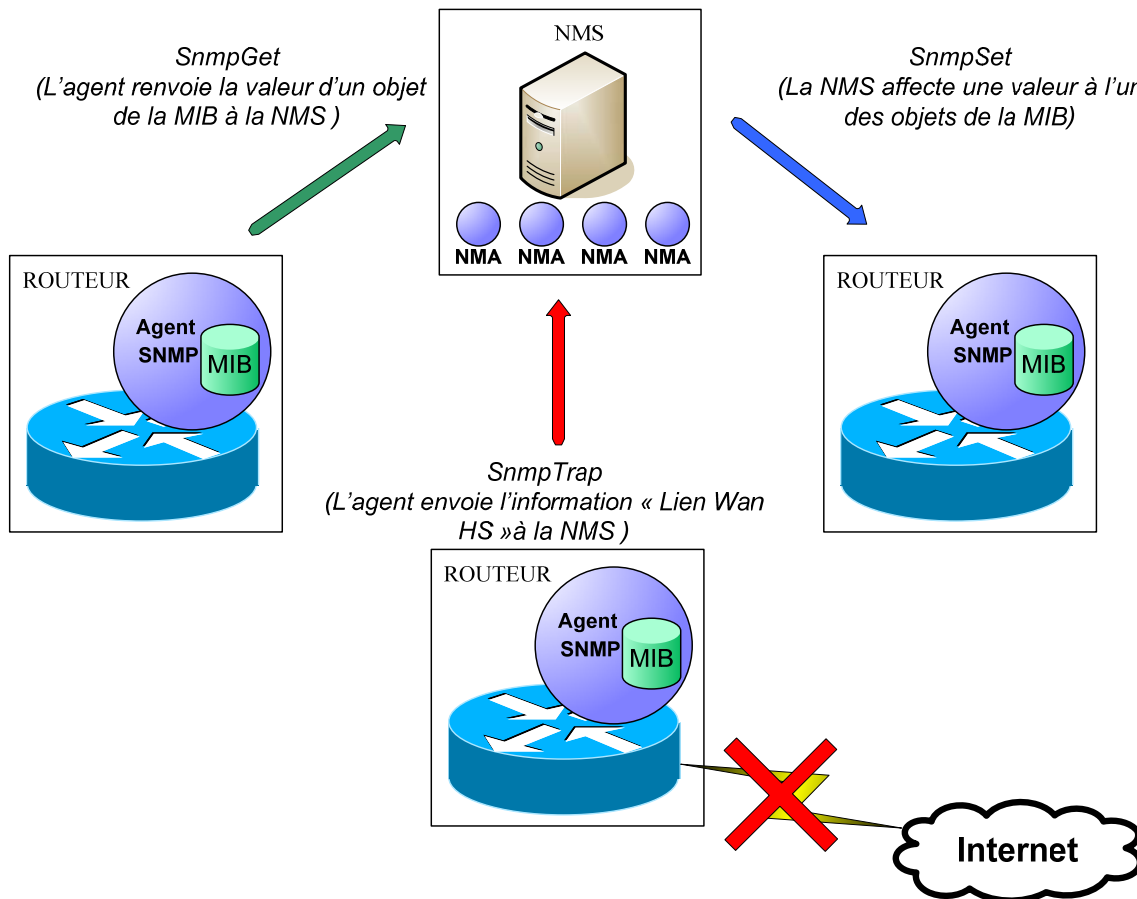


Figure 3 - Les types de messages SNMP

8.4.3. MIB

La MIB est organisée en arborescence définie par la norme SMI³. SMI spécifie également les types de données utilisés pour stocker un objet (entier, chaîne de caractère), la manière dont ces objets sont nommés etc. Chaque élément final de la MIB représente un attribut de l'équipement réseau concerné.

C'est un référentiel contenant une somme considérable d'informations concernant l'équipement. Il existe des MIB standards et propriétaires :

La MIB SMI d'origine est composée de 8 groupes et de 114 objets. Nous en sommes actuellement à la version 2 de la MIB aussi appelée MIB-II.

Les MIB propriétaires sont propres aux équipements du constructeur.

Ci-dessous, un exemple de MIB-II :

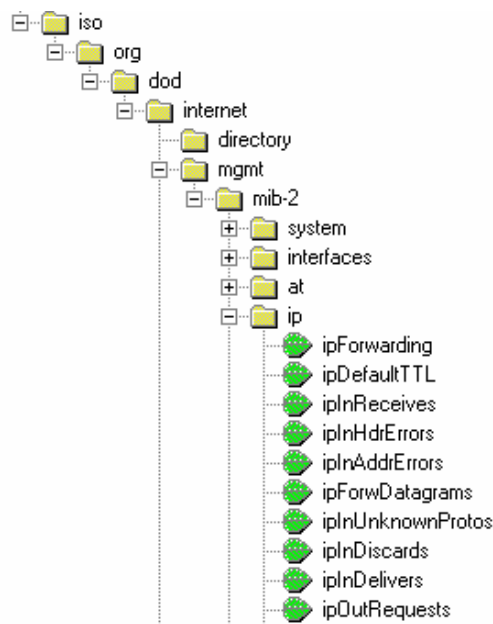


Figure 4 - Représentation logicielle d'une MIB

Chaque feuille de la MIB est identifiée par une OID⁴.

Une OID est une information constituée de valeurs décimales pointées. (Exemple : 1.3.6.1.2.1.4.3).

Chaque valeur décimale de l'OID identifie l'une des branches de la MIB.

Exemple pour l'objet « **ipInReceives** »:

iso(1).org(3).dod(6).internet(1).mgmt(2).mib-2(1).ip(4).ipInReceives(3)

Le schéma ci dessous présente les différents groupes de la MIB ainsi que leurs OID:

³ Structure of Management Information

⁴ Object Identifier

Laboratoire SUPINFO des Technologies Cisco

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

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

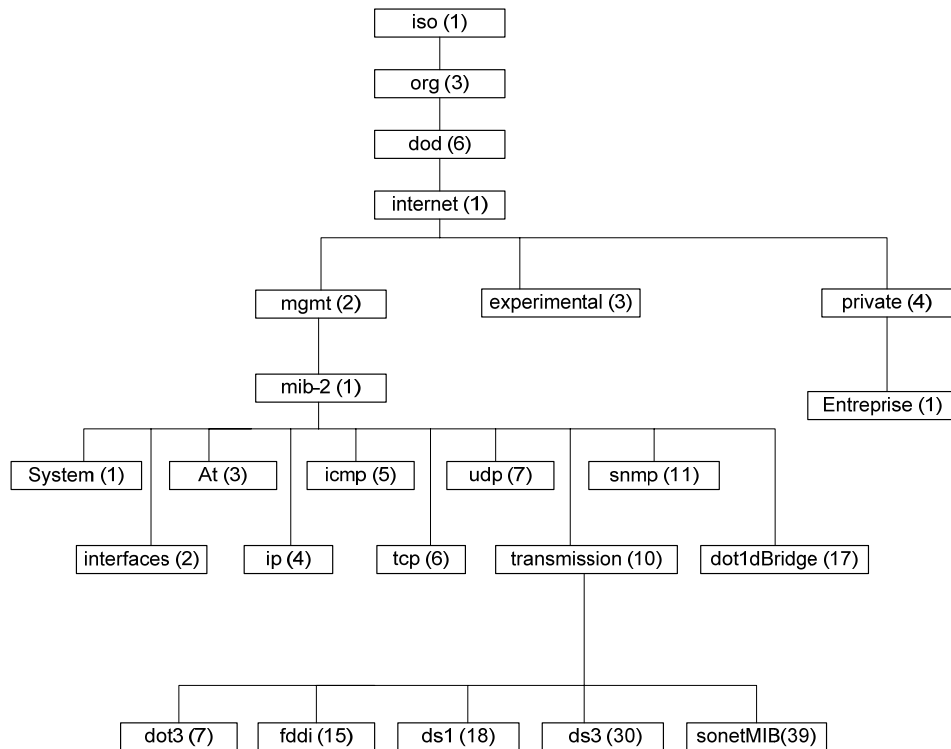


Figure 5 - Représentation des groupes de la MIB et de leur OID

8.4.4. Configuration

Voici les commandes de configuration nécessaires à la communication entre les équipements réseaux et la NMS :

- **snmp-server community {communauté} ro**
 - o Mode de configuration globale
 - o Autorise l'accès en lecture seule à la communauté spécifiée
- **snmp-server community {communauté} rw**
 - o Mode de configuration globale
 - o Autorise l'accès en lecture et écriture à la communauté spécifiée
- **snmp-server location {emplacement}**
 - o Mode de configuration globale
 - o Configure la description de l'emplacement du routeur
- **snmp-server contact {chaîne de caractère}**
 - o Mode de configuration globale
 - o Configure les informations relatives aux personnes à contacter si besoin est
- **snmp-server host {IP de la NMS} {communauté}**
 - o Mode de configuration globale
 - o Spécifie une NMS qui recevra les Traps SNMP
- **snmp-server enable traps snmp [authentication][linkup][linkdown][coldstart] [warmstart]**
 - o Mode de configuration globale
 - o Spécifie le(s) événement(s) qui déclencheront l'envoi des traps

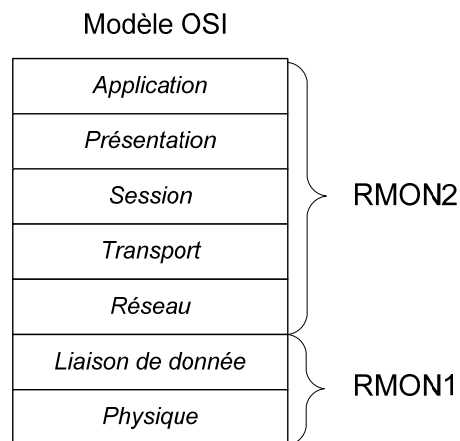
8.4.5. RMON

RMON définit une MIB de surveillance qui complète MIB-II. Cette MIB contient des informations de statistiques obtenues en analysant chaque trame d'un segment du réseau. Pour se faire, des dispositifs de surveillance matérielle (sonde RMON) sont placés sur les segments à surveiller. Ces dispositifs permettent de créer des alarmes définies par l'utilisateur, mais surtout de rassembler une multitude de statistiques vitales grâce à l'analyse approfondie de chaque trame d'un segment.

Avec RMON, l'administrateur peut obtenir des informations relatives à la globalité d'un segment LAN (pourcentage de collisions sur le segment, stations émettant le plus de broadcast etc....). L'administrateur n'a plus pour limite la vision d'information locale et propre à une station exécutant un agent SNMP classique. RMON n'a pas nécessité la modification du protocole SNMP, il n'a suffi pour intégrer RMON que de rajouter des entités dans la MIB. Il existe en deux versions :

RMON1 – Fonctionnant au niveau des couches 1 et 2 du modèle OSI.

RMON2 – Fonctionnant au niveau des couche 3 à 7 du modèle OSI.



- **Groupe des matrices de trafic**

Stocke les erreurs et les statistiques d'utilisation relatives aux paires de nœuds qui communiquent sur le réseau. Il s'agit, par exemple, des erreurs, des octets et des paquets.

- **Groupe des filtres**

Définit un ensemble de filtres afin d'identifier et capturer un flux de paquets correspondant à un schéma distinct.

- **Groupe d'interception des paquets**

Définit la méthode de mise en tampon interne des paquets qui répondent aux critères de filtrage.

- **Groupe des événements**

Consigne des événements à l'intention de l'administrateur. Il s'agit par exemple de rapports personnalisés s'appuyant sur le type d'alarme.

- **Groupe de répertoire des protocoles**

Contient une liste de protocoles supportés par la sonde RMON2, ce groupe est essentiel lorsqu'un agent RMON2 désire savoir quel protocole de communication utilise la sonde RMON2, surtout lorsque les constructeurs des agents et des sondes diffèrent.

- **Groupe de distribution des protocoles**

Contient les données collectées par la sonde, regroupées par protocoles.

- **Groupe de mappage des adresses IP**

Conserve les informations de mappages des adresses mac avec leurs adresses IP.

- **Groupe des hôtes réseaux**

Conserve les statistiques de la couche réseau relatives à une adresse IP.

- **Groupe des matrices de la couche réseau**

Contient des statistiques de la couche réseau concernant les échanges entre deux adresses IP.

- **Groupe des hôtes applicatifs**

Contient des statistiques concernant les protocoles de la couche application d'un hôte.

- **Groupe des matrices de trafic applicatif**

Stocke des statistiques de la couche application relative aux échanges entre deux hôtes.

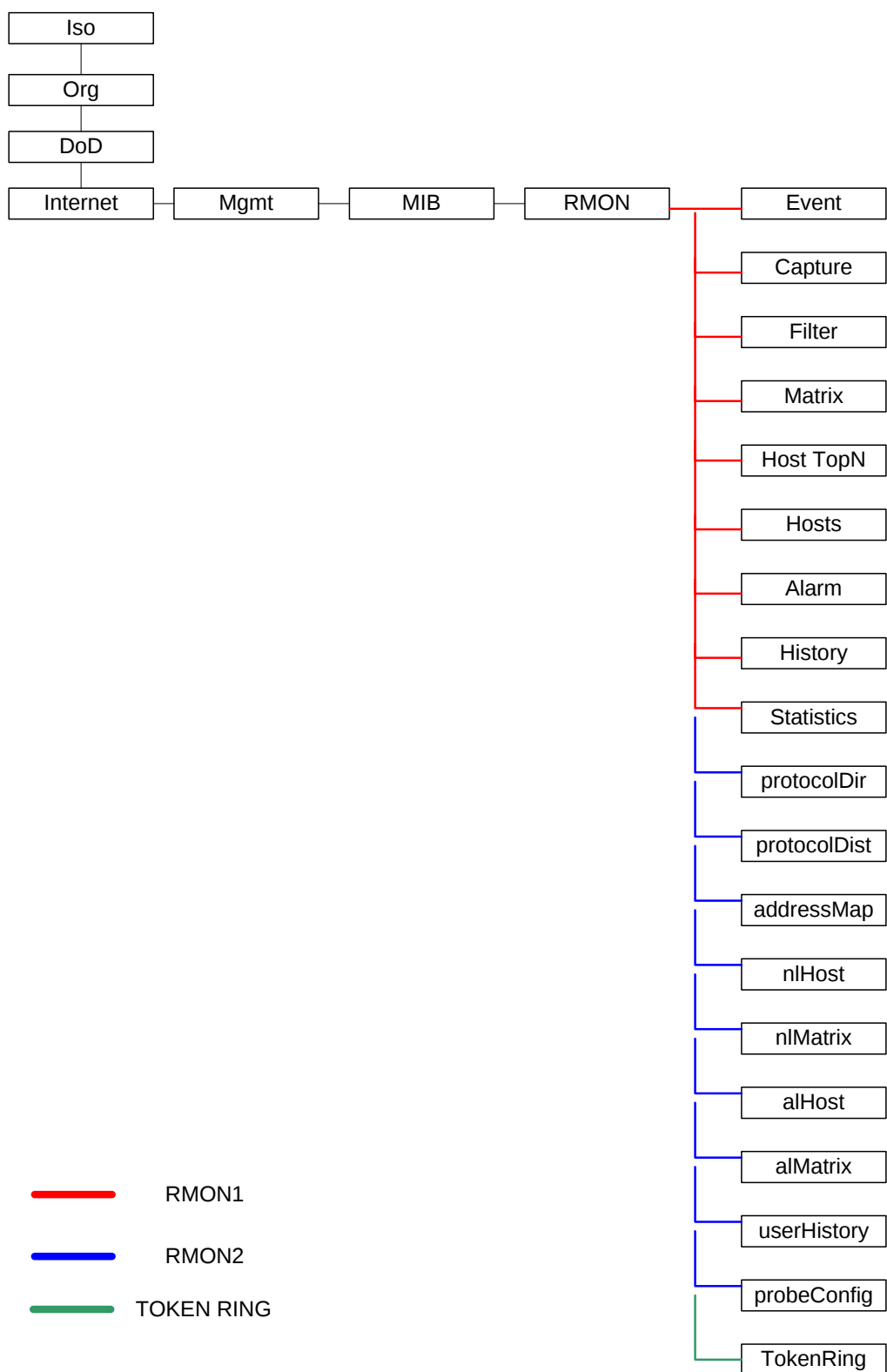
- **Groupe de l'historique des utilisateurs**

Permet à l'administrateur réseau d'archiver les données relatives à n'importe quel hôte du segment, un serveur web ou autres, ...

- **Groupe de configuration de la sonde**

Permet à une application d'un constructeur de configurer à distance la sonde RMON2 d'un autre constructeur.

Ci-dessous, les groupes : RMON1, RMON2, et Token Ring :



8.5. Syslog

8.5.1. Fonctionnement

Syslog est un utilitaire de consignation d'évènements Cisco basé sur l'utilitaire Syslog d'Unix. A l'origine, Syslog avait été développé pour le logiciel Sendmail uniquement. Mais l'utilité de ce dernier était telle que beaucoup d'autres applications se sont mises à l'utiliser. Syslog fonctionne sur un modèle client - serveur. Le port utilisé sur le serveur est le port UDP/514 et la taille des messages ne peut excéder 1024 octets. En 2001, les spécifications de Syslog ont été définies dans la RFC 3164.

Sur un routeur ou commutateur Cisco, les évènements Syslog peuvent être envoyés sur une NMS. Les messages envoyés seront alors de type « non sollicités » (Traps).

Chaque message syslog est horodaté, contient un niveau de gravité ainsi qu'un message de consignation. Ces messages sont parfois la seule manière de résoudre un problème sur les équipements. Il existe 8 niveaux de gravité dans les Traps Syslog (0 à 7). Le niveau 0 étant le plus critique (7 le moins).

Un équipement réseau n'enverra au serveur Syslog que des messages dont la gravité est supérieure (inférieure en chiffre) au seuil défini.

Par défaut, le niveau de gravité est à 6 sur les IOS Cisco. On aura donc tous les messages disponibles excepté ceux de débogage.

Niveau de gravité	Description
0	Urgences
1	Alertes
2	Critique
3	Erreurs
4	Avertissements
5	Notifications
6	Informatifs
7	Déboguages

Niveau par défaut de Cisco IOS →

Par défaut, Cisco IOS adopte le niveau de gravité 6. Ce paramètre est configurable.

8.5.2. Configuration

Pour que la NMS puissent recevoir les traps Syslog d'un équipement, il faut qu'une application serveur Syslog (CiscoWorks2000, Kiwi Syslog...) soit configurée sur celle-ci.

Il faut également configurer le routeur pour l'envoi des événements sur la NMS. Ci-dessous, les différentes commandes de configurations nécessaires sur un routeur 2620xm :

- **logging on**
 - o Mode de configuration globale
 - o Active la consignation des événements
- **logging {nom d'hôte} | {adresse IP de la station}**
 - o Mode de configuration globale
 - o Spécifie au routeur la station NMS recevant les traps Syslog
- **logging trap {debugging | informational | notification | warnings | errors | critical | alerts | emergencies}**
 - o Mode de configuration globale
 - o Configure le niveau de gravité (optionnel)
- **service timestamps log datetime**
 - o Mode de configuration globale
 - o Horodate les messages syslog (optionnel)