

Support de cours

RTEL

Guy Pujolle

Les réseaux de transfert

Les réseaux sont nés du besoin de transporter une information d'une personne à une autre. Pendant longtemps, cette communication s'est faite directement par l'homme, comme dans le réseau postal, ou par des moyens sonores ou visuels. Il y a un peu plus d'un siècle, avec l'apparition du télex puis du téléphone, le concept de réseau est né.

Empruntant d'abord des lignes terrestres de télécommunications, essentiellement composées de fils de cuivre, l'information s'est ensuite également propagée par le biais des ondes hertziennes et de la fibre optique. Il convient d'ajouter à ces lignes de communication le *réseau d'accès*, aussi appelé *boucle locale*, qui permet d'atteindre l'ensemble des utilisateurs potentiels.

Un réseau est donc composé d'un *réseau cœur* et d'un réseau d'accès. Aujourd'hui, on peut dire qu'un réseau est un ensemble d'équipements et de liaisons de télécommunications autorisant le transport d'une information, quelle qu'elle soit, d'un point à un autre, où qu'il soit.

Nous allons examiner maintenant les caractéristiques des topologies et des équipements nécessaires à la réalisation de tels réseaux.

Dans un réseau à transfert de paquets, deux technologies s'opposent : le *routing* et la *commutation*. Ces deux technologies, décrites sommairement précédemment, sont détaillées dans un chapitre ultérieur. Nous nous contentons ici d'introduire leur contexte opérationnel.

Les réseaux sont généralement de type *maillé*, c'est-à-dire qu'il existe toujours au moins deux chemins distincts pour aller d'un point à un autre. En cas de coupure d'une liaison, il y a toujours moyen de permettre la communication. La figure 1 illustre un réseau maillé à transfert de paquets.

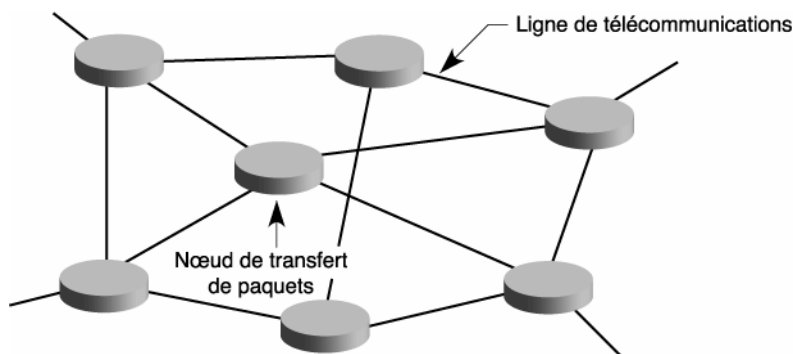


Figure 1. Réseau maillé à transfert de paquets.

Il existe d'autres topologies de réseau que les réseaux maillés, par exemple les réseaux en boucle ou en *bus*. Ces topologies sont généralement adaptées à des situations particulières.

Les réseaux peuvent utiliser des paquets de longueur constante ou variable. La longueur constante permet de déterminer aisément où se trouvent les différents champs et facilite le transfert d'un paquet de l'entrée du nœud à sa sortie. En contrepartie, la fragmentation des données à transporter en petits morceaux demande davantage de

puissance. De plus, si les données à acheminer se limitent à quelques octets, il faut remplir le paquet artificiellement avec des octets dits de remplissage pour qu'il atteigne sa longueur constante.

Les paquets de taille variable s'adaptent bien au contexte applicatif puisqu'ils peuvent être plus ou moins longs suivant la taille du message à transporter. Si le message est très grand, comme dans un transfert de fichier de plusieurs millions d'octets, le paquet a une taille maximale. Si le message est très court, il peut ne donner naissance qu'à un seul paquet, très court également.

Le traitement d'un paquet de longueur variable demande plus de temps que celui d'un paquet de longueur constante, car il faut déterminer les emplacements des champs. Quant au transfert lui-même, il ne peut s'appuyer sur une structure physique simple du nœud.

Les techniques de transfert de paquets

- *ATM (Asynchronous Transfer Mode)* est une technique de transfert de paquets dans laquelle les trames sont très petites et de longueur fixe. L'ATM correspond à une *commutation de trames*.
- Internet utilise la technique de transfert IP (*Internet Protocol*), dans laquelle les paquets sont de longueur variable. Les paquets IP peuvent éventuellement changer de taille lors de la traversée du réseau. La technologie IP fait référence à un *roulage de paquets*.
- Ethernet utilise des trames de longueur variable mais selon une technique différente des deux précédentes. Le transfert Ethernet peut être soit un roulage, soit une commutation.

Les différentes techniques de transfert de paquets ou de trames (*voir encadré*) ont leur origine dans des communautés au départ complètement séparées. L'ATM provient de la communauté des opérateurs de télécommunications et des industriels associés. Les protocoles IP et Ethernet sont nés du besoin des informaticiens de relier leurs machines les unes aux autres pour transférer des fichiers informatiques.

Ces techniques s'appliquent aux différentes catégories de réseaux — depuis les réseaux étendus, appelés *WAN*, jusqu'aux réseaux domestiques, ou *PAN*, en passant par les *MAN* et *LAN*, — en fonction de la distance qui sépare les points les plus éloignés de ces réseaux. La figure 2 illustre ces différentes catégories.

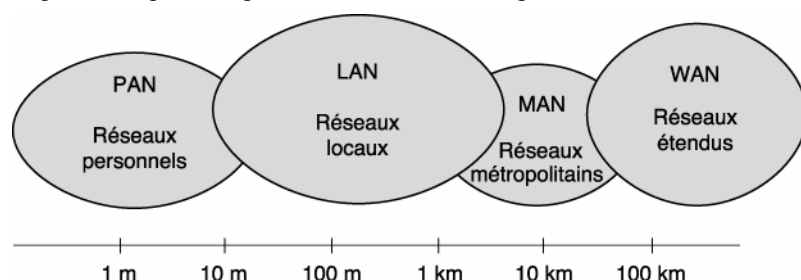


Figure 2. Taille des différentes catégories de réseaux numériques.

Les réseaux de transfert possèdent d'autres propriétés, qui permettent également de les classer, comme les *contrôles de flux* et de *congestion*.

Le contrôle de flux consiste à réguler les flux dans le réseau de façon que le réseau reste fluide et qu'il n'y ait pas de perte de paquets. Le contrôle de congestion s'effectue au moyen d'algorithmes capables de faire sortir le réseau d'un état de congestion. En règle générale, lorsque ces contrôles sont réalisés par l'équipement terminal, il s'agit d'un réseau informatique, puisque les algorithmes de contrôle se trouvent dans une machine informatique située à la périphérie du réseau. Lorsque ces contrôles sont effectués dans le cœur du réseau, il s'agit d'un réseau de télécommunications.

Les catégories de réseaux

Les réseaux proviennent de divers horizons : téléphonie, transport de données et télévision. Chacune de ces catégories d'application tend aujourd'hui à intégrer les autres. Par exemple, les opérateurs de téléphonie se sont intéressés au début des années 80 à l'intégration des données et aujourd'hui à l'intégration de la vidéo. Les réseaux d'interconnexion d'ordinateurs, qui ont donné naissance à Internet, s'intéressent depuis quelque temps à l'intégration de la parole et de la vidéo. De même, les opérateurs de télévision sur câble souhaitent faire transiter de la parole et proposer des connexions à Internet.

Les sections qui suivent détaillent les caractéristiques de ces grandes catégories de réseaux et analysent le rôle que chacune d'elles espère jouer dans les routes et autoroutes de l'information de demain.

Les opérateurs de télécommunications

Les industriels des télécommunications n'ont pas la même vision des architectures de réseau que les opérateurs informatiques ou les câblo-opérateurs. Cela tient aux exigences de leur application de base, la parole téléphonique, auxquelles il est difficile de satisfaire.

La parole téléphonique

Les contraintes de la parole téléphonique sont celles d'une application *temps réel* : l'acheminement de la parole ne doit pas dépasser 300 millisecondes (ms). De ce fait, les signaux doivent être remis au destinataire à des instants précis. Le nom d'application isochrone qui est donné à la parole téléphonique précise bien cette demande forte de *synchronisation*.

La *numérisation* de la parole utilise le *théorème d'échantillonnage*. Ce théorème détermine le nombre d'échantillons nécessaires à une reproduction correcte de la parole sur un support donné. Il doit être au moins égal au double de la bande passante.

Comme la parole téléphonique possède une *bande passante* de 3 200 Hz, ce sont au moins 6 400 échantillons par seconde qui doivent être acheminés au récepteur. La normalisation appelée *MIC* s'appuie sur 8 000 échantillons par seconde, soit un échantillon toutes les 125 microsecondes (μ s). Chaque échantillon est ensuite codé, c'est-à-dire qu'une valeur numérique est donnée à la valeur de fréquence de l'échantillon. La figure 3 illustre ce processus. Le codage est effectué sur 8 bits en Europe et sur 7 bits en Amérique du Nord, ce qui donne des débits respectifs de 64 et 56 Kbit/s.

En réception, l'appareil qui effectue le décodage, le *codec*, doit recevoir les échantillons, composés d'un octet en Europe (7 bits en Amérique du Nord, comme nous venons de le voir), à des instants précis. La perte d'un échantillon de temps en temps n'est pas catastrophique. Il suffit de remplacer l'octet manquant par un octet estimé à partir du précédent et du suivant. Cependant, il ne faut pas que ce processus se répète trop souvent, faute de quoi la qualité de la parole se détériore.

Une autre contrainte des applications de téléphonie concerne le temps de transit à l'intérieur du réseau. Il s'agit en réalité de deux contraintes distinctes mais parallèles : la *contrainte d'interactivité* et la contrainte due aux *échos*.

Les contraintes du transit de la parole téléphonique

La contrainte d'interactivité évalue à 300 ms le retard maximal que peut prendre un signal pour qu'on ait l'impression que deux utilisateurs se parlent dans une même pièce. Si cette limite est dépassée, l'application devient du talkie-walkie.

Dans un réseau symétrique, cette contrainte est de 600 ms aller-retour. Cette limite se réduit cependant à une valeur de moins de 60 ms aller-retour si un phénomène d'écho se produit. Dans les fils métalliques, qui convoient jusqu'au terminal le signal téléphonique, ce dernier est rarement numérisé mais plutôt transporté de façon analogique. Les équipements traversés provoquent des échos qui font repartir le signal en sens inverse.

Les échos ne sont pas perceptibles à l'oreille si le signal revient en moins de 56 ms. Au-delà de cette valeur, en revanche, un effet sonore indésirable rend la conversation pénible. Il existe des appareils qui suppriment l'écho, mais leur utilisation est limitée par leur coût d'installation relativement élevé et par leur portée, limitée à 2 000 à 3 000 km.

La valeur maximale du temps de transit autorisé dans le réseau pour respecter la contrainte d'interactivité (300 ms) est plus de dix fois supérieure à celle liée à la contrainte d'écho (28 ms). Il en résulte que les réseaux sujets aux échos, comme le sont la plupart des réseaux des opérateurs de télécommunications, doivent mettre en œuvre une technique de transfert particulièrement efficace. Dans les réseaux des opérateurs informatiques, dotés de terminaux de type PC, qui annulent les échos, la contrainte de temps de traversée du réseau se situe à 300 ms.

Les choix technologiques des réseaux de télécommunications

Étant donné les contraintes très fortes de l'application de téléphonie, les opérateurs de télécommunications se sont dirigés vers des technologies *circuit*. Ces dernières consistent à mettre en place un chemin dans le réseau avant d'envoyer les informations de l'utilisateur, chemin qui sera suivi par l'ensemble des informations.

La première génération de technologie circuit a été la *commutation de circuits*, toujours présente dans les réseaux téléphoniques actuels. On établit le circuit, après quoi la parole peut y transiter. L'avantage apporté par cette technique est bien évidemment le confort dans lequel l'application se déroule, puisque les ressources du circuit lui sont totalement dédiées. Dans la première génération, le signal reste analogique de bout en bout.

La deuxième génération de technologie circuit conserve la commutation de circuits mais passe au numérique : la parole est numérisée, et ce sont des octets qui transitent. La numérisation s'effectue le plus souvent à l'entrée du réseau, la boucle locale restant encore fortement analogique.

Cette génération, que l'on utilise toujours aujourd'hui, a été améliorée par une meilleure utilisation de la boucle locale. Sur le fil métallique qui sert à faire passer la parole analogique, on est capable, grâce à des *modems*, de faire transiter des flux de plus en plus puissants, pouvant atteindre 2 Mbit/s sur des distances de plusieurs kilomètres.

De ce fait, les opérateurs de télécommunications ont pu intégrer les données en même temps que la parole téléphonique. Ce passage n'a pas été aussi immédiat et est passé par un stade intermédiaire, le *RNIS (réseau numérique à intégration de services)*, consistant à fournir deux lignes simultanées, l'une pour le téléphone, l'autre pour les données. Nous présentons en détail cette solution intermédiaire dans une autre section plus loin, car elle est toujours commercialisée, malgré une certaine désaffection.

Avec la troisième génération, on travaille directement en mode paquet. Le terminal place toutes les informations, qu'elles soient téléphoniques, vidéo ou de données, dans des paquets et émet ces paquets sur la boucle locale vers le réseau cœur.

Cette technologie porte l'intégration à son paroxysme puisqu'il n'est plus possible de distinguer les applications entre elles au moment de leur transport. Les besoins actuels de *qualité de service* impliquent toutefois d'en revenir à une classification permettant d'identifier les paquets prioritaires et de leur octroyer une gestion différenciée.

Le futur des réseaux de télécommunications

Les opérateurs de télécommunications reconnaissent que l'avenir appartient à une troisième génération de réseau utilisant le paquet IP. À la différence des opérateurs informatiques, ils souhaitent toujours mettre en place un chemin puis faire transiter les paquets IP d'une façon sécurisée et dans l'ordre. Pour cela, ils ont besoin d'un réseau de signalisation capable de tracer le chemin. La révolution en cours dans les télécoms réside finalement dans la découverte que le meilleur réseau de signalisation est... le réseau IP.

Le futur des réseaux de télécommunications verra associés un réseau à commutation de paquets et des chemins qui seront ouverts par un réseau de type Internet. Ce dernier sera alors considéré comme un réseau de signalisation, avec pour objectif d'indiquer la meilleure route numérique à prendre.

Les technologies sous-jacentes sont la commutation de circuits, avec une signalisation nommée *CCITT n°7*, puis le passage à des techniques de *circuit virtuel*, avec les protocoles *X.25*, *relais de trames*, *ATM (Asynchronous Transfer Mode)*, qui utilisent des signalisations spécifiques, avant de converger vers *MPLS (MultiProtocol Label Switching)*, qui utilise une signalisation IP.

Les opérateurs de réseaux informatiques

Les opérateurs de réseaux informatiques correspondent en grande partie aux *ISP (Internet Service Provider)*. Développée pour le monde Internet, la technologie qu'ils utilisent consiste à encapsuler l'information dans des paquets dits *IP (Internet Protocol)*.

Les paquets IP comportent l'adresse complète du destinataire. Le paquet n'est donc jamais perdu dans le réseau puisqu'il connaît sa destination. Dans la commutation, à l'inverse, la trame ou le paquet ne sont munis que d'une référence, laquelle ne donne aucune information sur la destination. En cas de problème, il faut demander à la signalisation d'ouvrir un nouveau chemin.

Dans les réseaux informatiques, il n'y a pas besoin de signalisation. Cela allège considérablement le coût du réseau. À chaque nœud, il suffit de remonter à la couche IP pour retrouver l'adresse complète du destinataire et aller consulter la *table de routage*, qui indique la meilleure porte de sortie du nœud en direction du destinataire.

L'avantage évident des technologies informatiques réside dans la grande simplicité du réseau. Ce dernier est construit autour d'équipements, appelés *routeurs*, dont la seule fonction est de router les paquets vers la meilleure porte de sortie possible. Cette porte pouvant varier en fonction du trafic, la table de routage doit être mise à jour régulièrement.

Une autre caractéristique des réseaux informatiques consiste en un contrôle effectué uniquement par les équipements terminaux, sans aucune participation des routeurs internes au réseau. Les *algorithmes de contrôle* se trouvent dans les PC et autres machines terminales.

Les algorithmes de contrôle sont le plus souvent fondés sur des *fenêtres de contrôle*, c'est-à-dire des valeurs maximales du nombre de paquets qui peuvent être envoyés sans *acquittement*. L'équipement terminal doit détecter

d'après le temps de retour de l'acquittement si le réseau est congestionné ou non et augmenter ou diminuer sa fenêtre de contrôle.

Les choix technologiques des réseaux informatiques

Les solutions à la disposition des opérateurs informatiques pour garantir une qualité de service sont assez limitées.

Une première solution consiste à surdimensionner le réseau de telle sorte que les paquets ne soient jamais retardés par des attentes dans les routeurs. Cette solution est acceptable si la technologie offre, à bon prix, des capacités suffisantes. C'était le cas au début des années 2000, au cours desquelles l'essor de la fibre optique et des techniques de *multiplexage en longueur d'onde*, a permis d'atteindre des débits de plusieurs téraoctets par seconde. L'augmentation des capacités est aujourd'hui freinée par la saturation du support optique, car il ne reste pratiquement plus de bande passante disponible dans la fibre optique. Si la solution de surdimensionnement a pu être acceptable au début des années 2000, elle devient de plus en plus difficile à réaliser du fait que la technologie ne fait plus de progrès fondamentaux. Un rattrapage des débits s'opère toutefois avec l'arrivée massive de techniques d'accès à haut débit, comme l'*ADSL*.

Une seconde solution consiste à ne surdimensionner le réseau que pour les clients ayant besoin de temps de réponse garantis. Pour y parvenir, il faut classer les flux. La solution proposée par le monde Internet revient à spécifier trois grandes classes de flux : une classe de plus haute priorité, appelée premium ou platinum, une classe de basse priorité, qui n'offre aucune garantie, et une classe intermédiaire, qui propose une garantie sur le taux de perte mais pas sur le temps de transit à l'intérieur du réseau.

L'opérateur réseau doit pratiquer des coûts de connexion tels que tous les clients ne choisissent pas la classe de plus haute priorité et que cette dernière ne dépasse pas une quinzaine de pour cent de la capacité du réseau. À ces conditions, les clients prioritaires ne sont pas gênés par les autres et peuvent espérer une très bonne qualité de service. Cette solution revient moins cher que le surdimensionnement généralisé pour l'ensemble des utilisateurs. Elle est cependant plus complexe, car il faut introduire un coût pour différencier les clients, faute de quoi tous les clients réclameraient la plus haute priorité.

L'introduction de priorités nécessite une nouvelle génération de routeurs capables de les gérer, c'est-à-dire de reconnaître la valeur indiquée dans un champ de priorité du paquet IP et de placer les paquets de plus haute priorité en tête des files d'attente.

Cette solution est bien couverte par la proposition *DiffServ (Differentiated Services)*, qui classe les clients en trois classes. La classe intermédiaire peut comporter jusqu'à douze sous-classes.

Le futur des réseaux informatiques

Pour les opérateurs de réseaux informatiques, le futur est assez clair : accroître les débits. L'idée est de continuer à utiliser des routeurs mais à en augmenter la capacité de traitement pour aller vers ce que l'on appelle des *gigarouteurs*, ou même maintenant des térarouteurs. Un gigarouteur peut traiter un milliard de paquets par seconde et un térarouteur mille milliards de paquets par seconde. Pour cela, il faut trouver des solutions de traitement des adresses de façon à déterminer la bonne ligne de sortie en un laps de temps le plus court possible.

Les liaisons entre les gigarouteurs seront assurées par des techniques à très haut débit, comme le *Gigabit Ethernet*, le 10 Gigabit Ethernet ou les lignes *SONET* à 2,5 Gbit/s, 10 Gbit/s, voire 40 Gbit/s.

Les opérateurs vidéo

On désigne sous le terme d'opérateurs vidéo les diffuseurs et câblo-opérateurs qui mettent en place les réseaux terrestres et hertziens de diffusion des canaux de télévision. Les câblo-opérateurs se chargent de la partie terrestre, et les télédiffuseurs de la partie hertzienne. Ces infrastructures de communication permettent de faire transiter vers l'équipement terminal les canaux vidéo. Étant donné la *largeur de bande* passante réclamée par ces applications, ces canaux demandent des débits très importants.

La diffusion de programmes de télévision s'effectue depuis de longues années par le biais d'émetteurs hertziens, avec des avantages et des inconvénients. De nouvelles applications vidéo ont fait leur apparition ces dernières années, dont la qualité vidéo va d'images saccadées de piètre qualité jusqu'à des images animées en haute définition.

La classification admise pour les applications vidéo est la suivante :

- **Visioconférence.** Se limite à montrer le visage des correspondants. Sa définition est relativement faible puisqu'on diminue le nombre d'images par seconde pour gagner en débit. Le *signal* produit par une visioconférence se transporte aisément sur un canal numérique à 128 Kbit/s, et sa compression est simple à réaliser. On peut abaisser le débit jusqu'à 64 Kbit/s, voire moins, si l'on ne redoute pas une sérieuse baisse de la qualité.
- **Télévision numérique.** Sa qualité correspond ordinairement à un canal de 4 ou 5 MHz de bande passante en *analogique*. La numérisation sans compression de ce canal, en utilisant par exemple le théorème d'échantillonnage, produit un débit de plus de 200 Mbit/s. Après compression, le débit peut descendre à 2 Mbit/s, pratiquement sans perte de qualité. On peut, avec une compression poussée, aller jusqu'à des débits de 64 Kbit/s, mais avec une qualité fortement dégradée. De plus, à de tels débits, les erreurs en ligne deviennent gênantes, car elles perturbent l'image au moment de la décompression. L'optimum est un compromis entre une forte compression et un taux d'erreur de 10^{-9} (en moyenne une erreur tous les 109 bits transmis), ce qui ne détruit qu'une infime fraction de l'image, sans nuire à sa vision. Le principal standard pour la transmission d'un canal de télévision numérique est aujourd'hui MPEG-2.
- **Télévision haute définition.** Demande des transmissions à plus de 500 Mbit/s si aucune compression n'est effectuée. Après compression, la valeur peut tomber à 35 Mbit/s, voire 4 Mbit/s.
- **Vidéoconférence** (à ne pas confondre avec la visioconférence). Approche la qualité du cinéma et demande des débits considérables. C'est la raison pour laquelle ce type de vidéo ne sera intégré que plus tard dans les applications multimédias.

Câble coaxial et fibre optique

Les applications utilisant la vidéo sont nombreuses : elles vont de la télésurveillance à la vidéo à la demande en passant par la messagerie vidéo et la télévision.

Les réseaux câblés, installés par les diffuseurs sur la partie finale du réseau de distribution, utilisent un support physique de type *câble coaxial*, le *CATV*.

Ce câble à très grande bande passante peut également être utilisé pour acheminer aux utilisateurs des informations diversifiées, comme la parole ou les données informatiques, en plus de l'image. Aujourd'hui, ces réseaux câblés sont exploités en analogique et très rarement en numérique. À long terme, ils pourraient absorber plusieurs dizaines de mégabits par seconde, ce qui permettrait de véhiculer sans problème les applications multimédias.

Les câblo-opérateurs ont l'avantage de pouvoir atteindre de nombreux foyers et de constituer ainsi une porte d'entrée vers l'utilisateur final. Le câblage CATV est une des clefs de la diffusion généralisée de l'information. C'est pourquoi il a été privilégié pendant de nombreuses années par les opérateurs de télécommunications.

La fibre optique tend aujourd'hui à remplacer le câble coaxial par son prix attractif et sa bande passante encore plus importante.

Internet

L'Internet provient du concept d'*Inter-Network*, c'est-à-dire de la volonté de relier des réseaux hétérogènes entre eux par l'adoption d'un protocole unique, *IP (Internet Protocol)*. Le protocole IP a pour objectif de normaliser la façon de transporter les paquets dans un réseau. Internet est un réseau routé dans lequel les nœuds de transfert sont appelés des routeurs. Les routeurs sont munis d'une table de routage qui permet aux paquets entrants de trouver la meilleure sortie possible.

Internet est un réseau sans signalisation, c'est-à-dire sans chemin. Chaque paquet contient l'adresse complète du destinataire et est donc autonome dans le réseau. On appelle parfois ces paquets autonomes des *datagrammes*.

Une caractéristique du réseau Internet, que nous avons déjà mentionnée, concerne le contrôle du réseau effectué par la machine terminale. La fenêtre de contrôle est déterminée par le temps aller-retour d'un paquet, depuis son envoi jusqu'à l'arrivée de l'acquiescement. Si ce temps aller-retour devient trop important, l'équipement terminal interprète cet événement comme une congestion et diminue de façon drastique la fenêtre de contrôle.

La qualité de service d'Internet est la même pour l'ensemble des utilisateurs. On la nomme *best effort* pour indiquer que le réseau fait au mieux par rapport aux demandes des utilisateurs. Cette qualité de service ne permet pas de

prendre en charge des applications à fortes contraintes puisque le réseau est partagé à part égale entre tous les utilisateurs. Plus il y a d'utilisateurs connectés, plus les performances de tous se dégradent.

Une solution pour apporter de la qualité de service dans Internet consiste à surdimensionner le réseau. Dans ce cas, la qualité de service est bonne pour tous les utilisateurs. En effet, le réseau étant à peu près vide, puisqu'il est surdimensionné, chaque client le traverse en un temps très court. Cette solution est cependant difficile à évaluer économiquement. Si elle limite le nombre d'ingénieurs réseau nécessaires à la gestion du réseau, son coût peut devenir très élevé du fait que les composants à très haut débit restent chers.

Les solutions préconisées par le monde Internet pour obtenir de la qualité de service sont de trois types :

- Rendre le flot applicatif dynamique, de sorte à l'adapter aux possibilités du réseau. Si le débit du réseau diminue, on compresse plus fortement l'application.
- Adapter le réseau à l'application au moyen de la technique *IntServ (Integrated Service)*. Chaque flot se voit affecter des ressources réseau adaptées à sa demande. Cette solution a le défaut de ne pas passer l'échelle, c'est-à-dire d'être limitée en nombre de flots.
- Adapter le réseau à l'application mais avec la capacité de passer à l'échelle. C'est la technique *DiffServ*, dans laquelle les clients sont regroupés en grandes classes de service. À chaque classe correspondent des ressources.

Il existe d'autres solutions pour réaliser un réseau Internet contrôlé. L'une d'elles consiste à utiliser le protocole IP dans un environnement privé, appelé réseau *intranet*.

Une autre solution est d'utiliser un réseau d'opérateur de télécommunications dans lequel les paquets IP entrants sont immédiatement encapsulés dans une trame puis commutés sur un chemin. À la sortie, la trame est décapsulée, et le paquet IP est remis à l'utilisateur.

Les réseaux ATM

La technique de transfert ATM (*Asynchronous Transfer Mode*) s'est imposée dans les années 90 pour l'obtention de la qualité de service. Sa première caractéristique est de limiter la taille des paquets à une valeur constante de 53 octets de façon à garantir son traitement rapide dans les nœuds.

Cette solution offre le meilleur compromis pour le transport des applications qui transitent sur les réseaux. Le transport des applications isochrones s'obtient plus facilement du fait de la petite taille des paquets, ou *cellules*, qui engendre des temps de paquetsation et de dépaquetsation faibles.

Si les cellules permettent de transporter facilement les données *asynchrones*, c'est au prix d'une fragmentation assez poussée. Lorsque le contrôle de flux est strict, le temps de transport dans le réseau est à peu près égal au temps de propagation, ce qui permet de retrouver simplement la synchronisation en sortie.

Les réseaux ATM sont des réseaux commutés de *niveau trame*, c'est-à-dire qui commutent des trames. Un système de signalisation spécifique permet d'ouvrir et de fermer les circuits virtuels.

Des qualités de service sont affectées aux circuits virtuels pour donner la possibilité aux utilisateurs de faire transiter des applications avec contraintes.

Les réseaux Ethernet

Les réseaux Ethernet proposent une architecture différente. En particulier, ils définissent un autre format de trames, qui s'est imposé par l'intermédiaire des réseaux locaux, ou LAN (*Local Area Network*). Le format de la trame Ethernet est illustré à la figure 4.

Les adresses contenues dans la *trame Ethernet* se composent de deux champs de 3 octets chacun, le premier indiquant un numéro de constructeur et le second un numéro de série. La difficulté liée à ces adresses, spécifiques de chaque carte Ethernet introduite dans un PC, consiste à déterminer l'emplacement du PC. C'est la raison pour laquelle la trame Ethernet est utilisée dans un univers local, où il est possible de diffuser la trame à l'ensemble des récepteurs, le récepteur de destination reconnaissant son adresse et gardant la copie reçue.

Pour utiliser la trame Ethernet dans de grands réseaux, il faut soit trouver une autre façon d'interpréter l'adresse sur 6 octets, soit ajouter une adresse complémentaire.

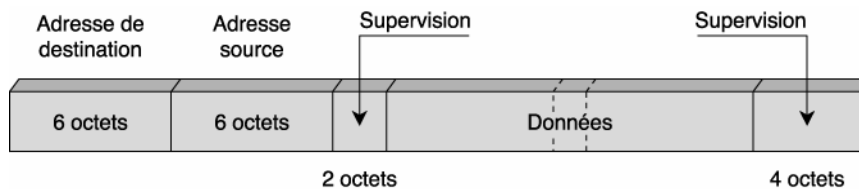


Figure 4. Format de la trame Ethernet.

Les techniques de transfert

La commutation de circuits

Un circuit peut être comparé à un tuyau placé entre un émetteur et un récepteur. Ce tuyau peut être constitué de fils métalliques, de fibre optique ou d'onde hertzienne. Le circuit n'appartient qu'aux deux entités qui communiquent. Le circuit le plus simple correspond à un tuyau posé entre deux points. Appelons-le circuit élémentaire. Il est possible de réaliser des circuits plus complexes en ajoutant des circuits élémentaires les uns derrière les autres. Cela donne un nouveau tuyau, dans lequel les différentes parties peuvent être réalisées à partir de matériaux différents (métal, fibre, fréquence). Le circuit doit être ouvert pour que les informations puissent transiter. Il reste ouvert jusqu'au moment où l'un des deux participants interrompt la communication. Cela a pour effet de relâcher les ressources affectées à la réalisation du circuit. Si les deux correspondants n'ont plus de données à se transmettre pendant un certain temps, la liaison reste inutilisée, et les ressources ne peuvent être employées par d'autres utilisateurs.

La commutation de circuits désigne le mécanisme consistant à rechercher les différents circuits élémentaires pour réaliser un circuit plus complexe. Cette opération se réalise grâce à la présence de nœuds, appelés commutateurs de circuits ou *autocommutateurs*, dont le rôle consiste à choisir un tuyau libre en sortie pour le rabouter au tuyau entrant, permettant ainsi de mettre en place le circuit nécessaire à la communication entre deux utilisateurs.

Un réseau à commutation de circuits consiste en un ensemble d'équipements, les autocommutateurs, et de liaisons interconnectant ces autocommutateurs, dont le but consiste à mettre en place des circuits à la demande des utilisateurs. Un tel réseau est illustré à la figure 5.

Pour mettre en place un circuit, il faut propager un ordre demandant aux autocommutateurs de mettre bout à bout des circuits élémentaires. Ces commandes et leur propagation s'appellent la *signalisation*.

La signalisation peut être dans la bande ou hors de la bande. Une signalisation dans la bande indique que la commande d'ouverture d'un circuit transite d'un autocommutateur à un autre en utilisant le circuit ouvert à la demande de l'utilisateur. La construction du circuit se fait par la propagation d'une commande dotée de l'adresse du destinataire (par exemple, le numéro de téléphone qui a été composé sur le cadran) et empruntant le circuit en cours de construction.

La signalisation hors bande indique le passage de la commande de signalisation dans un réseau différent du réseau à commutation de circuits dont elle est issue. Ce réseau externe, appelé *réseau sémaphore*, relie tous les autocommutateurs entre eux de façon que la commande d'ouverture puisse transiter d'un autocommutateur à un autre.

Ces deux types de signalisation sont illustrés à la figure 6.

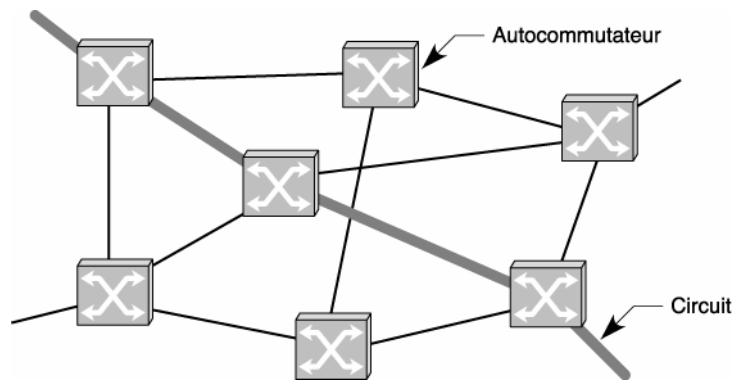


Figure 5. Réseau à commutation de circuits.

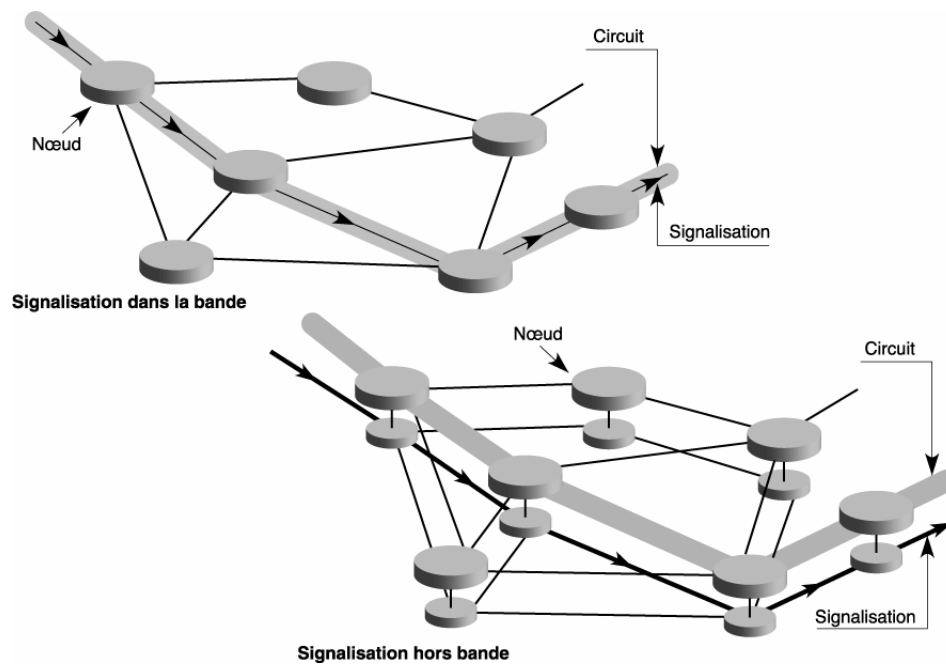


Figure 6. Signalisation dans la bande et signalisation hors bande.

Les réseaux sémaphores modernes

Les réseaux à commutation de circuits actuels utilisent une signalisation encore améliorée. La demande d'ouverture de circuit n'établit pas tout de suite le circuit. Elle est d'abord transmise vers le récepteur par le réseau sémaphore de façon soit à déclencher la sonnerie d'un téléphone, soit à s'apercevoir que le destinataire est occupé. Dans ce dernier cas, une commande de retour part vers l'émetteur par le réseau sémaphore pour déclencher la tonalité d'occupation. Lorsque le récepteur décroche son téléphone, une commande de signalisation repart, mettant en place le circuit. On voit ainsi que l'utilisation d'un réseau sémaphore, c'est-à-dire d'une signalisation hors bande, permet d'utiliser beaucoup mieux les circuits disponibles. Dans les anciens réseaux téléphoniques, avec leur signalisation dans la bande, il fallait quelques secondes pour mettre en place le circuit et s'apercevoir, par exemple, que le destinataire était occupé, d'où encore quelques secondes pour détruire le circuit inutile. Par le réseau sémaphore, cette opération s'effectue en moins de 500 ms, avec deux commandes acheminées et en n'utilisant aucun circuit.

Le transfert de paquets

Cette section décrit brièvement comment les informations sont paquetisées et les paquets acheminés par un réseau de transfert contenant des nœuds. Le transfert de paquets correspond à la technique utilisée pour réaliser cet acheminement. Deux méthodes principales sont mises en œuvre pour cela : le routage et la commutation. Lorsque le routage est choisi, les nœuds s'appellent des routeurs, et le réseau correspond à un réseau à *routage de paquets*.

Lorsque le choix porte sur la commutation, les nœuds s'appellent des commutateurs et le réseau correspond à un réseau à *commutation de paquets*.

Le rôle d'un nœud de transfert peut se résumer à trois fonctions :

- l'analyse de l'en-tête du paquet et sa traduction ;
- la commutation ou routage vers la bonne ligne de sortie ;
- la transmission des paquets sur la liaison de sortie choisie.

Un nœud de transfert, qui peut être un commutateur ou un routeur, est illustré à la figure 7. Le schéma indique notamment le choix à effectuer par la file d'entrée du nœud pour diriger au mieux les paquets vers l'une des trois files de sortie de cet exemple.

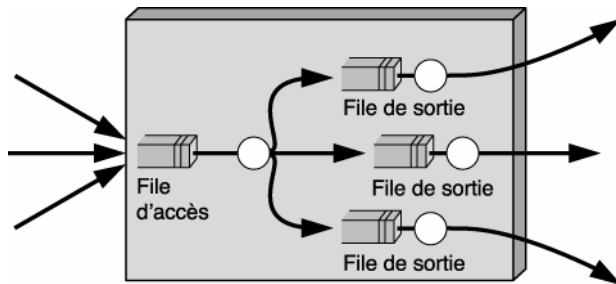


Figure 7. Un nœud de transfert.

La première file d'attente du nœud de transfert examine l'en-tête de chaque paquet pour identifier le port de sortie. Cette première file d'attente est parfois appelée file de commutation puisque les paquets sont, en quelque sorte, commutés vers une ligne de sortie. Il est toutefois préférable d'appeler cette fonction un transfert, de façon à ne pas confondre routeur et commutateur. Les paquets sont donc transférés individuellement, et toutes les fonctions du nœud de transfert sont effectuées au rythme du paquet.

Il existe une grande diversité de solutions permettant de réaliser un nœud de transfert. Dans tous les cas, il faut réaliser une fonction de stockage, qui peut se trouver à l'entrée, à la sortie ou le long de la chaîne de transfert.

Un débat sans fin oppose les mérites respectifs des routeurs et des commutateurs parce qu'ils symbolisent deux manières opposées d'acheminer l'information à l'intérieur d'un réseau maillé. Les deux solutions présentent bien sûr des avantages et des inconvénients. Parmi les avantages, citons notamment la souplesse pour le routage et la puissance pour la commutation.

Les techniques de transfert de l'ATM ou du *relais de trames* utilisent une commutation. Internet préfère le routage. Ethernet se place entre les deux, avec un routage quasiment fixe, qui ressemble à une commutation, d'où le nom de commutation Ethernet.

Les routeurs

Dans un routeur, le paquet qui arrive doit posséder l'adresse complète du destinataire, de sorte que le nœud puisse décider de la meilleure ligne de sortie à choisir pour l'envoyer vers un nœud suivant. Une décision de routage survient donc, consistant à aller consulter une table de routage, dans laquelle sont répertoriées toutes les adresses susceptibles d'être atteintes sur le réseau. La décision de router prend du temps : non seulement il faut trouver la bonne ligne de la table de routage correspondant à l'adresse du destinataire, mais, surtout, il faut gérer cette table de routage, c'est-à-dire la maintenir à jour pour que les routes soient les meilleures possibles.

Les différents types d'adresses

Une adresse complète correspond à l'adresse d'un utilisateur du réseau, et plus généralement à un moyen pour déterminer où se trouve cet utilisateur. Une adresse téléphonique désigne un utilisateur fixe ou mobile, mais la mobilité ne permet plus de déterminer l'emplacement de l'émetteur. L'adresse Internet est une adresse d'utilisateur qui ne permet de déterminer que très partiellement l'emplacement géographique de l'utilisateur.

Le routage est une solution de transfert de l'information agréable et particulièrement flexible, qui permet aux paquets de contourner les points du réseau en congestion ou en panne. Le paquet possède en lui tout ce qu'il lui faut pour continuer sa route seul. C'est la raison pour laquelle on appelle parfois les paquets des *datagrammes*.

Le routage peut poser des problèmes de différents ordres. Un premier inconvénient vient de la réception des paquets par le récepteur dans un ordre qui n'est pas forcément celui de l'envoi par l'émetteur. En effet, un paquet peut en

doubler un autre en étant aiguillé, par exemple, par une route plus courte, découverte un peu tardivement à la suite d'une congestion. Un deuxième inconvénient provient de la longueur de l'adresse, qui doit être suffisamment importante pour pouvoir représenter tous les récepteurs potentiels du réseau. La première génération du protocole Internet, IPv4, n'avait que 4 octets pour coder l'adresse, tandis que la seconde, IPv6, en offre 16. La troisième difficulté réside dans la taille de la table de routage. Si le réseau a beaucoup de clients, le nombre de lignes de la table de routage peut être très important, et de nombreux *paquets de supervision* sont nécessaires pour la maintenir à jour.

Pour réaliser des routages efficaces, il faut essayer de limiter le nombre de lignes des tables de routage — nous verrons qu'il en va de même pour les tables de commutation —, si possible à une valeur de l'ordre de 5 000 à 10 000, ces valeurs semblant être le gage d'une capacité de routage appréciable.

Les commutateurs

Les commutateurs acheminent les paquets vers le récepteur en utilisant des *références*, que l'on appelle aussi identificateurs, étiquettes ou « labels », de circuit ou de chemin.

Les tables de commutation sont des tableaux, qui, à une référence, font correspondre une ligne de sortie. Noter que seules les communications actives entre utilisateurs comportent une entrée dans la table de commutation. Cette propriété limite la taille de la table. De façon plus précise, le nombre de lignes d'une table de commutation est égal à 2^n , n étant le nombre d'éléments binaires donnant la référence.

L'avantage de la technique de commutation provient de la puissance offerte pour commuter les paquets du fait de l'utilisation de références. Ces dernières réduisent la taille de la table de commutation, car seules les références actives y prennent place. De surcroît, le nœud n'a pas à se poser de questions sur la meilleure ligne de sortie, puisqu'elle est déterminée une fois pour toutes. Un autre avantage de cette technique vient de ce que la zone portant la référence demande en général beaucoup moins de place que l'adresse complète d'un destinataire. Par exemple, la référence la plus longue, celle de l'ATM, utilise 28 bits contre 16 octets pour l'adresse IPv6).

La difficulté engendrée par les commutateurs est liée au besoin d'ouvrir puis de refermer les chemins ainsi que les références posées pour réaliser ces chemins. Pour cela, il faut faire appel à une signalisation du même type que celle mise au point pour les réseaux à commutation de circuits. Un paquet de signalisation part de l'émetteur en emportant l'adresse complète du récepteur. Ce paquet contient également la référence de la première liaison utilisée. Lorsque ce paquet de commande arrive dans le premier nœud à traverser, il demande à la table de routage du nœud de lui indiquer la bonne ligne de sortie. On voit donc qu'un commutateur fait aussi office de routeur pour le paquet de signalisation et qu'il doit comporter en son sein les deux systèmes intégrés.

Le temps nécessaire pour ouvrir la route lors de la signalisation n'est pas aussi capital que le temps de transfert des paquets sur cette route : le temps pour mettre en communication deux correspondants n'a pas les mêmes contraintes que le temps de traversée des paquets eux-mêmes. Par exemple, on peut prendre deux secondes pour mettre en place un circuit téléphonique, mais il est impératif que les paquets traversent le réseau en moins de 300 ms.

Le fonctionnement d'une commutation montre toute son efficacité lorsque le flot de paquets à transmettre est important. À l'inverse, le routage est plus efficace si le flot est court.

Circuit virtuel

Le chemin déterminé par les références s'appelle un *circuit virtuel*, par similitude avec un circuit classique. Les paquets transitent toujours par le même chemin, les uns derrière les autres, comme sur un circuit, mais le circuit est dit virtuel parce que d'autres utilisateurs empruntent les mêmes liaisons et que les paquets doivent attendre qu'une liaison se libère avant de pouvoir être émis. La figure 8 illustre un circuit virtuel.

Lors d'une panne d'une ligne ou d'un nœud, un reroutage intervient pour redéfinir un chemin. Ce reroutage s'effectue à l'aide d'un paquet de signalisation, qui met en place les nouvelles références que doivent suivre les paquets du circuit virtuel qui a été détruit.

En règle générale, l'état d'un circuit virtuel est spécifié en « dur » (*hard-state*). Cela signifie que le circuit virtuel ne peut être fermé qu'explicitement. Tant qu'une signalisation ne demande pas la fermeture, le circuit virtuel reste ouvert.

La même technique, mais avec des états « mous » (*soft-state*) — que l'on se gardera de considérer comme un circuit virtuel, même si leurs propriétés sont identiques —, détruit automatiquement les références et, par conséquent, le chemin, si ce dernier n'est pas utilisé pendant un temps déterminé au départ. Si aucun client n'utilise le chemin, il doit, pour rester ouvert, être rafraîchi régulièrement : un paquet de signalisation spécifique indique aux nœuds de garder encore valables les références.

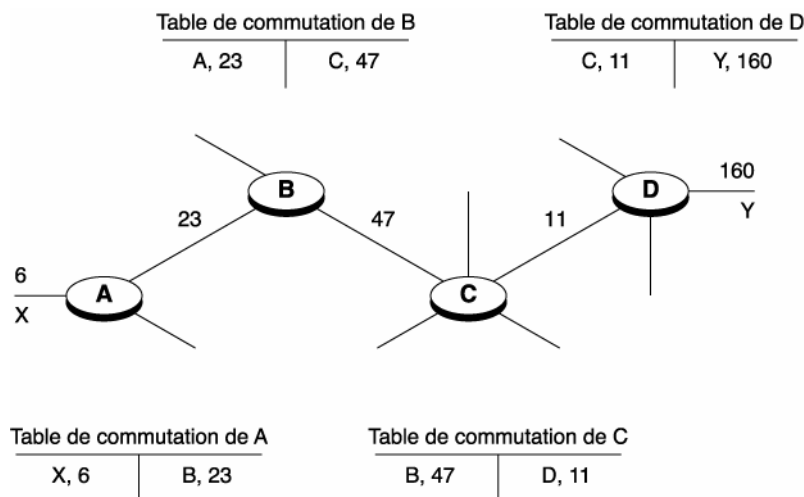


Figure 8. Un circuit virtuel.

La réservation en Soft-State

Les futurs réseaux Internet utiliseront sûrement ce mécanisme de chemin en Soft-State. En effet, à la signalisation ouvrant la route peut être ajoutée une réservation partielle de ressources. Cela ressemble presque à une commutation de circuits, sauf que, au lieu de réserver complètement les ressources nécessaires à la mise en place du circuit, on réserve seulement partiellement ces ressources en espérant que, statistiquement, tout se passe bien. Cette probabilité est évidemment plus grande que si aucune ressource n'était réservée.

Le routage-commutation

On constate depuis quelque temps une tendance à superposer dans un même équipement un commutateur et un routeur. La raison à cela est que certaines applications demandent plutôt un routage, tandis que d'autres réclament une commutation. Par exemple, la navigation dans une base de données Web distribuée au niveau mondial est préférable dans un environnement routé. En revanche, le transfert d'un gros fichier s'adapte mieux à une commutation.

Ces constatations ont incité beaucoup d'industriels à optimiser l'acheminement des paquets en proposant des solutions mixtes. Les routeurs-commutateurs présentent une architecture double, avec une partie routeur et une partie commutateur. L'application choisit si son flot doit transiter *via* une commutation ou un routage. Les routeurs-commutateurs doivent donc gérer pour cela à la fois un protocole de routage et un protocole de commutation, ce qui a l'inconvénient d'augmenter sensiblement le prix de ces équipements.

Les différents types de routeurs-commutateurs

En règle générale, les routeurs-commutateurs utilisent une commutation de type ATM ou Ethernet et un routage de type IP. Cela se traduit par des nœuds complexes, aptes à utiliser plusieurs politiques de transfert de paquets.

Les routeurs-commutateurs sont aussi appelés des LSR (*Label Switch Router*). Chaque grand équipementier propose sa propre solution pour intégrer les deux techniques simultanément, d'où le nombre de noms relativement différents affectés à cette technologie : IP-Switch, Tag-switch, commutateur ARIS, etc.

L'IETF, organisme de normalisation du monde Internet, a pris les choses en main pour essayer de faire converger toutes ces solutions et a donné naissance au protocole MPLS (*MultiProtocol Label Switching*).

Une question soulevée par les réseaux à transfert de paquets consiste à se demander si l'on peut réaliser une commutation de circuits sur une commutation de paquets ? Cette solution est en général difficile, quoique possible. Supposons que l'on soit capable de limiter le temps de traversée d'un réseau à transfert de paquets à une valeur T . Les données provenant du circuit sont encapsulées dans un paquet, qui entre lui-même dans le réseau pourvu de sa date d'émission t . À la sortie, le paquet est conservé jusqu'à la date $t + T$. Les données du paquet sont ensuite remises sur le circuit, qui est ainsi reconstitué à l'identique de ce qu'il était à l'entrée. Toute la difficulté consiste à assurer un délai de traversée borné par T , ce qui est en général impossible sur un réseau classique à transfert de paquets. Des solutions à venir devraient introduire des contrôles et des priorités pour garantir un délai de transport majoré par un temps acceptable pour l'application. Dans la future génération d'Internet, Internet NG, ou *Internet Next Generation*, la commutation de cellules et le transfert de paquets IP incorporeront des techniques de ce style pour réaliser le transport de tous les types d'informations.

Le transfert de trames et de cellules

Historiquement, les réseaux à commutation de circuits ont été les premiers à apparaître : le réseau téléphonique en est un exemple. Le transfert de paquets a pris la succession pour optimiser l'utilisation des lignes de communication dans les environnements informatiques. Récemment, deux nouveaux types de commutation, le transfert de trames et le transfert de cellules, sont apparus. Apparentés à la commutation de paquets, dont ils peuvent être considérés comme des évolutions, ils ont été mis au point pour augmenter les débits sur les lignes et prendre en charge des applications multimédias. Le réseau de transmission comporte des nœuds de commutation ou de routage, qui se chargent de faire progresser la trame ou la cellule vers le destinataire.

La commutation de trames se propose d'étendre la commutation de paquets. La différence entre un paquet et une trame est assez ténue, mais elle est importante. Une trame est un paquet dont on peut reconnaître le début et la fin par différents mécanismes. Un paquet doit être mis dans une trame pour être transporté. En effet, pour envoyer des éléments binaires sur une ligne de communication, il faut que le récepteur soit capable de reconnaître où commence et où finit le bloc transmis. La commutation de trames consiste à commuter des trames dans le nœud, ce qui a l'avantage de pouvoir les transmettre directement sur la ligne, juste après les avoir aiguillées vers la bonne porte de sortie. Dans une commutation de paquets, au contraire, le paquet doit d'abord être récupéré en décapsulant la trame qui a été transportée sur la ligne, puis la référence doit être examinée afin de déterminer, en se reportant à la table de commutation, la ligne de sortie adéquate. Il faut ensuite encapsuler de nouveau dans une trame le paquet pour émettre ce dernier vers le nœud suivant.

La commutation de trames présente l'avantage de ne remonter qu'au niveau trame au lieu du niveau paquet. Pour cette raison, les commutateurs de trames sont plus simples, plus performants et moins chers à l'achat que les commutateurs de paquets.

Plusieurs catégories de commutation de trames ont été développées en fonction du protocole de niveau trame choisi. Les deux principales concernent le relais de trames et la commutation Ethernet. Dans le relais de trames, on a voulu simplifier au maximum la commutation de paquets. Dans la commutation Ethernet, on utilise la trame Ethernet comme bloc de transfert. De ce fait, l'adressage provient des normes de l'environnement Ethernet.

Cette commutation de trames peut être considérée comme une technique intermédiaire en attendant soit l'arrivée des techniques à commutation de cellules, soit des extensions des techniques utilisées dans le réseau Internet.

La commutation de cellules est une commutation de trames très particulière, propre aux réseaux ATM, dans lesquels toutes les trames possèdent une longueur fixe de 53 octets. Quelle que soit la taille des données à transporter, la cellule occupe toujours 53 octets. Si les données forment un bloc de plus de 53 octets, un découpage est effectué, et la dernière cellule n'est pas complètement remplie. Plus précisément, la cellule comporte 48 octets de données et 5 octets de supervision. Le mot commutation indique la présence d'une référence dans l'en-tête de 5 octets. Cette référence tient sur 24 ou 28 octets, suivant l'endroit où l'on émet : d'une machine utilisateur vers un nœud de commutation ou d'un nœud de commutation vers un autre nœud de commutation.

Les techniques de transfert hybrides

Les différentes méthodes de transfert présentées dans ce support de cours peuvent se superposer pour former des techniques de transfert hybrides. En général, les superpositions concernent l'encapsulation d'un niveau paquet dans un niveau trame. Le protocole de niveau trame s'appuie essentiellement sur une commutation et celui de niveau paquet sur un routage. Cette solution permet de définir des nœuds de type routeur-commutateur. Si l'on remonte au niveau paquet, un routage a lieu ; si l'on ne remonte qu'au niveau trame, une commutation est effectuée.

La figure 9 illustre une architecture hybride de routeurs-commutateurs. Les données remontent jusqu'au niveau paquet pour être routées ou bien au niveau trame pour être commutées. Le choix de remonter au niveau trame ou au niveau paquet dépend en général de la longueur du flot de paquets d'un même utilisateur. Lorsque le flot est court, comme dans une navigation sur le World-Wide Web, chaque paquet détermine par lui-même son chemin. Il faut aller rechercher l'adresse dans le paquet pour pouvoir le router. En revanche, lorsque le flot est long, il est intéressant de mettre en place des références dans les nœuds pour commuter les paquets du flot. Le premier paquet est en général routé, et il pose des références pour les trames suivantes, qui sont alors commutées.

On peut très bien envisager un routage de niveau trame et une commutation de niveau paquet. Dans ce cas, certaines informations remontent au niveau paquet pour être commutées, alors que d'autres peuvent être routées directement au niveau trame. Ce cas de figure peut se trouver dans les réseaux qui ont une double adresse, par exemple, une adresse complète de niveau trame et une référence pour le niveau paquet. Lorsqu'une trame arrive dans un nœud, celui-ci récupère l'adresse de niveau trame. Si l'adresse de ce niveau est connue, le routage peut s'effectuer. En revanche, si l'adresse de niveau 2 ne correspond pas à une adresse connue du nœud, il est possible

de décapsuler la trame pour récupérer le paquet et d'examiner la référence de ce niveau déterminant la direction à prendre. Il en est ainsi des réseaux locaux Ethernet, qui, outre l'adresse de niveau trame, portent des paquets X.25 munis d'une référence de niveau paquet.

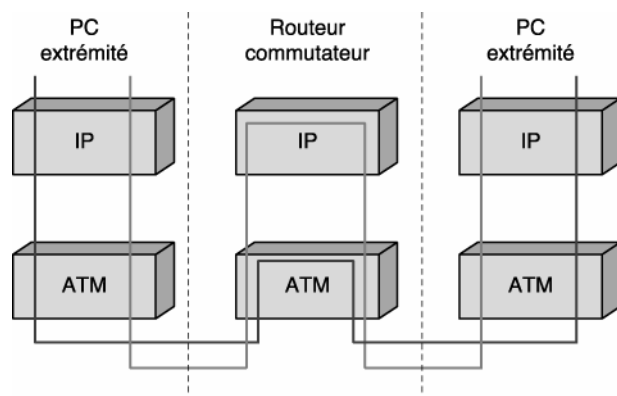


Figure 9. Une architecture hybride de routeurs-commutateurs.

Les réseaux télécoms : ATM et MPLS

Les réseaux à intégration de services

Du début des années 70 jusqu'à la fin des années 90, les services de transmission de données se développent sur le principe des réseaux spécialisés : à un usage correspond un réseau spécifique. L'utilisateur qui a besoin de communiquer avec chacun de ces réseaux est obligé d'avoir autant de raccordements que de réseaux ou d'applications à atteindre.

Cette multitude de raccordements différents et indépendants n'est optimale ni du point de vue de l'utilisateur, ni du point de vue de l'exploitant de télécommunications. De cette constatation naît le concept d'*intégration de services*.

Le RNIS bande étroite correspond à une évolution du réseau téléphonique. Au début des années 80, le réseau téléphonique achève sa numérisation : toutes les conversations téléphoniques sont numérisées à l'aide d'un codec en entrée du réseau, sous la forme de flots à 64 Kbit/s. Cette numérisation permet d'utiliser les lignes numériques à 64 Kbit/s à d'autres fins que pour le simple service téléphonique.

Le RNIS bande étroite n'est pas un réseau supplémentaire entrant en concurrence avec les réseaux existants, comme le téléphone traditionnel, les réseaux X.25 ou les liaisons spécialisées. C'est la réutilisation du réseau existant, devenu numérique, pour introduire des services de type informatique. La principale évolution concerne d'ailleurs la partie du réseau qui dessert l'utilisateur, le réseau d'accès, devenant également numérique pour permettre la continuité numérique d'un utilisateur émetteur vers un utilisateur récepteur.

Le RNIS bande étroite correspond avant tout à une interface d'accès universel, que l'on appelle l'interface S, entre l'utilisateur et le commutateur de l'opérateur.

À partir du commutateur de l'opérateur, les informations se dirigent vers un sous-réseau correspondant au type de flux à transmettre. Le réseau de signalisation correspond au réseau physique qui transporte les commandes du réseau. Cette architecture est illustrée à la figure 10.

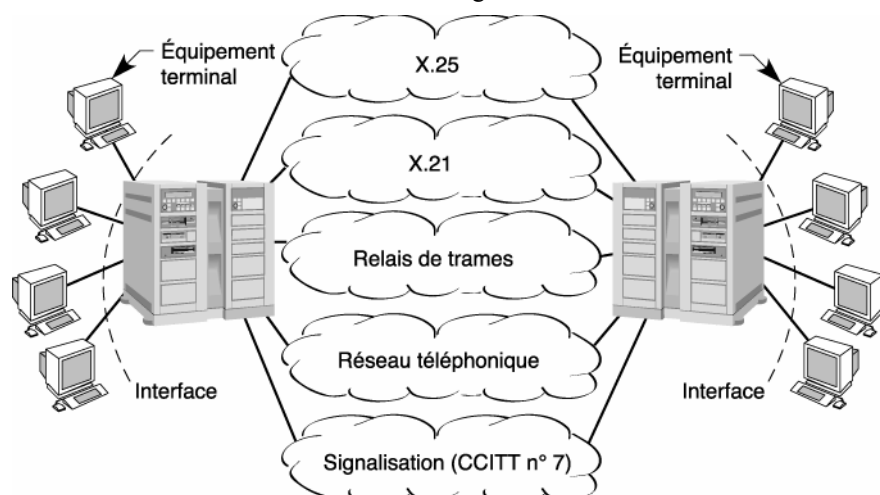


Figure 10. Architecture du RNIS bande étroite.

Pour échanger des messages avec le commutateur, un protocole de liaison est mis en œuvre, dont le rôle est d'assurer la transmission de messages sans erreur entre le réseau et le client. Ce protocole, le LAP-D, fait partie de la famille HDLC.

Le RNIS large bande part d'un principe totalement différent de celui du RNIS bande étroite. Pour gagner en coût, il consiste à multiplexer l'ensemble des informations voix, données et vidéo sur un seul et même réseau. La difficulté provient de la technique de transfert à adopter, qui doit à la fois tenir compte du temps réel, d'importants débits et plus généralement d'une qualité de service dépendant de l'application. De plus, cette technique de transfert doit permettre de très hauts débits, d'où le nom de large bande.

Les avantages d'un réseau large bande ayant une seule technique de transfert sont multiples :

- Investissements et coûts d'exploitation et de maintenance inférieurs à ceux du RNIS bande étroite. En effet, bien que le RNIS bande étroite offre une interface unique d'accès, l'interface S, les services doivent tenir compte des sous-interfaces avec les canaux B et D : une sous-interface en *mode circuit* et une sous-interface en *mode paquet*, ces deux sous-interfaces restant disjointes.
- Utilisation optimale des ressources disponibles pour la gestion dynamique du partage des ressources en fonction des besoins.
- Forte capacité d'adaptation aux nouveaux besoins des utilisateurs.

Le RNIS large bande offre, d'une part, des services qui lui sont propres et doit, d'autre part, assurer la continuité des services offerts par le RNIS bande étroite. Le RNIS large bande est ainsi capable de supporter tous les types de services et d'applications, avec des contraintes variées de débit et de qualité de service.

L'évolution vers le RNIS large bande a été décidée pour répondre à la demande croissante de services haut débit. Son déploiement a pu commencer grâce à l'émergence de technologies telles que la transmission par fibre optique, qui permet d'atteindre plusieurs gigabits par seconde, les équipements de commutation rapide, pour suivre le rythme de la fibre optique, et l'arrivée de techniques rapides sur le réseau d'accès.

Deux grandes méthodes se sont affrontées pendant les années 80 pour devenir la norme de transport des réseaux RNIS large bande : le transfert STM (*Synchronous Transfer Mode*) et le transfert ATM (*Asynchronous Transfer Mode*). Le mode de transfert synchrone temporel, ou STM, est fondé sur le multiplexage temporel et sur la commutation de circuits.

STM (*Synchronous Transfer Mode*)

Le mode de transfert STM a longtemps été considéré comme la solution adéquate pour les réseaux RNIS large bande. Ce mode de transfert découpe le temps en trame, chaque trame se découplant à son tour en tranches, elles-mêmes allouées aux utilisateurs. Une même tranche de toutes les trames qui passent peut être assignée à un appel, ce dernier étant identifié par la position de la tranche. Plus le débit demandé par un utilisateur est élevé, et plus le besoin de se voir allouer plusieurs tranches est important. La technique STM jouit d'une excellente réputation pour les services à débit constant. En revanche, la bande passante pour les services à débit variable est gaspillée. Cette solution est illustrée à la figure 11.

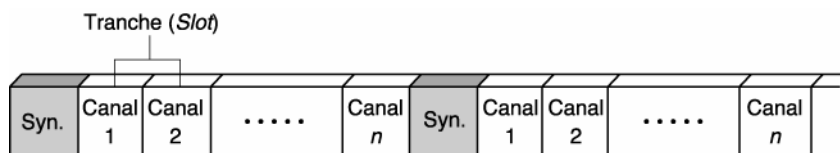


Figure 11. La méthode de transfert STM.

La technique ATM (*Asynchronous Transfer Mode*), fondée sur le *multiplexage temporel asynchrone* et sur la commutation de paquets, permet une meilleure utilisation des ressources lors du transport de données asynchrones. L'information est transportée par des cellules dont la taille est fixe. Ces cellules sont transmises à l'intérieur d'un circuit virtuel.

La différence fondamentale entre les modes ATM et STM réside dans la manière d'allouer les tranches de temps aux machines terminales. L'ATM utilise une technique asynchrone. Les services à débit variable peuvent ainsi être pris en charge par l'ATM sans gaspillage de bande passante. Les critères d'efficacité et de flexibilité ont été déterminants lors du choix de l'ATM comme mode de transfert pour les réseaux télécoms du futur. Pour améliorer les performances des applications à débit variable, les techniques de commutation ont été modifiées dans les directions suivantes :

- pas de contrôle d'erreur sur le champ d'information ;
- pas de contrôle de flux au niveau de la cellule ;
- intégrité du séquençement des cellules ;
- bloc d'information de longueur fixe (cellule) ;
- routage par un circuit virtuel en mode avec connexion.

L'utilisation d'une longueur fixe pour les cellules et les fonctions simplifiées permet d'effectuer une commutation de cellules à très haut débit.

Les dernières évolutions du monde IP introduisent tout un ensemble de protocoles destinés à favoriser la qualité de service. C'est là un changement de cap considérable par rapport à la première génération, qui se contentait d'un service best effort. La deuxième génération des réseaux IP peut s'imposer dans ce sens et œuvrer en faveur de son choix, à la place de l'ATM, dans les réseaux large bande. Les *gigarouteurs* et les technologies rapides de transport des paquets IP, qu'il s'agisse de *POS (Packet Over Sonet)* ou de *WDM (Wavelength Division Multiplexing)*, forment l'ossature de ces futurs réseaux large bande.

SONET (*Synchronous Optical NETwork*)

Proposée par les Américains, la norme SONET n'a d'abord concerné que l'interconnexion des réseaux téléphoniques des grands opérateurs (PTT, grands opérateurs américains, etc.). La difficulté de cette norme résidait dans l'interconnexion de lignes de communication qui ne présentaient pas du tout les mêmes standards en Europe, au Japon et aux Amériques.

La hiérarchie des débits étant également différente sur les trois continents, il a fallu trouver un compromis pour le niveau de base. C'est finalement un débit à 51,84 Mbit/s qui l'a emporté pour former le premier niveau de SONET, appelé STS-1 (*Synchronous Transport Signal, level 1*), les niveaux situés au-dessus du niveau 1 (STS-N) étant des multiples du niveau de base.

Sonet décrit la composition d'une trame synchrone émise toutes les 125 μ s. La longueur de cette trame dépend de la vitesse de l'interface. Ces diverses valeurs sont présentées dans le tableau ci-dessous et classées suivant la rapidité du support optique OC (*Optical Carrier*).

OC-1	51,84 Mbit/s
OC-3	155,52 Mbit/s
OC-9	466,56 Mbit/s
OC-12	622,08 Mbit/s
OC-24	1 244,16 Mbit/s
OC-48	2 488,32 Mbit/s
OC-96	4 976,64 Mbit/s
OC-192	9 953,28 Mbit/s

La trame SONET comprend, dans les trois premiers octets de chaque rangée, des informations de synchronisation et de supervision. Les cellules ATM ou les paquets IP sont émis à l'intérieur de la trame SONET. L'instant de début de l'envoi d'une cellule ou d'un paquet ne correspond pas forcément au début de la trame SONET et peut se situer n'importe où dans la trame. Des bits de supervision précèdent ce début de trame, de sorte que l'on ne perde pas de temps pour l'émission d'une cellule ou d'un paquet.

Les réseaux ATM

La technologie de transfert ATM a été choisie en 1988 pour réaliser le réseau de transport du RNIS large bande.

L'ATM désigne un mode de transfert asynchrone, utilisant des trames spécifiques et faisant appel à la technique de *multiplexage asynchrone par répartition dans le temps*. Le flux d'information multiplexé est structuré en petits blocs, ou cellules. Ces dernières sont assignées à la demande, selon l'activité de la source et les ressources disponibles.

La commutation de cellules est une commutation de trames assez particulière, dans laquelle toutes les trames possèdent une longueur à la fois constante et très petite. La cellule est formée d'exactly 53 octets, comprenant 5 octets d'en-tête et 48 octets de données. Sur les 48 octets provenant de la couche supérieure, jusqu'à 4 octets peuvent concerner la supervision (*voir figure 12*). Les 5 octets de supervision sont détaillés à la figure 13.

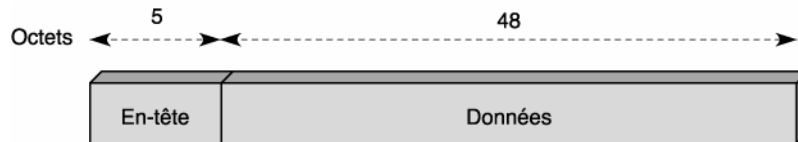
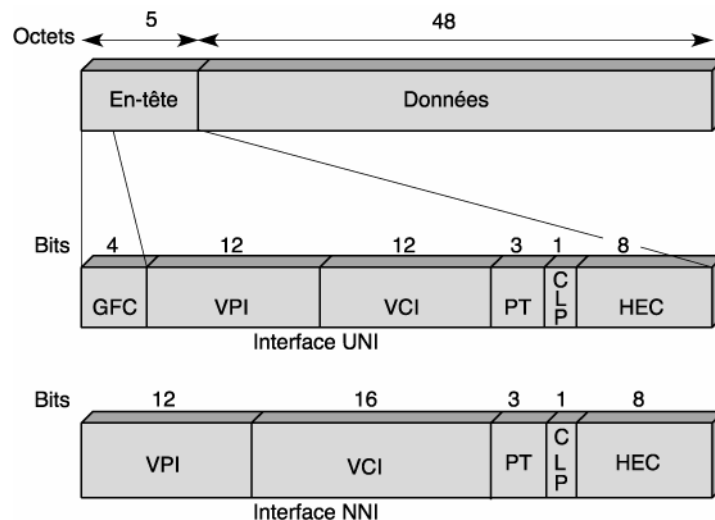


Figure 12. La cellule ATM.

La longueur de la zone de données de 48 octets est le résultat d'un accord entre les Européens, qui souhaitaient 32 octets, et les Américains, qui en désiraient 64.

La très faible longueur de la cellule est facilement explicable. Prenons pour cela l'exemple de la transmission de la parole téléphonique, qui demande une liaison de 64 Kbit/s. C'est une application isochrone, qui possède les deux contraintes suivantes :

- Une synchronisation très forte des données : un octet part de l'émetteur toutes les 125 μ s, et les octets doivent être remis au codeur-décodeur de l'autre extrémité toutes les 125 μ s.
- Un délai de propagation qui doit rester inférieur à 28 ms si l'on veut éviter tous les problèmes liés à la transmission de signaux sur une longue distance (suppression des échos, adaptation, etc.).



CLP : Cell Loss Priority
GFC : Generic Flow Control
HEC : Header Error Control
NNI : Network Node Interface
PT : Payload Type
UNI : User Network Interface
VCI : Virtual Channel Identifier
VPI : Virtual Path Identifier

Figure 13. Les octets d'en-tête de la cellule ATM.

Le temps de transit des octets pour la parole sortant d'un combiné téléphonique se décompose de la façon suivante :

- Un temps de remplissage de la cellule par les octets qui sortent du combiné téléphonique toutes les 125 μ s. Il faut donc exactement 6 ms pour remplir la cellule de 48 octets de longueur.
- Un temps de transport de la cellule dans le réseau.
- Encore 6 ms pour vider la cellule à l'extrémité, puisque l'on remet au combiné téléphonique un octet toutes les 125 μ s.

Comme le temps total ne doit pas dépasser 28 ms, on voit que, si l'on retranche le temps aux extrémités, il n'y a plus que 16 ms de délai de propagation dans le réseau lui-même. En supposant que le signal soit transmis sur un câble électrique à la vitesse de 200 000 km/s, la distance maximale que peut parcourir un tel signal est de 3 200 km. Cette distance peut bien évidemment être augmentée si l'on ajoute des équipements adaptés pour la suppression des échos, l'adaptation, etc.

Comme le territoire nord-américain est particulièrement étendu, il a fallu, aux États-Unis, mettre en place tous ces types de matériels dès les premières générations de réseaux téléphoniques. Les Américains ont préconisé une meilleure utilisation de la bande passante du RNIS large bande par l'allongement de la zone de données des cellules par rapport à la partie supervision. En Europe, pour éviter d'avoir à adapter les réseaux terrestres, on aurait préféré une taille de cellule plus petite, de 32 voire de 16 octets, de façon à gagner du temps aux extrémités. Ces contraintes sont illustrées à la figure 14.

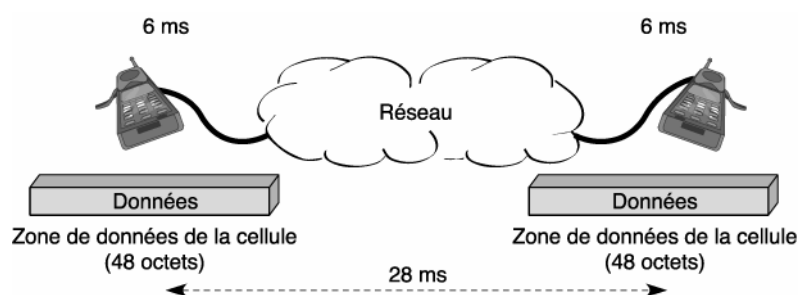


Figure 14. Les contraintes de propagation du signal dans un réseau ATM.

Caractéristiques des réseaux ATM

La première caractéristique importante des réseaux ATM est l'utilisation du mode avec connexion pour la transmission des cellules. Une cellule n'est transmise que lorsqu'un circuit virtuel est ouvert, ce circuit virtuel étant marqué à l'intérieur du réseau par des références laissées dans chaque nœud traversé.

La structure de la zone de supervision est illustrée à la figure 13. Elle comporte tout d'abord deux interfaces différentes, suivant que la cellule provient de l'extérieur ou passe d'un nœud de commutation à un autre à l'intérieur du réseau :

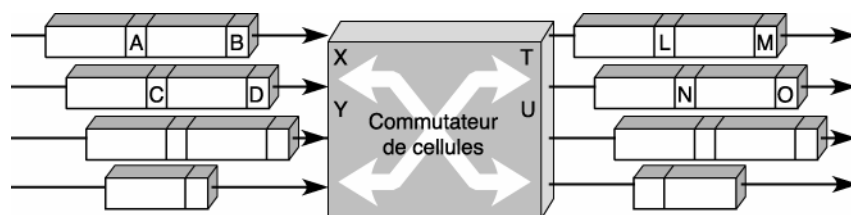
- l'interface NNI (*Network-Node Interface*), se situant entre deux nœuds du réseau ;
- l'interface UNI (*User Network Interface*), permettant l'entrée ou la sortie du réseau.

La première partie de la zone de supervision comporte deux valeurs : le numéro *VCI* (*Virtual Channel Identifier*, ou identificateur de *voie virtuelle*) et le numéro *VPI* (*Virtual Path Identifier*, ou identificateur de conduit virtuel). Ces numéros identifient une connexion entre deux extrémités du réseau. L'adjonction de ces deux numéros correspond à la référence du circuit virtuel, à l'instar de ce qui se passe dans la norme X.25 de niveau 3. En d'autres termes, la référence identifiant le circuit virtuel comporte deux parties : le numéro de conduit virtuel (*Virtual Path*) et le numéro de voie virtuelle (*Virtual Channel*).

L'ATM travaille en *mode commuté* et utilise un mode avec connexion, solutions prévisibles dans le cadre d'un environnement télécoms. Avant toute émission de cellules, un circuit virtuel de bout en bout doit être mis en place. Plus spécifiquement, la norme ATM précise qu'une structure de conduit virtuel doit être mise en place et identifiée par l'association d'une voie virtuelle et d'un conduit virtuel. On retrouve cette technique dans les réseaux X.25 possédant un circuit virtuel matérialisé dans les nœuds intermédiaires.

Le routage de la cellule de supervision mettant en place le circuit virtuel peut s'effectuer grâce à des tables de routage. Ces tables déterminent vers quel nœud doit être envoyée la cellule de supervision qui renferme l'adresse du destinataire final. Une autre solution consiste à ouvrir des circuits virtuels au préalable — si possible plusieurs — entre chaque point d'accès et chaque point de sortie. Cette cellule de supervision définit, pour chaque nœud traversé, l'association entre la référence du port d'entrée et la référence du port de sortie. Ces associations sont regroupées dans la table de commutation.

La figure 15 illustre l'association effectuée entre le chemin d'entrée dans un nœud de commutation et le chemin de sortie de ce même commutateur. Par exemple, si une cellule se présente à la porte d'entrée X avec la référence A, elle est transmise à la sortie T avec la référence L. La deuxième ligne du tableau de commutation constitue un autre exemple : une cellule qui entre sur la ligne X avec la référence B est envoyée vers la sortie U, accompagnée de la référence N de sortie.



Ligne d'entrée	Référence d'entrée	Ligne de sortie	Référence de sortie
X	A	T	L
X	B	U	N
Y	C	T	M
Y	D	T	O

Figure 15. Commutation des cellules dans un nœud de commutation.

Les références permettant de commuter les cellules sont appelées, comme nous l'avons vu, VCI et VPI pour ce qui concerne la voie et le conduit. Dans un commutateur ATM, on commute une cellule en utilisant les deux références. Dans un *brasseur*, on ne se sert que d'une seule référence, celle du conduit. Par exemple, on peut commuter un ensemble de voies virtuelles en une seule fois en ne se préoccupant que du conduit. Dans ce cas, on a un *brasseur de conduits* (ou *Cross-Connect*), et l'on ne redescend pas au niveau de la voie virtuelle.

La figure 16 représente un circuit virtuel avec un commutateur ATM et un brasseur.

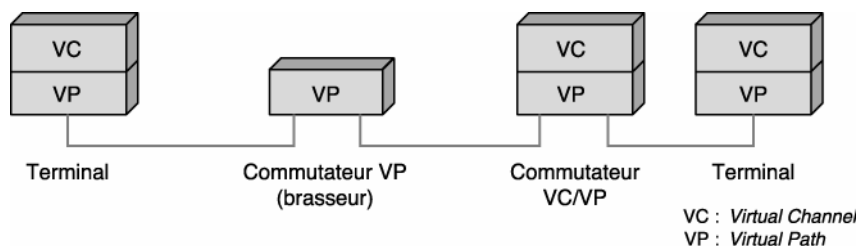


Figure 16. Circuit virtuel avec un brasseur et un commutateur ATM.

Dans un brasseur de conduits, on commute simultanément toutes les voies virtuelles à l'intérieur du conduit. On a donc intérêt à regrouper les voies virtuelles qui vont vers la même destination pour les intégrer dans un même conduit. Cela simplifie les problèmes de commutation à l'intérieur du réseau. La figure 17 illustre de façon assez symbolique un conduit partagé par un ensemble de voies. Le long du conduit, des brasseurs VP peuvent se succéder.



Figure 17. Multiplexage de VC dans un VP.

Les bits GFC (*Generic Flow Control*) servent au contrôle d'accès et au contrôle de flux sur la partie terminale, entre l'utilisateur et le réseau. Lorsque plusieurs utilisateurs veulent entrer dans le réseau ATM par un même point d'entrée, il faut en effet ordonner leurs demandes.

Le champ de contrôle comporte ensuite 3 bits PT (*Payload Type*), qui définissent le type de l'information transportée dans la cellule. On peut y trouver divers types d'informations pour la gestion et le contrôle du réseau. Les huit possibilités pour ce champ PT sont les suivantes :

- 000. Cellule de données utilisateur, pas de congestion : indication d'un niveau utilisateur du réseau ATM vers un autre utilisateur du réseau ATM = 0.
- 001. Cellule de données utilisateur, pas de congestion : indication d'un niveau utilisateur du réseau ATM vers un autre utilisateur du réseau ATM = 1.

- 010. Cellule de données utilisateur, congestion : indication d'un niveau utilisateur du réseau ATM vers un autre utilisateur du réseau ATM = 0.
- 011. Cellule de données utilisateur, congestion : indication d'un niveau utilisateur du réseau ATM vers un autre utilisateur du réseau ATM = 1.
- 100. Cellule de gestion pour le flux *OAM F5* de segment. Ces cellules correspondent à de l'information de gestion envoyée d'un nœud vers le nœud suivant.
- 101. Cellule de gestion pour le flux *OAM F5* de bout en bout. Ces cellules sont dirigées vers le bout du circuit virtuel et non à un nœud intermédiaire.
- 110. Cellule pour la gestion des ressources.
- 111. Réservee à des fonctions futures.

Vient ensuite le bit *CLP (Cell Loss Priority)*, qui indique si la cellule peut être perdue ($CLP = 1$) ou, au contraire, si elle est essentielle ($CLP = 0$). Ce bit aide au contrôle de flux. En effet, avant d'émettre une cellule dans le réseau, il convient de respecter un taux d'entrée, négocié au moment de l'ouverture du circuit virtuel. Il est toujours possible de faire entrer des cellules en surnombre, mais il faut les munir d'un indicateur permettant de les repérer par rapport aux données de base. Ces données en surnombre peuvent être perdues pour permettre aux informations entrées dans le cadre du contrôle de flux de passer sans problème.

La dernière partie de la zone de contrôle, le *HEC (Header Error Control)*, concerne la protection de l'en-tête. Ce champ permet de détecter et de corriger une erreur en mode standard. Lorsqu'un en-tête en erreur est détecté et qu'une correction n'est pas possible, la cellule est détruite.

En résumé, les réseaux ATM n'ont que peu d'originalité. On y retrouve de nombreux algorithmes déjà utilisés dans les réseaux classiques à commutation de paquets. Cependant, la hiérarchie des procédures et les protocoles utilisés sont assez différents de ceux de la première génération de réseaux. Cette architecture est illustrée à la figure 18.

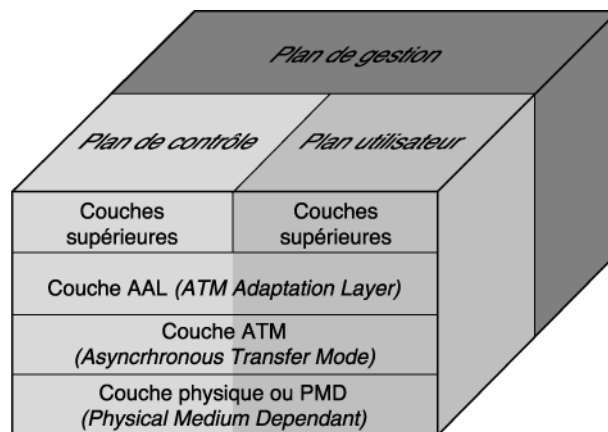


Figure 18. Architecture UIT-T de l'environnement ATM.

La couche la plus basse concerne les protocoles de niveau physique dépendant du médium, ou *PMD (Physical Medium Dependent)*. Cette couche est elle-même divisée en deux sous-couches :

- La couche *TC (Transmission Convergence)*, chargée du découplage du taux de transmission des cellules, de la génération et de la vérification de la zone de détection d'erreur de l'en-tête (le *HEC*), de la délimitation des cellules, de l'adaptation de la vitesse de transmission et, enfin, de la génération et de la récupération des cellules sur le support physique.
- La couche *PM (Physical Medium)*, chargée de la transmission sur le support physique et des problèmes d'horloge.

La deuxième couche gère le transport de bout en bout de la cellule ATM (*Asynchronous Transfer Mode*).

Enfin, le niveau *AAL (ATM Adaptation Layer)*, niveau d'adaptation à l'ATM, se charge de l'interface avec les couches supérieures. Cet étage est lui-même subdivisé en deux niveaux, l'un prenant en compte les problèmes liés directement à l'interfonctionnement avec la couche supérieure, et l'autre ceux concernant la fragmentation et le réassemblage des messages en cellules.

Dans cette couche *AAL*, quatre classes de services (A, B, C et D) ont été définies. À ces classes correspondent quatre classes de protocoles, numérotées de 1 à 4. En réalité, cette subdivision a été modifiée en 1993 par le regroupement des classes 3 et 4 et par l'ajout d'une nouvelle classe de protocole, la classe 5, qui définit un transport de données simplifié.

La première classe de service correspond à une émulation de circuit, la deuxième au transport de la vidéo, la troisième à un transfert de données en mode avec connexion et la dernière à un transfert de données en mode sans connexion.

Les classes de qualité de service de l'ATM

La qualité de service constitue un point particulièrement sensible de l'environnement ATM, puisque c'est l'élément qui permet de distinguer l'ATM des autres types de protocoles. Pour arriver à une qualité de service, il faut allouer des ressources. La solution choisie concerne l'introduction de classes de priorité.

Les cinq classes de qualité de service suivantes, dites classes de services, ont été déterminées :

- CBR (*Constant Bit Rate*), qui correspond à un circuit virtuel avec une bande passante fixe. Les services de cette classe incluent la voix ou la vidéo temps réel.
- VBR (*Variable Bit Rate*), qui correspond à un circuit virtuel pour des trafics variables dans le temps et plus spécifiquement à des *services rigides* avec de fortes variations de débit. Les services de cette classe incluent les services d'interconnexion de réseaux locaux ou le transactionnel. Il existe une classe VBR RT (*Real Time*), qui doit prendre en compte les problèmes de temps réel.
- ABR (*Available Bit Rate*), qui permet d'utiliser la bande passante restante pour des applications qui ont des débits variables et sont sensibles aux pertes. Un débit minimal doit être garanti pour que les applications puissent passer en un temps acceptable. Le temps de réponse n'est pas garanti dans ce service. Cette classe correspond aux *services élastiques*.
- GFR (*Guaranteed Frame Rate*), qui correspond à une amélioration du service ABR en ce qui concerne la complexité d'implantation de ce dernier sur un réseau. Le service GFR se fonde sur l'utilisation d'un trafic minimal. Si un client respecte son service minimal, le taux de perte de ses cellules doit être très faible. Le trafic dépassant le trafic minimal est marqué, et, si le réseau est en état de congestion, ce sont ces cellules qui sont perdues en premier. Le contrôle des paquets s'effectue sur la base de la trame : si une cellule de la trame est perdue, le mécanisme de contrôle essaie d'éliminer toutes les cellules appartenant à la même trame.
- UBR (*Unspecified Bit Rate*), qui correspond au best effort. Il n'existe aucune garantie ni sur les pertes, ni sur le temps de transport. Le service UBR, qui n'a pas de garantie de qualité de service, n'est d'ailleurs pas accepté par les opérateurs télécoms, qui ne peuvent se permettre de proposer un service sans aucune qualité de service.

La figure 19 illustre l'allocation des classes de services. Dans un premier temps, les classes CBR et VBR sont allouées avec des ressources permettant une garantie totale de la qualité de service des données qui transitent dans les circuits virtuels concernés. Pour cela, on peut allouer les ressources sans restriction, puisque tout ce qui n'est pas utilisé peut être récupéré dans le service ABR.

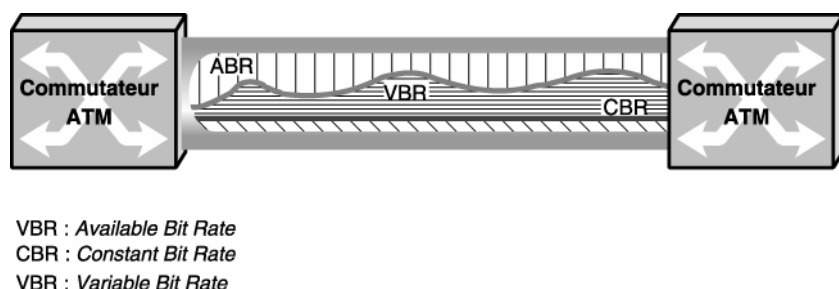


Figure 19. La réservation de classes de services entre deux nœuds.

La répartition des informations par classes s'effectue de la façon suivante : on affecte tout d'abord la bande passante au trafic CBR, et l'opérateur ne fait qu'additionner les bandes passantes demandées par les clients. On peut supposer que la bande passante ainsi réservée soit bien utilisée, sinon la place restant libre est réaffectée au trafic ABR. Une fois cette affectation réalisée, l'opérateur retient une bande passante pour y faire transiter le trafic VBR. Cette réservation correspond, dans la figure 19, à la somme des zones notées VBR et ABR VBR. Cette réservation est à la charge de l'opérateur, qui peut l'effectuer de différentes façons, par exemple en réservant la somme des débits crêtes ou, après calcul, par une surallocation, sachant qu'il existe peu de chance que tous les clients aient besoin du débit crête en même temps. Tout cela est du ressort de l'opérateur. Le client, lui, doit pouvoir considérer qu'il dispose quasiment du débit crête pour que les garanties de ce service puissent être réalisées à coup sûr.

Dans la réalité, l'utilisation de cette bande passante réservée est largement inférieure à la réservation faite par l'opérateur. La zone utilisée est, sur la figure, la zone non hachurée, notée VBR. La partie hachurée est la partie réservée mais non utilisée par le trafic VBR et qui est donc réaffectée au trafic ABR.

On comprend mieux maintenant pourquoi le contrôle de flux est indispensable au trafic ABR. En effet, le but de ce trafic est de remplir, au plus près possible des 100 p. 100, le tuyau global. Comme, à chaque instant, le volume de trafic avec garantie varie, il faut transmettre plus ou moins de trafic ABR. On doit donc être capable de dire à l'émetteur, à tout instant, quelle quantité de trafic ABR il faut laisser entrer pour optimiser l'utilisation des tuyaux de communication dans le réseau.

Comme le trafic ABR n'offre pas de garantie sur le temps de réponse, on peut se dire que si le contrôle de flux est parfait, on est capable de remplir complètement les voies de communication du réseau.

MPLS (*MultiProtocol Label Switching*)

L'environnement IP est devenu le standard de raccordement à un réseau pour tous les systèmes distribués provenant de l'informatique. De son côté, la technique de transfert ATM est la solution préférée des opérateurs pour relier deux routeurs entre eux avec une qualité de service. Il était donc plus que tentant d'empiler les deux environnements pour permettre l'utilisation à la fois de l'interface standard et de la puissance de l'ATM. Cette opération a donné naissance aux architectures dites IP sur ATM. La difficulté se situe au niveau de l'interface entre IP et ATM, avec le découpage des paquets IP en cellules, et lors de l'indication dans la cellule d'une référence correspondant à l'adresse IP du destinataire.

La technique regroupant ces possibilités a été mise au point par l'IETF (*Internet Engineering Task Force*), l'organisme de normalisation d'Internet pour l'ensemble des architectures et des protocoles de haut niveau (IP, IPX, AppleTalk, etc.), sous le nom de MPLS (*MultiProtocol Label Switching*). MPLS utilise les techniques de commutation et plus précisément de commutation de références, ou label-switching. Les techniques de commutation utilisées dans MPLS sont essentiellement ATM et Ethernet, mais elles pourraient être aussi de type relais de trames avec LAP-F. MPLS fait appel à un circuit virtuel permettant d'acheminer les paquets qui sont commutés dans les nœuds.

Des extensions de MPLS sont apportées par GMPLS (*Generalized MPLS*), qui introduit de nouveaux paradigmes de commutation sur les interfaces hertziennes et fibre optique. L'implantation de MPLS concerne uniquement le protocole IP, les autres protocoles ayant quasiment disparu au profit d'IP.

Les nœuds de transfert spécifiques utilisés dans MPLS sont appelés LSR (*Label Switched Router*). Ces LSR se comportent comme des commutateurs pour les flots de données utilisateur et comme des routeurs pour gérer la signalisation. Pour acheminer les trames utilisateur, on utilise des références, ou *labels*. À une référence d'entrée correspond une référence de sortie. La succession des références définit la route suivie par l'ensemble des trames contenant les paquets du flot IP. Toute trame utilisée en commutation, ou label-switching, peut être utilisée dans un réseau MPLS. La référence est placée dans un champ spécifique de la trame ou dans un champ ajouté dans ce but.

Dans cette architecture, la route est déterminée par le flot IP. Les LSR remplacent les routeurs en travaillant soit en mode routeur, pour tracer le chemin avec le paquet de signalisation, soit en mode commutation pour toutes les trames qui suivent le chemin tracé. Le paquet de signalisation est routé normalement, comme dans un réseau Internet. La route est déterminée par un algorithme de routage d'Internet.

Prenons l'exemple d'un ensemble de réseaux ATM interconnectés par des LSR pour expliquer l'ouverture du circuit virtuel, que l'on appelle un LSP (*Label Switched Path*). Une fois déterminé le premier routeur à traverser, le paquet IP de signalisation est subdivisé en cellules ATM pour traverser le premier sous-réseau, appelé LIS (*Logical Internet Subnetwork*).

Le paquet IP est recomposé au premier LSR, lequel décide de la route à suivre, toujours à l'aide d'un algorithme de routage classique d'Internet. En même temps, une ligne de la table de commutation est ajoutée pour permettre la commutation des cellules du même flot. Après avoir franchi le premier nœud, le paquet IP de signalisation est de nouveau fragmenté en cellules ATM, lesquelles sont émises vers le nœud suivant puis regroupées, et ainsi de suite.

Tous les paquets IP suivants appartenant au même flot sont subdivisés en cellules ATM par l'émetteur et commutés sur le chemin tracé. Ce dernier devient un circuit virtuel de bout en bout (LSP). Cette technique de routage-commutation est illustrée à la figure 20.

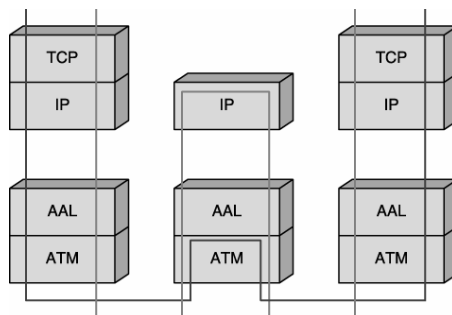


Figure 16-20. Architecture d'un réseau MPLS utilisant des sous-réseaux ATM.

Le chemin peut traverser des réseaux divers, aussi bien ATM qu'Ethernet ou relais de trames. La référence se trouve dans la zone VPI/VCI de la cellule ATM, dans la zone DLCI de la trame LAP-F d'un réseau relais de trames ou dans une zone ajoutée à la trame Ethernet, la zone de référence de dérivation, ou « shim label ».

Un déploiement au-dessus de réseaux Ethernet commutés est illustré à la figure 16-21.

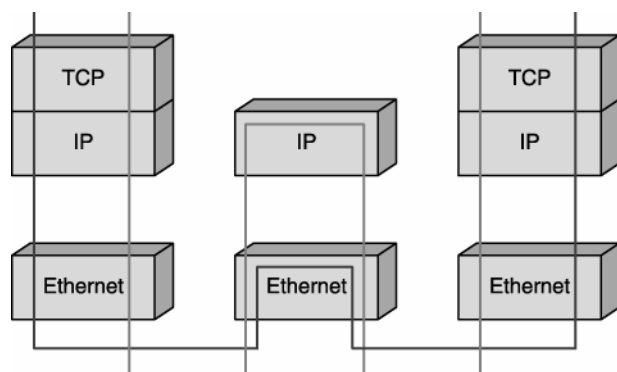


Figure 21. Architecture d'un environnement IP sur Ethernet.

La signalisation MPLS doit être introduite dans cette architecture pour ouvrir le chemin que suivront les trames Ethernet. Cette signalisation provient toujours de l'émission d'un paquet IP de supervision qui est routé dans les LSR. On peut ainsi réaliser des réseaux extrêmement complexes avec des segments partagés sur les parties locales, des liaisons commutées sur les longues distances ou entre les commutateurs Ethernet et des passages par des routeurs lorsqu'une remontée jusqu'au niveau IP est demandée.

Les principes de MPLS

Les caractéristiques les plus importantes de la norme MPLS sont les suivantes :

- Spécification des mécanismes pour transporter des flots de paquets IP avec diverses granularités des flots entre un LSR d'entrée et un LSR de sortie. La granularité désigne la grosseur du flot, lequel peut intégrer plus ou moins de flots utilisateur.
- Indépendance du niveau trame et du niveau paquet. En fait, seul le transport de paquets IP est pris en compte.
- Mise en relation de l'adresse IP du destinataire avec une référence d'entrée dans le réseau.
- Reconnaissance par les routeurs de bord des protocoles de routage, de type OSPF, et de signalisation, tels LDP (*Label Distribution Protocol*) ou RSVP.
- Utilisation de différents types de trames.

Quelques propriétés supplémentaires méritent d'être soulignées :

- L'ouverture du circuit virtuel est fondée sur la topologie, bien que d'autres possibilités soient également définies dans la norme.
- L'assignation des références est effectuée par l'aval, c'est-à-dire à la demande d'un nœud qui émet un message dans la direction de l'émetteur.
- La granularité des références est variable.
- Le stock des références est géré selon la méthode « dernier arrivé premier servi ».
- Il est possible de hiérarchiser les demandes.
- Un temporisateur TTL (*Time to Live*) est utilisé.
- Une référence est encapsulée dans la trame, incluant un TTL et une qualité de service (QoS).

Le point fort du protocole MPLS est la possibilité, illustrée à la figure 22, de transporter les paquets IP sur plusieurs types de réseaux commutés. Il est ainsi acceptable de passer d'un réseau ATM à un réseau Ethernet ou à un réseau relais de trames. En d'autres termes, il peut s'agir de n'importe quel type de trame, à partir du moment où une référence peut y être incluse.

Fonctionnement de MPLS

Les transmissions de données s'effectuent sur des chemins nommés LSP (*Label Switched Path*). Un LSP est une suite de références partant de la source et allant jusqu'à la destination. Les LSP sont établis avant la transmission des données (*control-driven*) ou à la détection d'un flot qui souhaite traverser le réseau (*data-driven*). Les

références qui sont incluses dans les trames sont distribuées en utilisant un protocole de signalisation, dont le plus important est LDP (*Label Distribution Protocol*), mais aussi RSVP (*Resource reSerVation Protocol*), éventuellement associé à un protocole de routage, comme BGP (*Border Gateway Protocol*) ou OSPF (*Open Shortest Path First*). Les trames acheminant les paquets IP transportent les références de nœud en nœud.

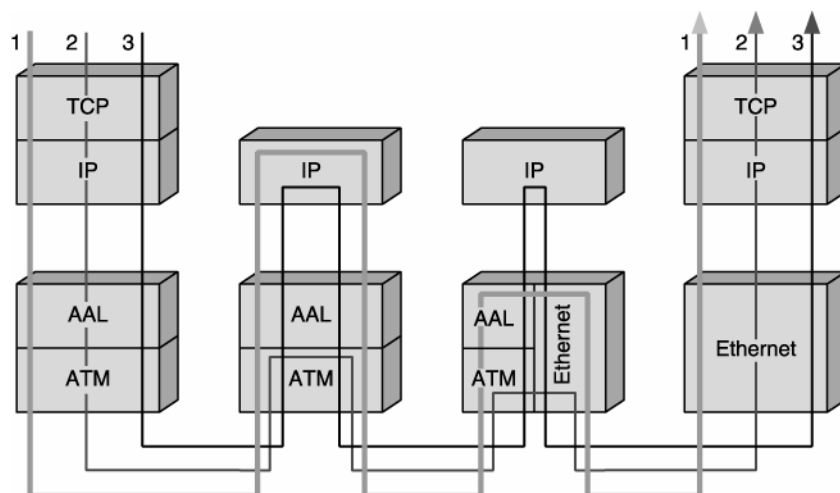


Figure 22. Un réseau MPLS avec des sous-réseaux distincts.

LSR (*Label Switched Router*) et LER (*Label Edge Router*)

Les nœuds qui participent à MPLS sont appelés LER (*Label Switched Router*) et LSR (*Label Edge Router*). Un LSR est un routeur dans le cœur du réseau qui participe à la mise en place du circuit virtuel par lequel les trames sont acheminées. Un LER est un nœud d'accès au réseau MPLS. Un LER peut avoir des ports multiples permettant d'accéder à plusieurs réseaux distincts, chacun pouvant avoir sa propre technique de commutation. Les LER jouent un rôle important dans la mise en place des références.

LSR (*Label Switched Router*)

Un équipement qui effectue à la fois une commutation sur une référence et un routage s'appelle un LSR. Les tables de commutation LSFT (*Label Switching Forwarding Table*) consistent en un ensemble de références d'entrées auxquelles correspondent des ports de sortie. À une référence d'entrée peuvent correspondre plusieurs files de sortie pour tenir compte des adresses multipoint.

La table de commutation peut être plus complexe. À une référence d'entrée peut correspondre le port de sortie du nœud dans une première sous-entrée mais aussi, dans une deuxième sous-entrée, un deuxième port de sortie correspondant à la file de sortie du prochain nœud qui sera traversé, et ainsi de suite. De la sorte, à une référence peuvent correspondre l'ensemble des ports de sortie qui seront empruntés lors de l'acheminement du paquet.

Les tables de commutation peuvent être spécifiques de chaque port d'entrée d'un LSR et regrouper des informations supplémentaires, comme une qualité de service ou une demande de ressources particulière.

FEC (*Forward Equivalence Class*)

Une FEC (*Forward Equivalence Class*) est une classe représentant un ensemble de flots qui partagent les mêmes propriétés. Toutes les trames d'une FEC ont le même traitement dans les nœuds du réseau MPLS. Les trames sont introduites dans une FEC au nœud d'entrée et ne peuvent plus être distinguées des autres flots à l'intérieur de la classe.

Une FEC peut être bâtie de différentes façons et avoir une adresse de destination bien déterminée, un même préfixe d'adresse, une même classe de service, etc. Chaque LSR possède une table de commutation qui indique les références associées aux FEC. Toutes les trames d'une même FEC sont transmises sur la même interface de sortie. Cette table de commutation est appelée LIB (*Label Information Base*).

Les références utilisées par les FEC peuvent être regroupées de deux façons :

- Par plate-forme. Les valeurs des références sont uniques sur l'ensemble des LSR d'un domaine, et les références sont distribuées sur un ensemble commun géré par un nœud particulier.

- Par interface. Les références sont gérées par interface, et une même valeur de référence peut se retrouver sur deux interfaces différentes.

MPLS et les références

Une référence en entrée permet de déterminer la FEC par laquelle transite le flot. Cette solution ressemble beaucoup à la notion de conduit virtuel dans le monde ATM, dans lequel les circuits virtuels sont multiplexés. Ici, nous avons la même chose, avec un multiplexage de tous les circuits virtuels à l'intérieur d'une FEC, de telle sorte que, dans ce conduit, l'on ne puisse plus distinguer les circuits virtuels. Le LSR examine la référence et envoie la trame dans la direction indiquée. On voit le rôle capital joué par les LER, qui assignent aux flots de paquets des références qui permettent de commuter les trames sur le bon LSP. La référence n'a de signification que localement, puisqu'il y a modification de sa valeur sur la liaison suivante.

Une fois le paquet classifié dans un FEC, une référence est assignée à la trame qui va le transporter. Cette référence détermine le point de sortie par le chaînage des références. Dans le cas des trames classiques, comme LAP-F du relais de trames ou d'ATM, la référence est positionnée dans le DLCI (*Data Link Connexion Identifier*) ou dans le VPI/VCI. Dans les autres cas, il faut ajouter un champ, qui est le « shim label ».

MPLS et l'ingénierie de trafic

Il est difficile de réaliser une ingénierie de trafic dans Internet. L'une des raisons à cela est que le protocole BGP (*Border Gateway Protocol*) n'utilise que des informations de topologie du réseau. Pour réaliser une ingénierie de trafic efficace, l'IETF a introduit dans l'architecture MPLS un routage à base de contrainte, CR-LDP (*Constraint-based Routing-LDP*), et un protocole de routage interne à état des liens étendu, RSVP-TE (*RSVP-Traffic Engineering*).

Comme nous l'avons vu, chaque trame encapsulant un paquet IP qui entre dans le réseau MPLS se voit ajouter par le LSR d'entrée, ou *Ingress LSR*, une référence au niveau de l'en-tête. Cette référence permet d'acheminer la trame dans le réseau, les chemins étant préalablement ouverts par un protocole de signalisation, RSVP ou LDP. À la sortie du réseau, le label ajouté à l'en-tête de la trame est supprimé par le LSR de sortie, ou *Egress LSR*. Au LSP (*Label Switched Path*), qui est le chemin construit entre le LSR d'entrée et le LSR de sortie, sont associés des attributs permettant de contrôler les ressources attribuées à ces chemins. Ces attributs sont détaillés au tableau ci-dessous. Ils regroupent essentiellement la bande passante nécessaire au chemin, son niveau de priorité, son aspect dynamique par l'intermédiaire du protocole utilisé pour son ouverture et sa flexibilité en cas de panne.

ATTRIBUT	DESCRIPTION
Bande passante	Besoins minimaux de bande passante à réserver sur le chemin du LSP
Attribut de chemin	Indique si le chemin du LSP doit être spécifié manuellement ou dynamiquement par l'algorithme CR-LDP (<i>Constraint-based Routing-LDP</i>).
Priorité de démarrage	Le LSP le plus prioritaire se voit allouer une ressource demandée par plusieurs LSP.
Priorité de préemption	Indique si une ressource d'un LSP peut lui être retirée pour être attribuée à un autre LSP plus prioritaire.
Affinité ou couleur	Exprime des spécifications administratives.
Adaptabilité	Indique si le chemin d'un LSP doit être modifié pour avoir un chemin optimal.
Flexibilité	Indique si le LSP doit être rerouté dans le cas d'une panne sur le chemin du LSP.

Tableau • Attributs des chemins LSP dans un réseau MPLS.

Les réseaux Ethernet

Il existe deux grands types de réseaux Ethernet : ceux utilisant le mode partagé et ceux utilisant le mode commuté. Dans le premier cas, un même support physique est partagé entre plusieurs utilisateurs, de telle sorte que le coût de connexion soit particulièrement bas. Dans le second, chaque machine est connectée à un nœud de transfert, et ce dernier commute ou route la trame Ethernet vers un nœud suivant. Ce chapitre décrit en détail ces différentes configurations. Il se conclut en présentant l'arrivée du multimédia dans les réseaux Ethernet.

- La trame Ethernet
- L'Ethernet partagé
- L'Ethernet commuté
- Les réseaux Ethernet partagés et commutés
- Ethernet et le multimédia

La trame Ethernet

Le bloc transporté sur un réseau Ethernet a été normalisé par l'organisme américain IEEE, après avoir été défini à l'origine par le triumvirat d'industriel Xerox, Digital et Intel.

Ce bloc appartient à la famille des trames, car il contient un champ capable de déterminer son début et sa fin. La structure de la trame Ethernet est illustrée à la figure 23. Deux possibilités de trames Ethernet coexistent, l'une correspondant à la version primitive du triumvirat fondateur et l'autre à la normalisation par l'IEEE.

Format de la trame IEEE



Format de l'ancienne trame Ethernet



DA : Adresse récepteur
FCS : Frame Check Sequence
LLC : Logical Link Control

SA : Adresse émetteur
SFD : Synchronous Frame Delimitation
Synch : Synchronization

Figure 23. Structure de la trame Ethernet.

La trame Ethernet commence par un *préambule* et une *zone de délimitation* tenant au total sur 6 octets. Le préambule est une suite de 56 ou 62 bits de valeur 1010...1010. Il est suivi, dans le cas de la trame IEEE, par la zone de début de message SFD (*Start Frame Delimiter*), de valeur 10101011, ou, pour l'ancienne trame, de 2 bits de synchronisation. Ces deux séquences sont en fait identiques, et seule la présentation diffère. Le drapeau de début de trame sur 6 octets au total est suffisamment long pour qu'il soit impossible de la retrouver dans la séquence d'éléments binaires qui suit.

préambule.— Zone située en tête de la trame Ethernet permettant au récepteur de se synchroniser sur le signal et d'en reconnaître le début.

zone de délimitation.— Zone située juste derrière le préambule d'une trame Ethernet et indiquant la fin de la zone de début de trame.

La trame contient l'adresse de l'émetteur et du récepteur, chacune sur 6 octets. Ces adresses sont dotées d'une forme spécifique du monde Ethernet, conçue de telle sorte qu'il n'y ait pas deux coupleurs dans le monde qui possèdent la même adresse. On parle d'un *adressage plat*, construit de la façon suivante : les trois premiers octets correspondent à un numéro de constructeur, et les trois suivants à un numéro de série. Dans les trois premiers octets, les deux bits initiaux ont une signification particulière. Positionné à 1, le premier bit indique une adresse de groupe. Si le deuxième bit est également à la valeur 1, cela indique que l'adresse ne suit pas la structure normalisée.

adressage plat.— Ensemble dans lequel les adresses n'ont aucune relation les unes avec les autres.

Regardons dans un premier temps la suite de la trame IEEE. La zone Longueur (*Length*) indique la longueur du champ de données provenant de la couche supérieure. Ensuite, la trame encapsule le bloc de niveau trame proprement dit, ou trame LLC (*Logical Link Control*). Cette trame encapsulée contient une zone PAD, qui permet de remplir le champ de données de façon à atteindre la valeur de 46 octets, qui est la longueur minimale que doit atteindre cette zone pour que la trame totale fasse 64 octets en incluant les zones de préambule et de délimitation.

Dans l'ancienne trame Ethernet, intervient un type, qui indique comment se présente la zone de données (*Data*) transportée par la trame Ethernet. Par exemple, si la valeur de cette zone est 0800 en hexadécimal, cela signifie que la trame Ethernet transporte un paquet IP.

La raison pour laquelle les deux types de trames sont compatibles et peuvent coexister sur un même réseau est expliquée à la question 4, à la fin de cette section.

La détection des erreurs est assurée par le biais d'un polynôme générateur $g(x)$ selon la formule :

$$g(x) = x^{32} + x^{26} + x^{22} + x^{16} + x^{12} + x^{11} + x^{10} + x^8 + x^7 + x^5 + x^4 + x^2 + x^1.$$

Ce polynôme donne naissance à une *séquence de contrôle* (CRC) sur 4 octets.

séquence de contrôle (CRC, *Cyclic Redundancy Checksum*).— Cette séquence, encore appelée centrale de contrôle, est obtenue comme le reste de la division de polynômes du polynôme généré par les éléments binaires de la trame par le polynôme générateur.

Pour se connecter au réseau Ethernet, une machine utilise un coupleur, c'est-à-dire une carte que l'on insère dans la machine et qui supporte le logiciel et le matériel réseau nécessaires à la connexion.

Comme nous l'avons vu, la trame Ethernet comporte un préambule. Ce dernier permet au récepteur de synchroniser son horloge et ses divers circuits physiques avec l'émetteur, de façon à réceptionner correctement la trame. Dans le cas d'un réseau *Ethernet partagé*, tous les coupleurs sur le réseau enregistrent la trame au fur et à mesure de son passage. Le composant électronique chargé de l'extraction des données incluses dans la trame vérifie la concordance entre l'adresse de la station de destination portée dans la trame et l'adresse du coupleur. S'il y a concordance, le paquet est transféré vers l'utilisateur après vérification de la conformité de la trame par le biais de la séquence de contrôle.

Ethernet partagé.— Réseau comportant un support physique commun à l'ensemble des terminaux. Le support commun peut être aussi bien un câble métallique qu'une fibre optique ou une fréquence hertzienne.

L'Ethernet partagé

Un réseau Ethernet partagé présente un support physique commun à l'ensemble des terminaux. Le support commun peut être aussi bien un câble métallique qu'une fibre optique ou une fréquence hertzienne. La raison de ce support unique s'explique par une volonté de simplification de l'infrastructure du réseau. La figure 24 illustre quelques supports partagés.

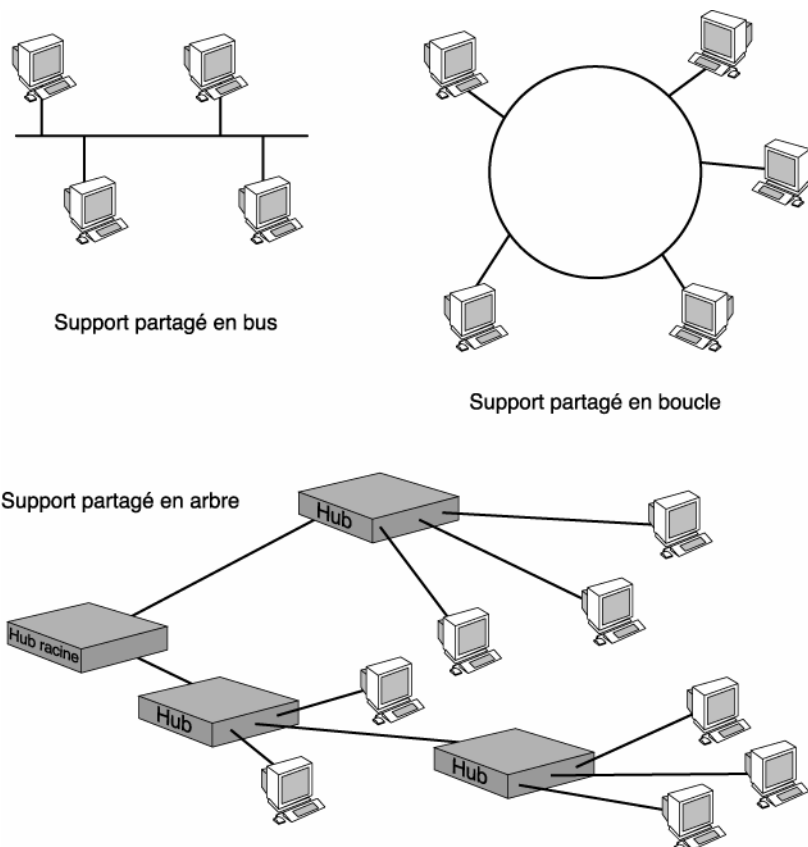


Figure 24. Supports partagés.

Si le support est partagé, il est nécessaire que le réseau dispose d'un protocole apte à décider quelle station a le droit d'émettre, sinon une collision des différents signaux émis par les stations est inévitable et engendre une perte de l'information. Ce protocole, appelé MAC (*Medium Access Control*), doit permettre à une station de savoir si elle a le droit d'émettre ou non.

La couche MAC qui gère le protocole d'accès au support physique se situe dans la couche 2, ou niveau trame, du modèle de référence. Le niveau trame, également appelé couche liaison dans l'ancien modèle, prend le nom de LLC (*Logical Link Control*) dans l'environnement des réseaux locaux partagés.

LLC (*Logical Link Control*).— Dans l'environnement des réseaux Ethernet partagés, couche équivalente à la couche liaison du modèle de référence originel.

Dans un autre réseau célèbre, le Token-Ring, la solution au problème des collisions consiste à attribuer à une trame particulière le rôle de jeton. Dans ce cas, seule la station qui possède cette trame a le droit de transmettre. On reproche à ce système son temps de latence, le temps de passage du jeton d'une station à une autre station devenant d'autant plus élevé que les stations sont éloignées.

Très différente, la solution Ethernet s'appuie sur une technique dite d'écoute de la porteuse et de détection des collisions. Cette technique, plus connue sous le nom de CSMA (*Carrier Sense Multiple Access*), consiste à écouter le canal avant d'émettre. Si le coupleur détecte un signal sur la ligne, il diffère son émission à une date ultérieure.

Cette méthode réduit considérablement le risque de collision, sans toutefois le supprimer complètement. Si, durant le temps de propagation entre les deux stations les plus éloignées (période de vulnérabilité), une trame est émise sans qu'un coupleur la détecte, il peut y avoir superposition de signaux. De ce fait, il faut réémettre ultérieurement les trames perdues.

La technique Ethernet complète d'accès et de détection des collisions s'appelle CSMA/CD (CD, pour *Collision Detection*). Elle est illustrée à la figure 22. C'est la méthode normalisée par l'ISO. À l'écoute préalable du réseau s'ajoute l'écoute pendant la transmission : un coupleur prêt à émettre et ayant détecté le canal libre transmet et continue à écouter le canal. De ce fait, s'il se produit une collision, il interrompt dès que possible sa transmission et envoie des signaux spéciaux, appelés « jam sequence » en sorte que tous les coupleurs soient prévenus de la collision. Il tente de nouveau son émission ultérieurement, suivant un algorithme de redémarrage, appelé algorithme de *back-off*.

jam sequence.— Bits que l'on ajoute pour que la longueur de la trame Ethernet atteigne au moins 64 octets.

back-off.— Algorithme de redémarrage après collision destiné à éviter une nouvelle collision.

Ce système permet une détection immédiate des collisions et une interruption de la transmission en cours sans perte de temps. Les coupleurs émetteurs reconnaissent une collision en comparant le signal émis avec celui qui passe sur la ligne, c'est-à-dire par détection d'interférences, et non plus par une absence d'acquittement.

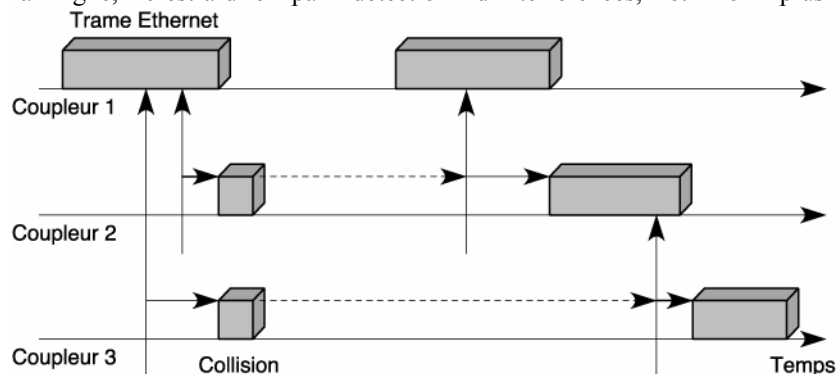


Figure 25. Principe du CSMA/CD.

Cette méthode de détection des conflits est relativement simple, mais elle nécessite des techniques de codage suffisamment performantes pour permettre de reconnaître facilement une superposition de signaux.

On peut calculer la longueur maximale d'un réseau Ethernet à partir de la longueur minimale de la trame Ethernet. En effet, il faut absolument éviter qu'une station puisse finir sa transmission en ignorant qu'il y a eu une collision sur sa trame, sinon elle penserait que la transmission s'est effectuée avec succès. C'est la raison pour laquelle l'IEEE a fixé la longueur minimale de la trame Ethernet à 64 octets. La station d'émission de la trame Ethernet ne peut donc se retirer avant d'avoir fini d'émettre ses 64 octets. Si le signal de collision doit lui revenir avant la fin de l'émission, le temps maximal correspondant, pour un réseau dont la vitesse est de 10 Mbit/s, est le suivant : $64 \times 8 = 512$ bits à $0,1 \mu\text{s}$ par bit, soit $51,2 \mu\text{s}$.

La figure 26 illustre le pire cas de retard dans la reconnaissance d'une collision : la station émet depuis une extrémité du support, et son signal doit se propager jusqu'à la station située à l'autre extrémité. Au moment où le signal arrive à cette station, celle-ci émet à son tour, ce qui provoque une collision. Le temps de propagation de la collision jusqu'à l'émetteur étant identique à celui de la propagation de l'émetteur vers le récepteur, il faut considérer un temps aller-retour avant que l'émetteur s'aperçoive de la collision.

Comme le temps maximal pendant lequel on est sûr que la station émettrice écoute correspond à $51,2 \mu\text{s}$, le délai de propagation aller-retour ne doit pas dépasser pas cette valeur.

La distance maximale entre deux points d'un réseau Ethernet partagé est donc déterminée par un temps d'aller-retour de $51,2 \mu\text{s}$, ce qui correspond à une valeur de $25,6 \mu\text{s}$ de propagation dans un seul sens. Cette distance maximale est donc de $5,12 \text{ km}$. Cependant, certains équipements, comme les répéteurs ou les hubs que nous verrons dans la suite, obligent à réduire cette distance. Dans le commerce la limite est souvent réduite aux environs de $2,5 \text{ km}$.

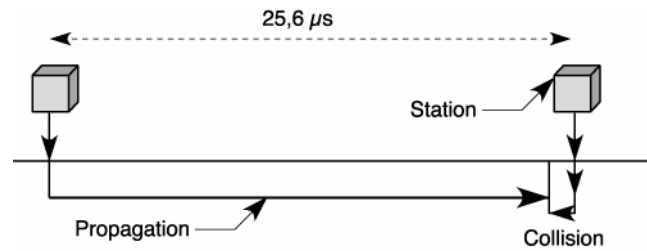


Figure 26. Le cas le pire pour la détection d'un signal de collision.

Si une collision se produit, le module d'émission-réception émet un signal pour, d'une part, interrompre la collision et, d'autre part, initialiser la procédure de retransmission. L'interruption de la collision intervient après l'envoi d'une séquence de bourrage (*jam*), qui vérifie que la durée de la collision est suffisante pour être remarquée par toutes les stations en transmission impliquées dans la collision.

Il est nécessaire de définir plusieurs paramètres pour expliquer la procédure de reprise sur une collision. Le temps aller-retour maximal correspond au temps qui s'écoule entre les deux points les plus éloignés du réseau local, à partir de l'émission d'une trame jusqu'au retour d'un signal de collision. Nous avons vu que cette valeur était de 51,2 μ s, soit de 512 temps d'émission d'un bit. Ethernet définit encore « une tranche de temps », qui est le temps minimal avant retransmission. Cette tranche vaut évidemment 51,2 μ s. Le temps avant retransmission d'une station dépend également du nombre n de collisions déjà effectuées par cette station. Le délai aléatoire de retransmission dans Ethernet est un multiple r de la tranche de temps, suivant l'algorithme : $0 \leq r < 2^k$, où $k = \text{minimum}(n, 10)$ et n le nombre de collisions déjà effectuées.

La reprise s'effectue après le temps $r \times 51,2 \mu$ s.

Si, au bout de seize essais, la trame est encore en collision, l'émetteur abandonne sa transmission. Une reprise s'effectue alors à partir des protocoles de niveaux supérieurs.

Lorsque deux trames entrent en collision pour la première fois, elles ont une chance sur deux d'entrer de nouveau en collision, puisque $r = 1$ ou 0. Il vaut cependant mieux essayer d'émettre, quitte à provoquer une collision de courte durée, plutôt que de laisser le support vide.

Un calcul simple montre que les temps de retransmission après une dizaine de collisions successives ne représentent que quelques millisecondes, c'est-à-dire un temps encore très court. La méthode CSMA/CD est toutefois une technique probabiliste, et il est difficile de bien cerner le temps qui s'écoule entre l'arrivée de la trame dans le coupleur de l'émetteur et le départ de la trame du coupleur récepteur jusqu'au destinataire. Ce temps dépend bien sûr du nombre de collisions, mais aussi, indirectement, du nombre de stations, de la charge du réseau et de la distance moyenne entre deux stations. Plus le temps de propagation est grand, plus le risque de collision est important. De nombreuses courbes de performances montrent le débit réel en fonction du débit offert (le débit provenant des nouvelles trames additionné au débit provoqué par les retransmissions). Des courbes de performances sont illustrées à la figure 27.

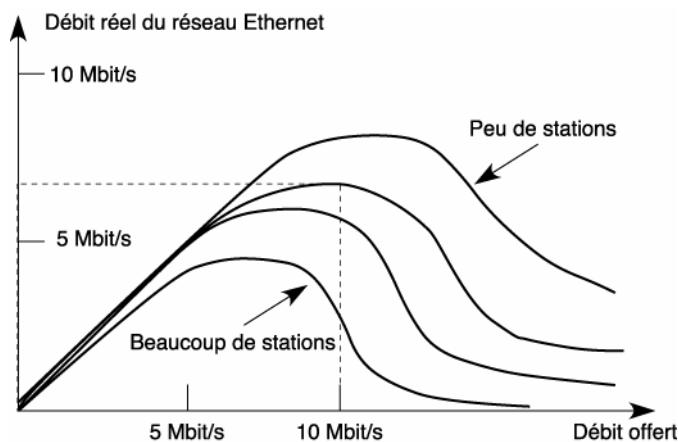


Figure 27. Performances du réseau Ethernet partagé.

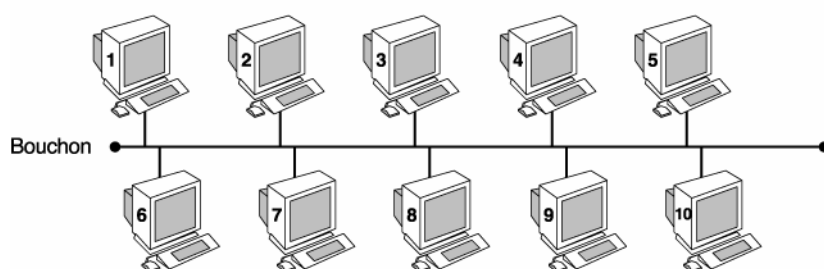
L'Ethernet commuté

L'Ethernet commuté consiste à utiliser la trame Ethernet dans un réseau de transfert dont les nœuds sont des commutateurs. On utilise dans ce cas un Ethernet particulier, ou Ethernet FDSE (*Full-Duplex Switched Ethernet*), sur les lignes de communication duquel il est possible d'envoyer des trames Ethernet dans les deux sens simultanément. L'avantage de la commutation Ethernet par rapport à l'Ethernet partagé est de ne pas imposer de distance maximale entre deux nœuds étant donné qu'il n'y a plus de risque de collision.

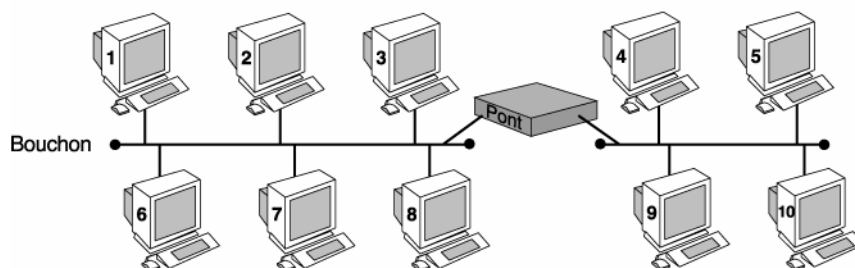
Avant d'en arriver à une technique purement commutée, on a commencé, historiquement, par découper les réseaux Ethernet en des sous-réseaux autonomes, reliés les uns aux autres par des passerelles, ou ponts, tout en essayant de conserver un trafic local.

Un pont est un organe intelligent capable de reconnaître l'adresse du destinataire et de savoir s'il faut ou non retransmettre la trame vers un autre segment Ethernet. Ajouté au milieu d'un réseau Ethernet, un pont découpe le réseau en deux Ethernet indépendants. De ce fait, le trafic est multiplié par le nombre de sous-réseaux. Les ponts ne sont dans ce cas que des commutateurs Ethernet, qui mémorisent les trames et les réémettent vers d'autres réseaux Ethernet. La logique extrême d'un tel système est de découper le réseau jusqu'à n'avoir qu'une seule station par réseau Ethernet. On obtient alors la commutation Ethernet. Ce cheminement est illustré à la figure 28.

Réseau partagé



Deux réseaux partagés reliés par un pont



Réseau commuté

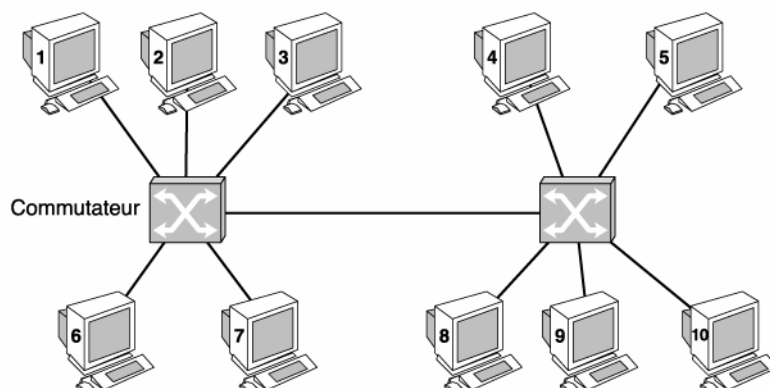


Figure 28. Passage vers la commutation Ethernet.

Dans la commutation Ethernet, chaque carte coupleur est reliée directement à un commutateur Ethernet, ce dernier se chargeant de rediriger les trames dans la bonne direction. La commutation demande une référence qui, *a priori*, n'existe pas dans le monde Ethernet. Aucun paquet de supervision n'ouvre le circuit virtuel en posant des

références. Le mot de commutateur peut être considéré comme inexact puisqu'il n'y a pas de référence. Il est cependant possible de parler de commutation, si l'on considère l'adresse du destinataire comme une référence. Le circuit virtuel est déterminé par la suite de références dont les valeurs sont déterminées par l'adresse du destinataire sur 6 octets. Il faut, pour réaliser cette commutation de bout en bout, que chaque commutateur ait la possibilité de déterminer la liaison de sortie en fonction de l'adresse du récepteur. L'algorithme qui permet de mettre en place les tables de commutation s'appelle **Spanning Tree**. Il détermine les routes à suivre par un algorithme de décision, qui prend en compte les liaisons par un poids.

Spanning Tree (arbre recouvrant).– Algorithme permettant de disposer les nœuds d'un réseau sous la forme d'un arbre avec un nœud racine. Les connexions à utiliser pour aller d'un point à un autre du réseau sont celles désignées par l'arbre. Cette solution garantit l'unicité du chemin et évite les duplications de paquets.

Cette technique de commutation peut présenter des difficultés liées à la fois à la gestion des adresses de tous les coupleurs raccordés au réseau pour déterminer les tables de commutation et à celle des congestions éventuelles au sein d'un commutateur. Il faut donc mettre en place des techniques de contrôle, qui limitent, sur les liaisons entre commutateurs, les débits provenant simultanément de tous les coupleurs Ethernet. On retrouve là tous les problèmes posés par une architecture de type commutation de trames.

L'environnement Ethernet s'impose actuellement par sa simplicité de mise en œuvre, tant que le réseau reste de taille limitée. Cette solution présente l'avantage de s'appuyer sur l'existant, à savoir les coupleurs et les divers réseaux Ethernet que de nombreuses sociétés ont mis en place pour créer leur réseau local. Comme tous les réseaux de l'environnement Ethernet sont compatibles entre eux, toutes les machines émettant des trames Ethernet peuvent facilement s'interconnecter. On peut ainsi réaliser des réseaux extrêmement complexes, avec des segments partagés sur les parties locales et des liaisons commutées sur les longues distances ou entre les commutateurs Ethernet.

L'un des avantages de cette technique est de ne plus présenter de limitation de distance puisque l'on est en mode commuté. Les distances entre machines connectées peuvent atteindre plusieurs milliers de kilomètres. Un autre avantage est offert par l'augmentation des débits par terminaux puisque, comme mentionné précédemment, la capacité en transmission peut atteindre 10, 100, 1 000 Mbit/s et même 10 Gbit/s par machine. Des débits de 40 puis 80 Gbit/s seront bientôt disponibles.

L'inconvénient majeur du système est de retourner à un mode commuté, dans lequel les nœuds de commutation doivent gérer un contrôle de flux et un routage et effectuer une gestion des adresses physiques des coupleurs. En d'autres termes, chaque commutateur doit connaître l'adresse MAC de tous les coupleurs connectés au réseau et savoir dans quelle direction envoyer les trames.

Dans les entreprises, le réseau Ethernet peut consister en une association de réseaux partagés et de réseaux commutés, tous les réseaux Ethernet étant compatibles au niveau de la trame émise. Si le réseau de l'entreprise est trop vaste pour permettre une gestion de toutes les adresses dans chaque commutateur, il faut alors diviser le réseau en domaines distincts et passer d'un domaine à un autre, en remontant au niveau paquet de l'architecture de référence, c'est-à-dire en récupérant l'information transportée dans la zone de données de la trame et en se servant de l'adresse de niveau paquet pour effectuer le routage. Cet élément de transfert n'est autre qu'un routeur effectuant une décapsulation puis une encapsulation.

La commutation peut se faire par port, et, dans ce cas, les coupleurs sont directement connectés au commutateur, ou par segment, et ce sont des segments de réseaux Ethernet qui sont interconnectés. Ces deux possibilités sont illustrées à la figure 29.

Du point de vue de la commutation elle-même, les deux techniques de commutation suivantes sont également disponibles dans les équipements :

- Le **store-and-forward**, dans lequel un paquet Ethernet est stocké en entier dans les mémoires du commutateur, puis examiné avant d'être retransmis sur une ligne de sortie.

Le **cut-through**, ou **fast-forward**, dans lequel le paquet Ethernet peut commencer à être retransmis vers le nœud suivant dès que la zone de supervision a été traitée, sans se soucier si la fin du paquet est arrivée ou non dans le nœud. Dans cette solution, il est possible qu'un même paquet Ethernet soit transmis simultanément sur plusieurs liaisons : le début du paquet sur une première liaison, la suite du paquet sur une deuxième et la fin sur une troisième liaison, comme illustré à la figure 30.

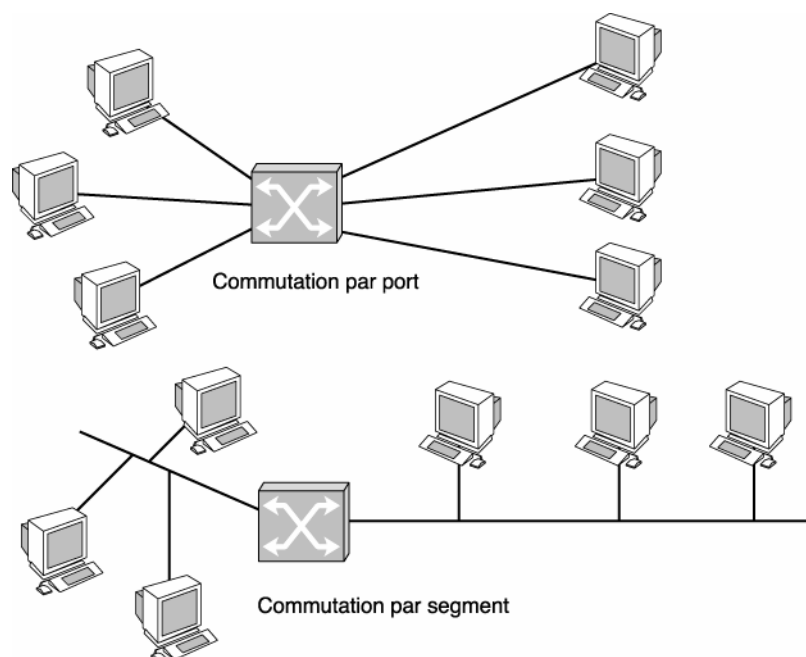


Figure 29. Les deux types de commutation.

La seconde solution présente plusieurs inconvénients. Elle ne permet pas de contrôler la correction du paquet, et la fin du paquet peut ne plus exister à la suite d'une collision.

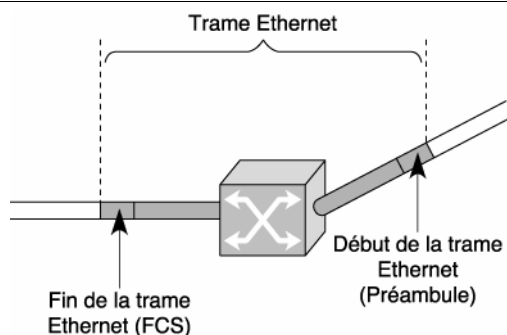


Figure 30. Le cut-through.

Les réseaux Ethernet partagés et commutés

Normalisée par le groupe de travail IEEE 802.3, la norme Ethernet est née de recherches effectuées au début des années 70 sur les techniques d'accès aléatoire. La société Xerox a développé la première des prototypes de cette solution. La promotion de ce produit Ethernet a été assurée en grande partie par le triumvirat Digital, Intel et Xerox. Aujourd'hui, toutes les grandes compagnies informatiques qui vendent des produits de réseau possèdent à leur catalogue tout l'arsenal Ethernet.

On peut caractériser les produits Ethernet par la technique d'accès CSMA/CD, avec un débit de 1, 10, 100, 1 000 Mbit/s et 10 Gbit/s. Cheapernet correspond à un Ethernet utilisant un câble plus fin (*thin cable*) mais en conservant les mêmes capacités de transmission. Starlan utilise une topologie très différente et à des vitesses de 1 Mbit/s, pour la première génération, 10 Mbit/s, pour la deuxième, 100 Mbit/s, pour la troisième, et 1 000 Mbit/s pour la quatrième génération. Fast Ethernet est le nom des réseaux à 100 Mbit/s, Gigabit Ethernet, ou 1GbE, celui des réseaux à 1 000 Mbit/s et 10 Gigabit Ethernet, ou 10GbE, celui des réseaux à 10 Gbit/s.

Le nombre de réseaux Ethernet normalisés par l'IEEE est impressionnant. La liste ci-après est empruntée à la nomenclature officielle.

- IEEE 802.3 10base5 (Ethernet jaune).

- IEEE 802.3 10base2 (Cheapernet, Thin Ethernet).
- IEEE 802.3 10broad36 (Ethernet large bande).
- IEEE 802.3 1base5 (Starlan à 1 Mbit/s).
- IEEE 802.3 10baseT, Twisted-pair (Ethernet sur paire de fils torsadés).
- IEEE 802.3 10baseF, Fiber Optic (Ethernet sur fibre optique) :
 - 10baseFL, Fiber Link ;
 - 10baseFB, Fiber Backbone ;
 - 10baseFP, Fiber Passive.
- IEEE 802.3 100baseT, Twisted-pair ou encore Fast Ethernet (Ethernet 100 Mbit/s en CSMA/CD), qui se décompose en :
 - 100baseTX ;
 - 100baseT4 ;
 - 100baseFX.
- IEEE 802.3z 1000baseCX, qui utilise deux paires torsadées de 150 Ohm.
- IEEE 802.3z 1000baseLX, qui utilise une paire de fibre optique avec une longueur d'onde élevée.
- IEEE 802.3z 1000base SX, qui utilise une paire de fibre optique avec une longueur d'onde courte.
- IEEE 802.3z 1000baseT, qui utilise quatre paires de catégorie 5 UTP.
- IEEE 802.ah EFM (Ethernet in the First Mile) pour les réseaux d'accès.
- IEEE 802.3 10Gbase T, qui utilise quatre paires torsadées.
- IEEE 802.3ak 10Gbase-CX4, qui utilise des paires de câbles coaxiaux.
- IEEE 802.11b/a/g pour les réseaux Ethernet hertziens.
- IEEE 802.17 pour les réseaux métropolitains.

La signification des sigles a évolué. La technique utilisée est citée en premier. IEEE 802.3 correspond à CSMA/CD, IEEE 802.3 Fast Ethernet à une extension de CSMA/CD, IEEE 802.9 à une interface CSMA/CD à laquelle on ajoute des canaux B, IEEE 802.11 à un Ethernet par voie hertzienne, etc. Vient ensuite la vitesse, puis la modulation ou non (base = bande de base et broad = broadband, ou large bande). Le sigle se termine par un élément qui était à l'origine la longueur d'un brin puis s'est transformé en type de support physique.

L'architecture de communication classique, que l'on trouve dans les entreprises, comporte comme épine dorsale, un réseau Ethernet sur lequel sont connectés des réseaux locaux de type capillaire (Cheapernet et Starlan). Ces derniers sont capables d'irriguer les différents bureaux d'une entreprise à des coûts nettement inférieurs à ceux du réseau Ethernet de base, ou Ethernet jaune (qui doit son nom à la couleur du câble coaxial utilisé). La figure 31 illustre l'architecture générale d'un environnement Ethernet.

Le câblage des réseaux capillaires

Les réseaux capillaires sont formés à partir du câblage partant du répartiteur d'étage. De plus en plus souvent, les nouveaux bâtiments sont précâblés avec une structure identique à celle des câblages du réseau téléphonique à partir du répartiteur d'étage. De nombreuses différences existent, cependant entre ces câblages, et notamment les suivantes :

- Câblage banalisé. Un même câble peut être utilisé pour raccorder un combiné téléphonique (ou les terminaux s'adaptant aux réseaux téléphoniques, comme les terminaux vidéotex) ou un terminal informatique par l'intermédiaire d'une prise spécialisée informatique (comme la prise hermaphrodite d'IBM).
- Câblage non banalisé. Une bien meilleure qualité est préconisée pour la partie informatique. Par exemple, le câble informatique se présente sous la forme de paires de fils torsadés blindés et est d'une qualité largement supérieure à celle du câblage téléphonique, de type 3. Ce dernier câble est réalisé à l'aide de une, deux, trois ou quatre paires, d'une qualité médiocre par rapport à celle dédiée à l'informatique.
- On peut réaliser divers types de réseaux capillaires à partir du système de câblage. Le choix de la qualité du câble est important si des contraintes de distance existent. Il vaut mieux limiter la distance entre le local technique et la périphérie à une cinquantaine de mètres. La recommandation américaine de l'ANSI propose une limitation à 295 pieds.

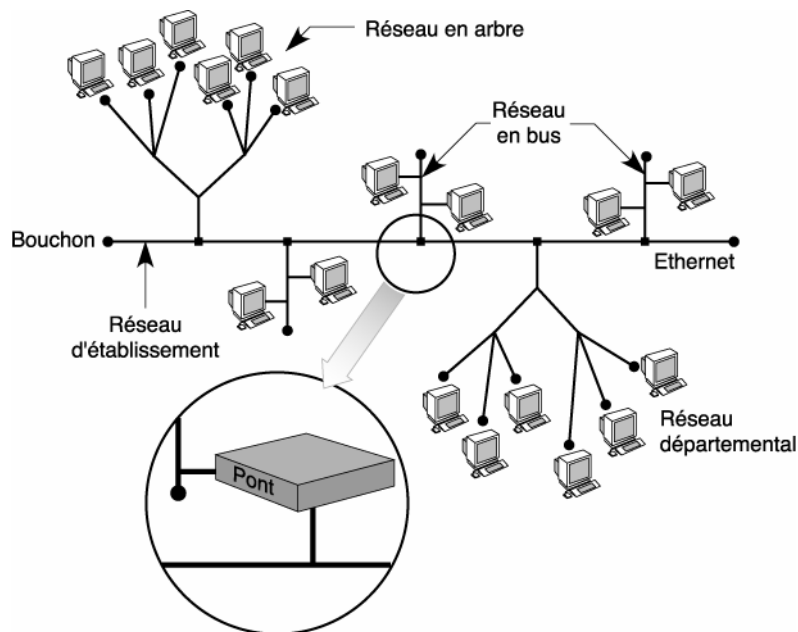


Figure 31. Architecture d'un réseau Ethernet d'entreprise.

Les caractéristiques d'un réseau Ethernet de base sont décrites dans la norme IEEE 8802.3 10base5. La topologie du réseau Ethernet comprend des brins de 500 m au maximum. Ces brins sont interconnectés par des répéteurs. Le raccordement des matériels informatiques peut s'effectuer tous les 2,5 m, ce qui permet jusqu'à 200 connexions par brin. Dans de nombreux produits, les spécifications indiquent que le signal ne doit jamais traverser plus de deux répéteurs et qu'un seul d'entre eux peut être éloigné. La régénération du signal s'effectue une fois franchie une ligne d'une portée de 1 000 m. On trouve une longueur maximale de 2,5 km correspondant à trois brins de 500 m et un répéteur éloigné (voir figure 32). Cette limitation de la distance à 2,5 km n'est cependant pas une caractéristique de la norme, la distance sans répéteur pouvant théoriquement atteindre 5,12 km.

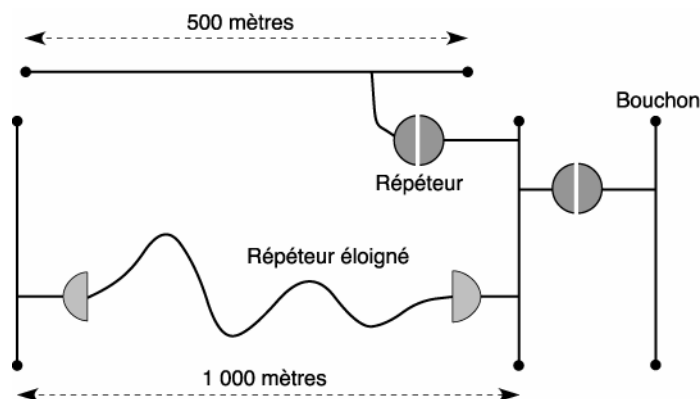


Figure 32. Topologie d'Ethernet.

Le groupe de travail IEEE 802.3u est à l'origine de la normalisation Fast Ethernet, l'extension à 100 Mbit/s du réseau Ethernet à 10 Mbit/s. La technique d'accès est la même que dans la version Ethernet à 10 Mbit/s, mais à une vitesse multipliée par 10. Les trames transportées sont identiques. Cette augmentation de vitesse se heurte au système de câblage et à la possibilité ou non d'y faire transiter des débits aussi importants. C'est la raison pour laquelle les trois sous-normes suivantes sont proposées pour le 100 Mbit/s :

- Le réseau IEEE 802.3 100baseTX, qui requiert deux paires non blindées (UTP) de catégorie 5 ou deux paires blindées (STP) de type 1.
- Le réseau IEEE 802.3 100baseT4, qui requiert quatre paires non blindées (UTP) de catégories 3, 4 et 5.
- Le réseau IEEE 802.3 100baseFX, qui requiert deux fibres optiques.

Les paires métalliques STP (*Shielded Twisted Pairs*) et UTP (*Unshielded Twisted Pairs*) correspondent à l'utilisation, pour le premier cas, d'un blindage qui entoure les paires et les protège et, pour le second, de paires non blindées qui peuvent être de moins bonne qualité mais dont la pose est grandement simplifiée.

La distance maximale entre les deux points les plus éloignés d'un réseau Fast Ethernet est fortement réduite par rapport à la version 10 Mbit/s. En effet, la longueur minimale de la trame est toujours de 64 octets, ce qui représente un temps de transmission de 5,12 μ s. On en déduit que la distance maximale qui peut être parcourue dans ce laps de temps est de l'ordre de 1 000 m, ce qui représente une longueur maximale d'approximativement 500 m. Comme le temps de traversée des hubs est relativement grand, la plupart des constructeurs ont limité la distance maximale à 210 m pour le Fast Ethernet.

L'avantage de cette solution est qu'elle permet une bonne compatibilité avec la version 10 Mbit/s, ce qui permet de relier sur un même hub à la fois des stations à 10 Mbit/s et à 100 Mbit/s. Le coût de connexion du 100 Mbit/s ne dépasse pas deux fois celui de l'Ethernet classique, dix fois moins rapide.

Les réseaux Fast Ethernet partagé servent souvent de réseaux d'interconnexion de réseaux Ethernet à 10 Mbit/s. Néanmoins la distance relativement limitée couverte par le Fast Ethernet partagé ne lui permet pas de recouvrir une entreprise de quelque importance.

Le Gigabit Ethernet partagé ne résout pas davantage ce problème. Comme nous le verrons, sa taille est similaire à celle du Fast Ethernet. Pour étendre la couverture du réseau Ethernet, la solution consiste à passer à des Fast Ethernet ou à des Gigabit Ethernet commutés. On trouve aujourd'hui dans les grandes entreprises des réseaux à transfert de trames Ethernet, qui utilisent des commutateurs Ethernet.

Le Gigabit Ethernet est la dernière évolution du standard Ethernet. Plusieurs améliorations ont été nécessaires par rapport au Fast Ethernet à 100 Mbit/s, notamment la modification du CSMA/CD. En effet, comme la longueur de la trame doit être compatible avec les autres options, la longueur minimale de 64 octets entraînerait un temps aller-retour maximal de 0,512 μ s et donc une distance maximale d'une cinquantaine de mètres dans le meilleur des cas, mais plus sûrement de quelques mètres. Pour que cette distance maximale soit agrandie, la trame émise sur le support doit avoir une longueur d'au moins 512 bits. Si la trame initiale est de 64 octets, le coupleur ajoute donc le complément en bit de « rembourrage » (*padding*). S'il s'agit là d'une bonne solution pour agrandir le réseau Gigabit, le débit utile n'en reste pas moins très faible si toutes les trames à transmettre ont une longueur de 64 octets.

La technique full-duplex commutée est la plus fréquente dans cette nouvelle catégorie de réseaux. Le Gigabit Ethernet accepte les répéteurs ou les hubs lorsqu'il y a plusieurs directions possibles. Dans ce dernier cas, un message entrant est recopié sur toutes les lignes de sortie. Des routeurs Gigabit sont également disponibles lorsqu'on remonte jusqu'au niveau paquet, comme avec le protocole IP. Dans ce cas, il faut récupérer le paquet IP pour pouvoir router la trame Ethernet. La figure 33 illustre une interconnexion de deux réseaux commutés par un routeur Gigabit.

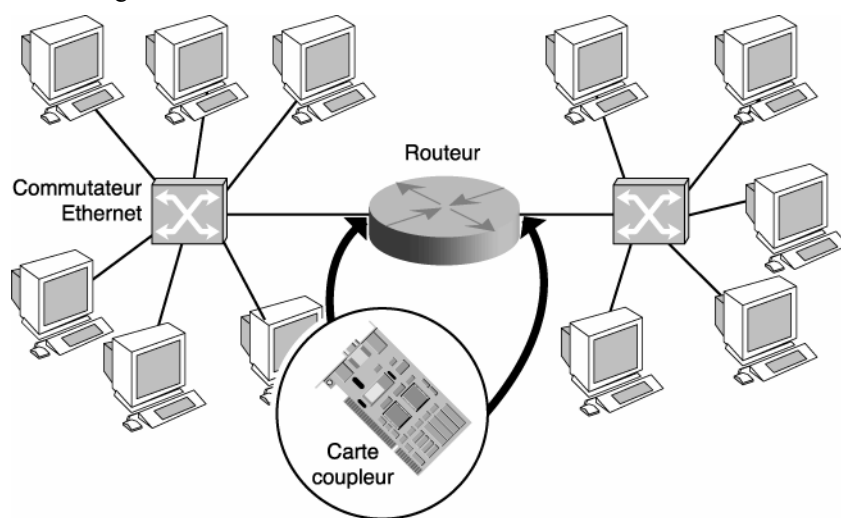


Figure 33. Interconnexion de deux réseaux Ethernet commutés par un routeur.

Ethernet et le multimédia

Ethernet n'a pas été conçu pour des applications multimédias mais informatiques. Pour se mettre à niveau et entrer dans le domaine du multimédia, l'environnement Ethernet a dû se transformer. Cette mutation concerne essentiellement l'Ethernet commuté. Pour réaliser des applications multimédias, l'IEEE a introduit une priorité de traitement des paquets dans les commutateurs Ethernet. Les paquets les plus prioritaires sont placés en tête des files d'attente, de telle sorte que des applications isochrones, comme la parole téléphonique, soient réalisables sur une

grande distance. On choisit, de préférence, des trames de la plus petite taille possible : 64 octets, contenant 46 octets de données.

Il est à noter que, dans le Gigabit Ethernet, on complète les trames de 64 octets jusqu'à ce qu'elles atteignent la valeur de 512 octets, cette extension permettant une distance maximale de 400 m entre les deux stations les plus éloignées. Cette technique de « rembourrage » (*padding*) n'est pas efficace pour la parole téléphonique, puisque seulement 46 octets sur 512 sont utilisés, ce qui ne rend pas le Gigabit Ethernet performant pour le transport de la parole téléphonique.

Pour remplir un paquet, il faut un temps de $46 \times 125 \mu s$, soit 5,75 ms, ce qui est un peu moins long que le temps nécessaire pour remplir une cellule ATM. En revanche, le paquet à transporter est beaucoup moins long en ATM (53 octets) qu'en Ethernet (64 octets, voire 512 octets dans le Gigabit Ethernet). Si le réseau est doté d'une technique de contrôle de flux permettant de ne pas perdre de paquet en utilisant d'une façon efficace les priorités, il est possible de transmettre de la parole dans Ethernet sans difficulté. Il en va de même pour la vidéo temps réel. Les applications temps réel, avec de fortes contraintes temporelles, sont réalisables sur les réseaux Ethernet.

Ethernet a été pendant très longtemps synonyme de réseau local. Cette limitation géographique s'explique par la technique d'accès qui est nécessaire sur les réseaux Ethernet partagé. Pour s'assurer que la collision a été bien perçue par la station d'émission avant que celle-ci ne se déconnecte, la norme Ethernet réclame que 64 octets au minimum soient émis, ce qui limite le temps aller-retour sur le support physique au temps de transmission de ces 512 bits. À partir du moment où l'on passe en commutation, la distance maximale n'a plus de sens. On utilise parfois le terme de WLAN (*Wide LAN*) pour indiquer que cette distance maximale atteint désormais le champ des réseaux étendus.

Les évolutions d'Ethernet ont été multiples pour que cette norme rejoigne les possibilités offertes par ses concurrents. Tout d'abord, la norme d'adressage Ethernet a été modifiée : de plat et absolu, cet adressage est devenu hiérarchique. Cette évolution est aujourd'hui consacrée par la norme 802.1q, qui étend la zone d'adressage grâce à un niveau hiérarchique supplémentaire. On appelle cette nouvelle solution de structuration du réseau un *VLAN (Virtual LAN)*.

VLAN (Virtual LAN).— Réseau logique dans lequel sont regroupés des clients qui ont des intérêts communs. La définition d'un VLAN a pendant longtemps été un domaine de diffusion : la trame émise par l'un des membres est automatiquement diffusée vers l'ensemble des autres membres du VLAN.

Les réseaux locaux virtuels ont pour but initial de permettre une configuration et une administration plus faciles de grands réseaux d'entreprise construits autour de nombreux ponts. Il existe pour cela plusieurs stratégies d'applications de réseaux virtuels.

La notion de VLAN introduit une segmentation des grands réseaux. Les utilisateurs sont regroupés suivant des critères à déterminer. Un logiciel d'administration doit être disponible pour la gestion des adresses et des commutateurs. Un VLAN peut être défini comme un domaine de *broadcast*, c'est-à-dire un domaine où l'adresse de diffusion atteint toutes les stations appartenant au VLAN. Les communications à l'intérieur d'un VLAN peuvent être sécurisées, et les communications entre deux VLAN distincts contrôlées.

Plusieurs types de VLAN ont été définis, suivant les regroupements des stations du système :

- Les VLAN de niveau physique, ou de niveau 1, regroupent les stations appartenant aux mêmes réseaux physiques ou à plusieurs réseaux physiques mais reliés par une gestion commune des adresses. La figure 34 présente un exemple de VLAN de niveau 1.
- Les VLAN de niveau liaison, ou plus exactement de niveau MAC, ou encore de niveau 2, ont des adresses MAC qui regroupent les stations appartenant au même VLAN. Elles peuvent se trouver dans des lieux géographiquement distants. La difficulté est de pouvoir réaliser une diffusion automatique sur l'ensemble des stations du VLAN.

L'adresse MAC (*Medium Access Control*) n'est pas autre chose que l'adresse sur 6 octets que nous avons décrite à propos de la trame Ethernet. Une station peut appartenir à plusieurs VLAN simultanément. La figure 35 illustre un VLAN de liaison, liaison étant l'ancien nom du niveau trame.

-

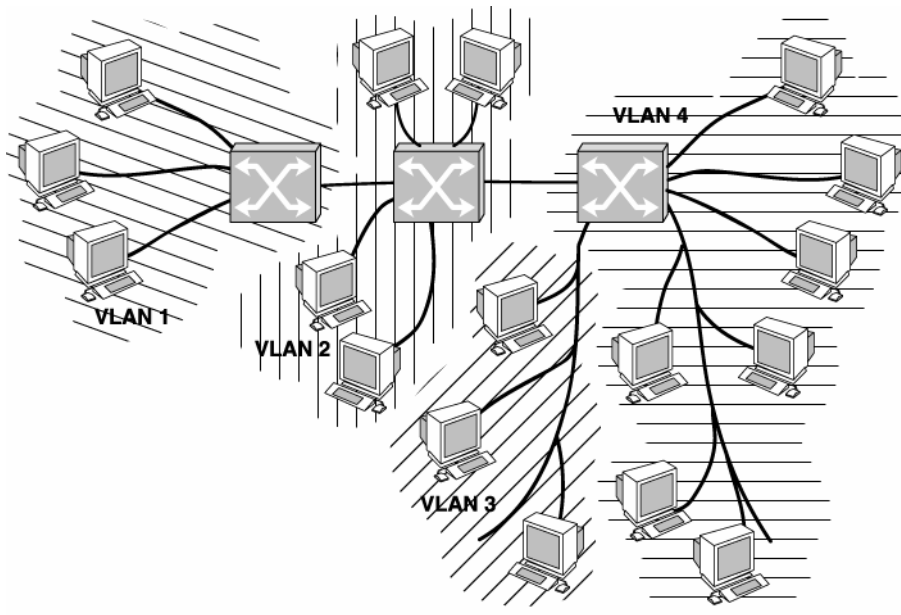


Figure 34. VLAN de niveau physique.

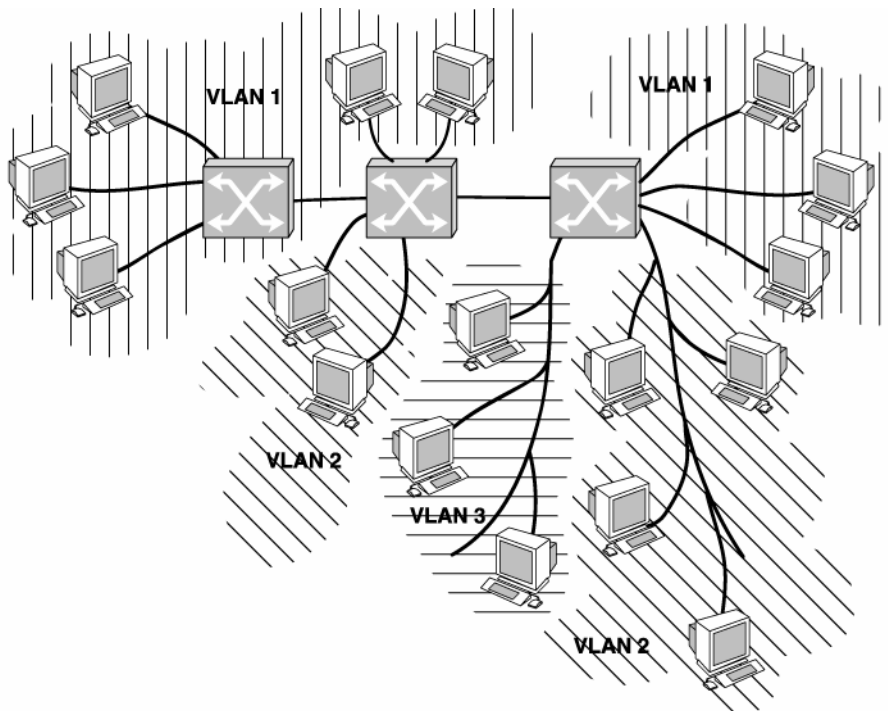


Figure 35. VLAN de niveau trame.

Les VLAN de niveau paquet, ou VLAN de niveau 3, correspondent à des regroupements de stations suivant leur adresse de niveau 3 (des adresses IP, par exemple). Il faut, dans ce cas, pouvoir faire correspondre facilement l'adresse de niveau paquet et celle de niveau trame. Ce sont les protocoles de type ARP (*Address Resolution Protocol*) qui effectuent cette correspondance d'adresses. Deux réseaux VLAN sont illustrés à la figure 36 ; la difficulté vient de la diffusion vers les seuls utilisateurs 1, 2 et 5 lorsqu'un membre du VLAN 1 émet et, de même, de la diffusion vers les seuls utilisateurs 3, 4, 6 et 7 lorsqu'un membre du VLAN 2 émet.

Lorsqu'un établissement de grande taille veut structurer son réseau, il peut créer des réseaux virtuels suivant des critères qui lui sont propres. Généralement, un critère géographique est retenu pour réaliser une communication simple entre les différents sites d'un établissement. L'adresse du VLAN doit être rajoutée dans la structure de la trame Ethernet (ou du paquet d'une autre technologie, puisque la structuration en VLAN ne concerne pas uniquement les environnements Ethernet).

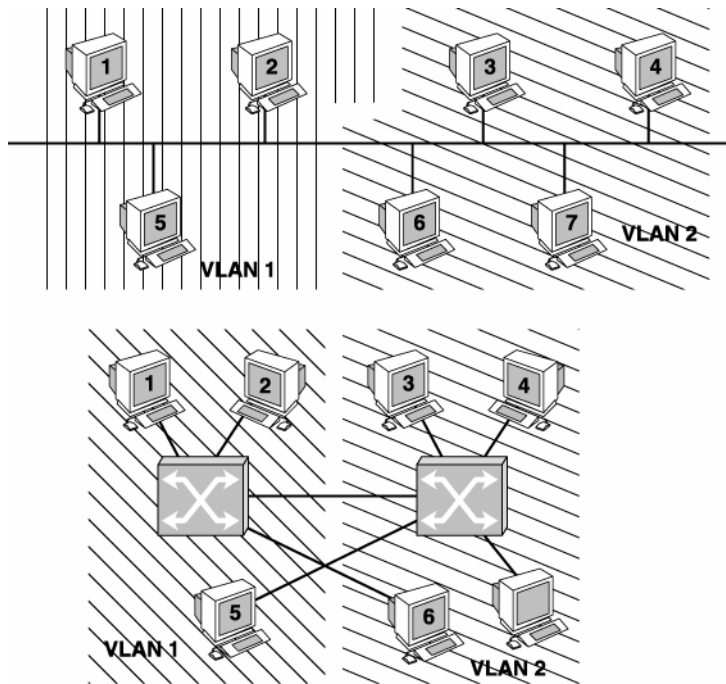


Figure 36. Deux topologies de VLAN

La norme VLAN Tagging

Cet aparté décrit la norme VLAN Tagging IEEE 802.1q, qui peut être utilisée suivant les différents schémas décrits précédemment. Le format de la trame Ethernet VLAN, défini dans les normes 802.3ac et 802.1q, est illustré à la figure 37.

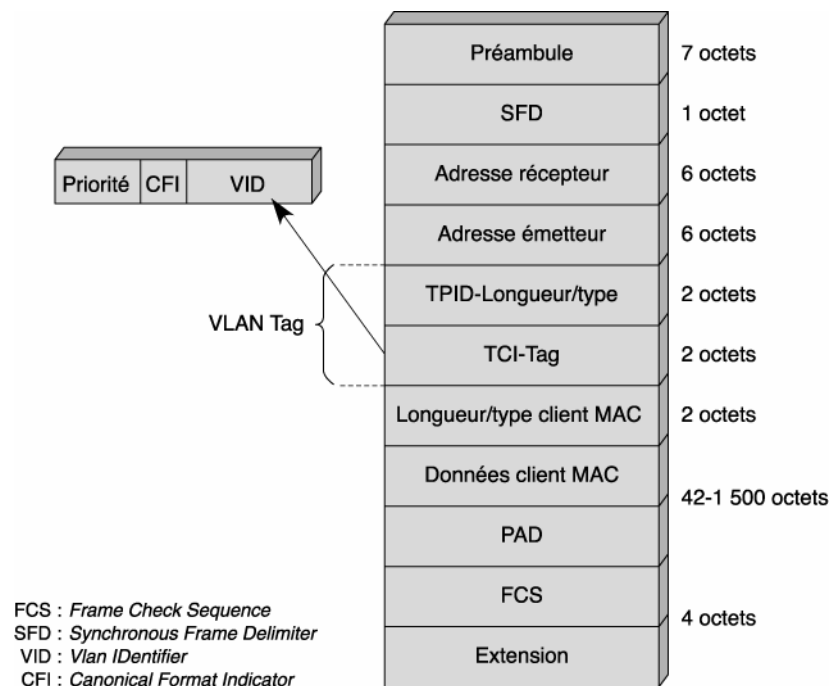


Figure 37. Format de la trame Ethernet VLAN.

L'identificateur VLAN (*VLAN Tag*) comporte 4 octets et prend en compte le champ longueur-type ainsi que le tag lui-même. Le VLAN Tag contient plus précisément un premier champ TPID (*VLAN Tag Protocol Identifier*) et un champ TCI (*Tag Control Information*). Il est inséré entre l'adresse source et le champ Longueur-type du client MAC. La longueur de la trame Ethernet passe à 1 522 octets (1 518 lorsque ce champ n'est pas présent). Le champ TPID prend la valeur 0x81-00, qui indique la présence du champ TCI.

Le TCI (*Tag Control Information*) contient lui-même les trois champs suivants :

- Un champ de priorité de 3 bits permettant jusqu'à 8 niveaux de priorité.

- Un champ d'un bit, le bit CFI (*Canonical Format Indicator*). Ce bit n'est pas utilisé pour les réseaux IEEE 802.3, et il doit être mis à 0 dans ce cas. On lui attribue la valeur 1 pour des encapsulations de trames Token-Ring.
- Un champ de 12 bits VID (*VLAN IDentifier*) indique l'adresse du VLAN.

Le rôle de l'élément priorité est primordial, car il permet d'affecter des priorités aux différentes applications multimédias. Cette fonctionnalité est définie dans la norme IEEE 802.1p.

À partir du moment où une commutation est mise en place, il faut ajouter un contrôle de flux, puisque les paquets Ethernet peuvent s'accumuler dans les nœuds de commutation. Ce contrôle de flux est effectué par le paquet Pause. C'est un contrôle de type *back-pressure*, dans lequel l'information de congestion remonte jusqu'à la source, nœud par nœud. À la différence des méthodes classiques, on indique au nœud amont une demande d'arrêt des émissions en lui précisant le temps pendant lequel il doit rester silencieux. Cette période peut être brève, si le nœud est peu congestionné, ou longue lorsque le problème est important. Le nœud amont peut lui-même estimer, suivant la longueur de la période de pause, s'il doit faire remonter un signal Pause ou non vers ses nœuds amont.

back-pressure.— Contrôle imposant une pression qui se propage vers la périphérie. Cette pression est exercée dans le cadre du contrôle de flux Ethernet par une commande Pause, qui demande au nœud amont de stopper ses transmissions pendant un laps de temps déterminé.

Comme on vient de le voir, Ethernet s'étend vers le domaine des WAN privés, en utilisant des techniques de commutation. Pour les réseaux locaux partagés, la tendance consiste plutôt à augmenter les débits, et ce grâce au Gigabit Ethernet et au 10 Gigabit Ethernet.

Le protocole MPLS

Le protocole MPLS (*MultiProtocol Label Switching*) a été choisi par l'IETF pour devenir le protocole d'interconnexion de tous les types d'architectures. Deux protocoles sous-jacents sont particulièrement mis en avant, ATM et Ethernet. Dans le cas d'Ethernet, une référence supplémentaire est ajoutée juste après l'adresse Ethernet sur 6 octets. Ce champ transporte la référence « shim », ou référence de « dérivation ». Cette dernière permet de faire transiter un paquet Ethernet d'un sous-réseau Ethernet à un autre sous-réseau Ethernet ou vers une autre architecture, ATM ou relais de trames.

La référence supplémentaire se place dans un troisième champ, situé entre l'adresse MAC et l'adresse VLAN.

Les réseaux d'accès

La boucle locale

La boucle locale, appelée également réseau de distribution, ou réseau d'accès, est l'une des parties les plus importantes pour un opérateur qui distribue de l'information à des utilisateurs. Elle constitue le capital de base de l'opérateur, en même temps que son lien direct avec le client.

Le coût global de mise en place et de maintenance d'un tel réseau est considérable. Il faut en général compter entre 500 et 3 000 euros par utilisateur pour installer le support physique entre le nœud de l'opérateur et la prise de l'utilisateur. Ce coût comprend l'infrastructure, le câble et les éléments extrémité de traitement du signal, mais il ne tient pas compte du terminal. Pour déterminer l'investissement de base d'un opérateur, il suffit de multiplier le coût d'installation d'une prise par la quantité d'utilisateurs raccordés. Le nombre de possibilités pour mettre à niveau un tel réseau à partir de l'existant est très important et continue à augmenter avec l'arrivée des techniques hertziennes sur la partie terminale, la plus proche de l'utilisateur.

La boucle locale correspond à la desserte de l'utilisateur : ce sont les derniers mètres ou kilomètres qui séparent le réseau du poste client. D'où le nom qu'on lui donne parfois de « dernier kilomètre », ou *last mile*. Les méthodes pour parcourir ce « dernier kilomètre » sont nombreuses et de type extrêmement varié. Pour les opérateurs historiques, c'est-à-dire ceux installés depuis longtemps et qui ont profité en général d'un monopole, la meilleure solution semble être l'utilisation d'un modem spécifique, permettant le passage de plusieurs mégabits par seconde

sur les paires métalliques de la boucle locale existante. La capacité dépend essentiellement de la distance entre l'équipement terminal et l'autocommutateur.

Comme on vient de le voir, la boucle locale correspond à la partie du réseau qui relie l'utilisateur au premier commutateur de l'opérateur. La valeur cible pour accéder au multimédia semble se situer aux alentours de 2 Mbit/s, ce qui est très inférieur aux prévisions effectuées il y a quelques années. Les progrès du codage et des techniques de compression sont à l'origine de cette nouvelle valeur. D'ici aux années 2005, une vidéo de qualité télévision devrait pouvoir être prise en charge avec un débit compris entre 64 Kbit/s et 512 Kbit/s. La parole sous forme numérique ne demande plus que quelques kilobits par seconde. Les compressions vont permettre, avec un débit de 1 à 2 Mbit/s, de rendre très acceptables les temps d'accès aux bases de données et de récupération des gros fichiers. Globalement, un débit de 2 Mbit/s devrait être suffisant pour assurer à un utilisateur un accès confortable aux informations multimédias.

La fibre optique

Une solution pour réaliser un réseau d'accès performant consiste à recâbler complètement le réseau de distribution. Le moyen le plus souvent évoqué pour cela est la fibre optique. Cette technique, qui donne de hauts débits jusqu'au terminal, est particulièrement bien adaptée au réseau numérique à intégration de services (RNIS) large bande. La boucle locale se présente sous la forme illustrée à la figure 38. Sa topologie est celle d'un *arbre optique passif*, ou PON (*Passive Optical Network*).

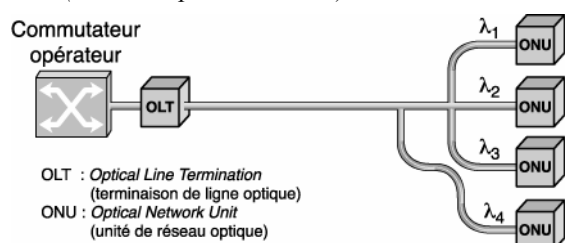


Figure 38. La boucle locale optique.

Cette solution est cependant assez onéreuse. Il est possible de réduire les coûts en ne câblant pas la portion allant jusqu'à la prise terminale de l'utilisateur. Il faut pour cela déterminer le point jusqu'où le câblage doit être posé. Les solutions offertes à l'opérateur sont les suivantes :

- Jusqu'à un point pas trop éloigné de l'immeuble ou de la maison qui doit être desservi, le reste du câblage étant effectué par l'utilisateur final (FTTC, *Fiber To The Curb*).
- Jusqu'à un répartiteur dans l'immeuble lui-même (FTTN, *Fiber To The Node*).
- Jusqu'à la porte de l'utilisateur (FTTH, *Fiber To The Home*).
- Jusqu'à la prise de l'utilisateur, à côté de son terminal (FTTT, *Fiber To The Terminal*).

Le prix augmentant fortement en fonction de la proximité avec l'utilisateur, la tendance actuelle consiste plutôt à câbler en fibre optique jusqu'à des points de desserte répartis dans le quartier et à choisir d'autres solutions moins onéreuses pour atteindre l'utilisateur. Avec l'aide de modems xDSL, le câblage métallique est capable de prendre en charge des débits de plusieurs mégabits par seconde sur les derniers kilomètres. La solution consiste donc à câbler en fibre optique jusqu'à un point situé à moins de 5 km de l'utilisateur. En ville, cette distance est très facile à respecter, mais hors des agglomérations, d'autres solutions doivent être recherchées.

Les réseaux optiques passifs

Sur un réseau optique passif (PON), il est possible de faire transiter des cellules ATM suivant la technique dite FSAN (*Full Service Access Network*). Les deux extrémités de l'arbre optique s'appellent OLT (*Optical Line Termination*) et ONU (*Optical Network Unit*). En raison de la déperdition d'énergie, il n'est pas possible de dépasser une cinquantaine de branches sur le tronc. La figure 39 illustre l'architecture d'un réseau optique passif.

Un superPON a également été défini, connectant jusqu'à 2 048 ONU sur un même OLT. Dans ce cas, le débit montant est de 2,5 Gbit/s.

Sur ces réseaux d'accès en fibre optique mis en place par les opérateurs, c'est le protocole ATM qui est en général retenu. Le système prend alors le nom de APON (*ATM Over PON*). La difficulté, avec les boucles passives optiques, comme celle de l'accès CATV, que nous examinerons ultérieurement, vient du partage de la bande passante montante, c'est-à-dire depuis l'utilisateur vers le réseau. En effet, si plusieurs centaines de clients se connectent simultanément, voire plusieurs milliers dans le cas des superPON (jusqu'à 20 000), la bande passante peut ne pas être suffisante. Sur la partie descendante, les canaux de vidéo sont diffusés. Ils utilisent chacun un canal sur le tronc de l'arbre. En revanche, les

canaux montants des utilisateurs sont tous différents et utilisent chacun un canal distinct . S'il y a 1 000 utilisateurs à la périphérie, 1 000 canaux séparés doivent atteindre la racine du câblage. Une technique d'accès MAC (*Medium Access Control*) est nécessaire pour prendre en charge cette superposition. Le multiplexage en longueur d'onde offre une solution simple, dans laquelle chaque utilisateur possède une longueur d'onde différente des autres utilisateurs. Cette solution ne peut cependant convenir que si le nombre de branches est limité. C'est pourquoi il est en général nécessaire de recourir à une technique de partage. De très nombreuses solutions permettent à l'ONU de faire une requête vers l'OLT, cette dernière réservant une bande passante aux clients demandeurs.

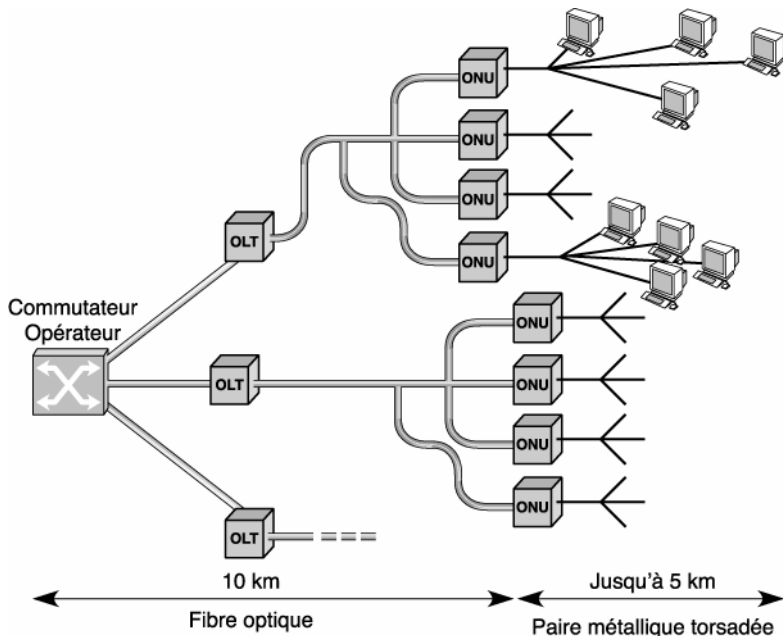


Figure 39. Architecture d'un PON.

Dans le sens descendant, les cellules ATM sont émises de façon classique sur le support physique. Dans le sens montant, une réservation est nécessaire. Elle s'effectue à l'intérieur de trames, divisées en tranches de 56 octets comportant une cellule et 3 octets de supervision. Au centre de la trame, une tranche particulière de 56 octets est destinée à la réservation d'une tranche de temps. La figure 40 illustre ces différentes zones de données.

Figure 40. Structure de la trame FSN.

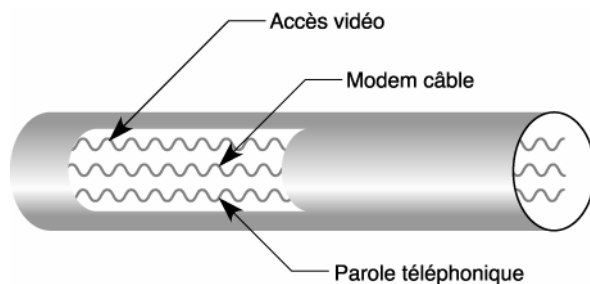
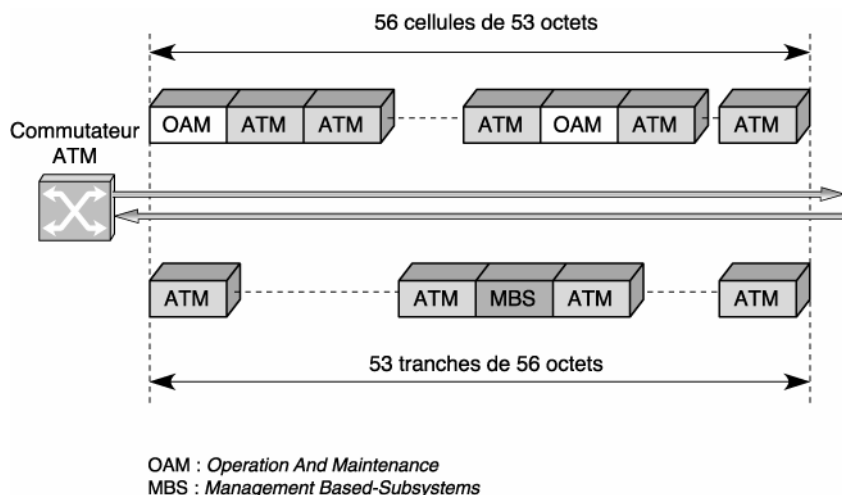


Figure 41. Multiplexage en fréquence dans un CATV.

Les réseaux câblés

Une autre solution pour obtenir un réseau d'accès à haut débit consiste à utiliser l'infrastructure des câblo-opérateurs, lorsqu'elle existe. Ce câblage a pendant longtemps été constitué de CATV (câble TV), dont la bande passante dépasse facilement les 800 MHz. Aujourd'hui, cette infrastructure est légèrement modifiée par la mise en place de systèmes HFC (*Hybrid Fiber/Coax*), qui associent la fibre optique jusqu'à la tête de retransmission et le CATV pour la desserte terminale, cette dernière pouvant représenter plusieurs kilomètres.

La technologie utilisée sur le CATV est de type multiplexage en fréquence : sur la bande passante globale, une division en sous-canaux indépendants les uns des autres est réalisée, comme illustré à la figure 41.



Cette solution présente de nombreux avantages, mais aussi quelques gros défauts. Son avantage principal réside dans la possibilité d'optimiser ce qui est transmis dans les canaux, puisque le contenu de chaque canal est indépendant de celui des autres. Dans ces conditions, le multimédia est facile à transporter. Il suffit d'affecter un média par *sous-bande*, chaque sous-bande ayant la possibilité d'être optimisée et de transporter les informations soit en analogique, soit en numérique.

Les canaux de télévision transitant dans des sous-bandes distinctes, rien n'empêche d'en avoir certains en numérique et d'autres en analogique. Une connexion de parole téléphonique peut être mise en place par une autre sous-bande. L'accès à un réseau Internet peut aussi être pris en charge par ce système, à condition d'utiliser un *modem câble*.

La faiblesse de cette solution provient du multiplexage en fréquence, qui n'utilise pas au mieux la bande passante et ne permet pas réellement une intégration des différents services qui transitent dans le CATV. Un multiplexage temporel apporte une meilleure utilisation de la bande passante disponible et intègre dans un même composant l'accès à l'ensemble des informations au point d'accès. Un transfert de paquets pourrait représenter une solution mieux adaptée, à condition de modifier complètement les composants extrémité. Cette dernière solution pourrait être utilisée sur les accès informatiques ou télécoms.

En résumé, il est possible d'acheminer une application multimédia sur le câble coaxial des câblo-opérateurs, mais avec le défaut de transporter les médias sur des sous-bandes en parallèle et non sur une bande unique.

La technologie HFC (*Hybrid Fiber/Coax*) se propose d'utiliser la fibre optique pour transporter des communications à haut débit jusqu'à une distance peu éloignée de l'utilisateur, en étant relayé par du câble coaxial jusqu'à la prise utilisateur. Par son énorme capacité, la fibre optique peut véhiculer autant de canaux que d'utilisateurs à atteindre, ce dont est incapable le câble CATV dès que le nombre d'utilisateurs devient important. Pour la partie câble coaxial, il faut trouver une solution de multiplexage des voies montantes vers la *tête de réseau*, de façon à faire transiter l'ensemble des demandes des utilisateurs vers le réseau. Cette solution est illustrée à la figure 42.

Prenons un exemple : si 10 000 prises doivent se connecter sur un arbre CATV, il est possible d'obtenir environ 4 Mbit/s sur la *voie descendante* et 80 Kbit/s sur la *voie montante* en respectant la division actuelle de la bande passante. Les vitesses sur la voie montante peuvent être considérées comme insuffisantes, mais il est possible, dans ce cas, d'utiliser plus efficacement la bande passante par un multiplexage permettant de récupérer les canaux inactifs.

Le norme IEEE 802.14

La transmission numérique sur un CATV s'effectue d'une manière unidirectionnelle, depuis la station terminale vers la tête de ligne et *vice versa*. La bande passante du CATV est subdivisée en une bande montante vers la tête de ligne et une bande descendante vers les équipements terminaux.

Comme cette partie du câblage peut desservir entre 500 et 2 000 utilisateurs depuis la tête de ligne, si chacun veut effectuer une application de télévision à la demande (VoD, pour *Video on Demand*), la bande passante n'est pas suffisante ou, du moins, chaque client doit-il se limiter à une partie de la bande passante. Pour permettre une meilleure adéquation de la bande passante, notamment aux applications interactives, le groupe de travail IEEE 802.14 a développé une proposition de partage de ce support par une technique dite MLAP (*MAC Level Access Protocol*), qui permet de distribuer le support entre les machines connectées.

La difficulté principale de ce système est de trouver une technique d'accès semblable à celle des réseaux locaux. Comme les tuyaux sont unidirectionnels, l'équipement le plus en aval ne peut écouter les émissions des autres stations qu'après un

long laps de temps, qui correspond à la propagation jusqu'à la tête de ligne et à celle en retour jusqu'à la station. Comme la portée du CATV peut atteindre plusieurs dizaines de kilomètres, il faut trouver une solution intermédiaire entre les techniques satellite et les méthodes utilisées dans les réseaux locaux.

Le protocole MLAP repose sur une succession d'actions découpées en cinq phases. Dans la première phase, la station que l'on examine est inactive. Dans la deuxième, elle devient active, c'est-à-dire qu'elle veut émettre des trames ; pour cela, elle avertit la tête de ligne par des primitives (UP.Frame et UP.REQ) et par un mécanisme d'accès aléatoire. À cet effet, la tête de ligne notifie à toutes les stations, par le biais du canal aval, les intervalles de temps pendant lesquels les stations peuvent émettre. Les canaux sont cependant utilisés dans un mode avec contention. L'allocation d'un canal ne se fait pas de façon unique du premier coup, et des collisions peuvent se produire. Associé à la tête de ligne, un contrôleur décide de modifier l'allocation des canaux en tenant compte de la demande de qualité de service des stations. L'algorithme est réinitialisé et les informations mises à jour ; une nouvelle allocation est alors déterminée. Les stations reçoivent une notification de la tête de ligne indiquant les nouveaux intervalles de temps qui leur sont alloués. Ce processus se poursuit jusqu'à ce que les stations aient leur canal réservé. Si une station modifie sa demande de bande passante ou de qualité de service, la nouvelle demande s'effectue par les canaux partagés : c'est la phase 5.

L'algorithme d'allocation des bandes passantes du contrôleur ne fait pas partie de la norme.

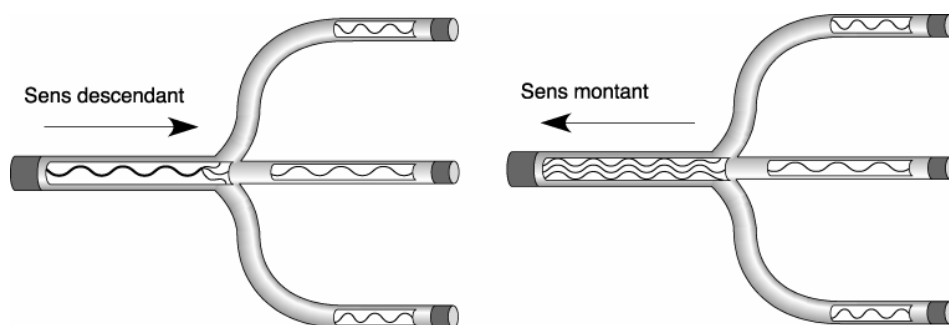


Figure 42. Le problème du multiplexage dans la boucle locale en CATV.

Les paires métalliques

Les *paires métalliques* forment l'ossature la plus classique de la boucle locale, principalement pour l'accès au réseau téléphonique. Lorsque l'accès se fait en analogique, ce qui est encore le plus souvent le cas, on peut utiliser une paire en full-duplex. Il est évidemment possible d'émettre des données binaires en utilisant un modem ; la vitesse peut alors atteindre quelques dizaines de kilobits par seconde.

La paire métallique peut devenir une liaison spécialisée si des répéteurs *ad hoc* sont placés à distance régulière. On atteint en général 2 Mbit/s ou 1,5 Mbit/s (*liaison T1*). Enfin, une paire métallique permet de mettre en place l'interface RNIS bande étroite ($2B + D_{16}$), dont la capacité, pour l'utilisateur, est de 144 Kbit/s. Il peut parfois s'agir de deux paires, voire de quatre.

La révolution est venue de nouveaux modems extrêmement puissants, les modems xDSL (*Digital Subscriber Line*), capables de véhiculer plusieurs mégabits par seconde. Ces modems permettent d'utiliser les paires métalliques du réseau d'accès pour réaliser une boucle locale à haut débit. Le débit visé est du même ordre de grandeur que celui des liaisons spécialisées à 2 Mbit/s.

Les modems ADSL (*Asymmetric Digital Subscriber Line*) sont les plus répandus. Leur vitesse est dissymétrique, c'est-à-dire plus lente entre le terminal et le réseau que dans l'autre sens. Les vitesses annoncées culminent, dans le sens équipement terminal vers réseau, à 250, 500 ou 750 Kbit/s. Dans l'autre sens, elles peuvent atteindre approximativement :

- 1,5 Mbit/s pour 6 km ;
- 2 Mbit/s pour 5 km ;
- 6 Mbit/s pour 4 km ;
- 9 Mbit/s pour 3 km ;
- 13 Mbit/s pour 1,5 km ;
- 26 Mbit/s pour 1 km ;
- 52 Mbit/s pour 300 m.

La technologie classique actuelle permet, sur la plupart des installations d'abonnés, d'émettre à la vitesse de 0,64 Mbit/s vers le réseau et de recevoir à 6 Mbit/s.

Un modem ADSL utilise une modulation d'amplitude quadratique, c'est-à-dire que 16 bits sont transportés à chaque signal. Avec une rapidité de modulation de 340 KB (kilobauds) et une atténuation de l'ordre d'une trentaine de décibels, on atteint plus de 5 Mbit/s.

Le succès de cette solution a entraîné l'apparition de nombreux dérivés. En particulier, la possibilité de faire varier le débit sur le câble a donné naissance à la variante RADSL (*Rate Adaptive DSL*). Pour les hauts débits, les variantes HDSL (*High bit rate DSL*) et VDSL (*Very high bit rate DSL*) peuvent être exploitées avec succès si le câblage, souvent en fibre optique, le permet.

Le téléphone et la télévision sur ADSL

Le téléphone sur les accès xDSL s'appelle VoDSL (*Voice over DSL*). Deux solutions peuvent être mises en œuvre :

- Dans le cas d'un opérateur historique, ou ILEC (*Incumbent Local Exchange Carrier*), on utilise la bande classique de la parole analogique pour y faire passer plusieurs canaux numériques compressés.
- Dans le cas d'un opérateur entrant, ou CLERC (*Competitive Local Exchange Carrier*), la parole téléphonique est mise en paquets, lesquels transitent sur la bande passante du modem xDSL.

La télévision transite quant à elle uniquement sur la bande transportant des paquets IP en utilisant le protocole IP Multicast pour permettre un acheminement optimisé. La compression est effectuée par des protocoles de type MPEG-2 ou MPEG-4.

Les accès hertziens

L'utilisation de la voie hertzienne représente une autre solution prometteuse à long terme pour la boucle locale. Dès maintenant, cette solution est envisageable pour les communications téléphoniques, surtout si le câblage terrestre n'existe pas. Lors de la construction d'une ville nouvelle, par exemple, la desserte téléphonique peut s'effectuer par ce biais. L'inconvénient réside dans l'étroitesse de la bande passante disponible dans le spectre des fréquences.

Cette solution, appelée WLL (*Wireless Local Loop*) ou WITL (*Wireless In The Loop*), est en plein essor. Deux grandes orientations peuvent être adoptées, suivant que l'on souhaite une mobilité du client ou non. Dans le premier cas, la communication doit continuer sans interruption tandis que le mobile se déplace. On parle alors de réseau de mobiles. Dans le second cas, la communication est fixe ou possède une mobilité réduite. Plus les fréquences utilisées sont hautes, et plus la directivité est importante, limitant la mobilité. Cette dernière solution, appelée BLR (*Boucle Locale Radio*) ou WLL (*Wireless Local Loop*) ou encore WITL (*Wireless In The Loop*), est en plein essor.

La présente section se concentre sur les solutions à mobilité restreinte, comme LMDS (*Local Multipoint Distribution System*) et les constellations de satellites.

Les accès hertziens peuvent aussi être utilisés pour des terminaux demandant des débits plus importants que les mobiles classiques. Le DECT (*Digital Enhanced Cordless Terminal*) se présente comme une solution potentielle, mais au prix d'une limitation de la mobilité du terminal, qui doit rester dans la même cellule. Cette norme ETSI de 1992 utilise une technique de division temporelle TDMA (*Time Division Multiple Access*), permettant de mettre en place des interfaces RNIS bande étroite.

La solution LMDS (*Local Multipoint Distribution System*) utilise des fréquences très élevées dans la bande Ka, c'est-à-dire au-dessus de 25 GHz. À de telles fréquences, les communications sont très directives. Comme il est difficile d'obtenir un axe direct entre l'antenne et le terminal, il faut placer les antennes dans des lieux élevés. Par ailleurs, de très fortes pluies peuvent légèrement perturber la propagation des ondes. L'avantage du LMDS est évidemment d'offrir d'importantes largeurs de bande, permettant des débits de 20 Mbit/s par utilisateur. En 1997, la FCC a alloué 1 300 MHz au service LMDS dans les bandes de fréquence 28 GHz et 31 GHz. Sa portée s'étend jusqu'à une dizaine de kilomètres.

Les accès satellite

Le réseau d'accès peut aussi utiliser les techniques de distribution directe par satellite. Les très grands projets qui ont été finalisés dans un premier temps ne visent que la téléphonie, par suite d'un manque flagrant de bande passante. Ce défaut est cependant partiellement compensé par le grand nombre de satellites défilant à basse altitude, qui permet une réutilisation partielle des fréquences. Plusieurs grands projets ont abouti techniquement, mais certains se sont effondrés financièrement. Deux possibilités bien différentes se font jour : soit le réseau satellite ne couvre que le réseau d'accès, soit il propose également le transport, effectué dans le premier cas par un opérateur terrestre. La constellation Global Star est un pur réseau d'accès, puisqu'il ne fait que diriger les communications vers des portes d'accès d'opérateurs terrestres. En revanche, Iridium (qui a fait faillite) permettait le routage de

satellite en satellite pour arriver jusqu'au destinataire ou pour accéder, sur la dernière partie de la communication, au réseau d'un opérateur.

La première génération de constellations de satellites ne concerne, comme on vient de le voir, que la téléphonie, en raison d'un manque important de bande passante. La deuxième génération visera le multimédia, avec des débits allant jusqu'à 2 Mbit/s et des qualités de service associées. Ces constellations seront des LEOS, des MEOS et des GEOS (*Low, Medium et Geostationary Earth Orbital Satellite*). Les plus connues sont Teledesic, qui regroupe derrière Bill Gates, Motorola, Boeing, etc., et SkyBridge, d'Alcatel.

LEOS, MEOS et GEOS (*Low, Medium et Geostationary Earth Orbital Satellite*) sont des satellites situés approximativement à 1 000, 13 000 et 36 000 km de la Terre. Les deux premières catégories concernent les satellites défilants, et la dernière les satellites qui semblent fixes par rapport à la Terre.

La figure 43 illustre une constellation basse orbite.



Figure 43. Une constellation de satellites.

Les satellites de télécommunications de la première génération sont géostationnaires, c'est-à-dire qu'ils décrivent une orbite circulaire autour de la Terre dans un plan voisin de l'équateur, avec une vitesse angulaire égale à celle de la rotation de la Terre sur elle-même (*voir figure 44*). Ils apparaissent ainsi comme sensiblement immobiles pour un observateur terrien, ce qui permet une exploitation permanente du satellite.

L'orbite d'un satellite géostationnaire se situe à 36 000 km de la Terre, ce qui confère à un signal un trajet aller-retour d'approximativement 0,27 s. Ce délai de propagation très grand a un impact important sur les techniques d'accès au canal satellite. À cette altitude, parmi trois satellites placés à 120° les uns des autres sur l'orbite géostationnaire, au moins l'un d'entre eux est visible d'un point quelconque de la Terre.

Le signal reçu par le satellite à une fréquence f_1 est retransmis à une fréquence f_2 vers l'ensemble des stations terrestres. Il se produit ainsi une diffusion des signaux. Ces deux propriétés — toutes les stations écoutent toutes les transmissions, et toutes les stations peuvent émettre — permettent d'implanter des schémas de contrôle spécifiques. Il faut noter une certaine ressemblance de ce système avec les réseaux partagés, qui possèdent généralement l'accès multiple et la diffusion. La véritable différence provient du délai de propagation, qui n'est pas du tout du même ordre de grandeur. Nous allons cependant trouver des variantes de la technique Ethernet, mais avec des caractéristiques différentes.

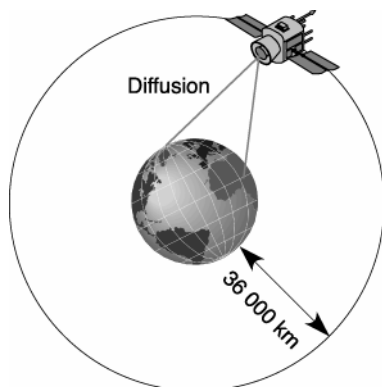


Figure 44. Un satellite géostationnaire.

Les limites à l'utilisation des satellites géostationnaires sont bien connues. Citons, notamment, les suivantes :

- La puissance d'émission des terminaux et du satellite doit être importante, à moins que l'antenne ne soit d'un diamètre important. Un terminal ayant une antenne de 3 dBW (affaiblissement de 3 décibels par watt) exige du satellite une antenne de 10 m de diamètre.
- La puissance demandée au satellite étant importante, ce dernier doit disposer de batteries importantes et donc de capteurs de grande surface.
- La couverture totale de la Terre est impossible, les zones situées au-dessus de 81° de latitude n'étant pas couvertes et celles situées au-dessus de 75° de latitude devant faire face à une capacité de communication fortement réduite.
- Comme il est difficile de réutiliser des fréquences, la capacité globale est très faible en comparaison des réseaux terrestres.
- Plus l'angle d'inclinaison est grand, et plus la trajectoire des ondes est perturbée par les obstacles, ce qui rend les communications très difficiles à partir de 50° de latitude.
- La communication de mobile à mobile entre deux stations qui ne sont pas situées dans la même zone de couverture requiert le passage par un réseau terrestre, les communications entre satellites stationnaires étant particulièrement complexes.

La conséquence de toutes ces limites est simple : le satellite lui-même est extrêmement lourd et demande une puissante fusée pour être mis sur orbite. En revanche, les satellites à basse orbite, plus légers, peuvent être lancés par de petites fusées ou par grappes de 6 à 10.

Les nouvelles générations de satellites ne sont d'ailleurs plus géostationnaires. Elles peuvent de la sorte profiter de la réutilisation potentielle des fréquences par une altitude beaucoup plus basse, ce qui permet d'obtenir de petites cellules. Parmi les nombreux avantages de cette évolution, citons la réutilisation d'environ 20 000 fois la même fréquence lorsque les satellites se situent à 1 000 km de la Terre et un coût de lancement et une puissance d'émission des signaux bien moindres étant donné la proximité de la Terre. Son inconvénient principal vient bien sûr du déplacement du satellite, puisque celui-ci n'est plus stationnaire. Les communications d'un utilisateur terrestre doivent donc régulièrement changer de satellite, ce que l'on nomme un *handover*, comme dans les réseaux de mobiles. Ces constellations sont décrites plus en détail à la prochaine section.

Deux catégories de constellations de satellites à basse orbite se font jour, qui autorisent une bonne réutilisation des fréquences et permettent de mettre en place des réseaux universels : celles qui jouent le rôle de réseau d'accès et celles qui font office de réseau complet, gérant des communications de bout en bout en utilisant éventuellement des liaisons intersatellite.

Suivant la trajectoire des satellites, l'inclinaison de l'orbite et l'orbite elle-même, on peut réaliser des réseaux spécifiques pour privilégier des régions particulières. La trajectoire correspond à la forme de l'orbite, qui peut être circulaire ou en ellipse. Dans le cas d'une ellipse, la vitesse relative par rapport à la Terre varie, et le satellite peut se positionner pour rester plus longtemps au-dessus de certaines zones. L'inclinaison définit l'orbite par rapport au plan équatorial. Des zones particulières peuvent être privilégiées suivant l'inclinaison. Les orbites polaires imposent aux satellites de passer au-dessus des pôles, les régions les plus proches du pôle étant mieux desservies que les régions équatoriales. Les orbites en rosette sont assez fortement inclinées par rapport au plan équatorial. Dans ce cas, les régions intermédiaires entre le pôle et l'équateur sont les mieux couvertes. Enfin, les orbites équatoriales favorisent bien sûr les pays situés autour de l'équateur.

Les fréquences radio sont divisées en bandes déterminées par l'IEEE (*Standard Radar Definitions*) sous forme de lettres. Les numéros de bandes et les noms sont donnés par l'organisme international de régulation des bandes de fréquences. La figure 45 illustre ces bandes.

La bande C est la première à avoir été utilisée pour les applications commerciales. La bande Ku accepte des antennes beaucoup plus petites (VSAT), de 45 cm de diamètre seulement. La bande Ka autorise des antennes encore plus petites, et c'est pourquoi la plupart des constellations de satellites l'utilisent. De ce fait, les terminaux peuvent devenir mobiles, grâce à une antenne presque aussi petite que celle des terminaux de type GSM. On qualifie ces terminaux de USAT (*Ultra Small Aperture Terminal*). En revanche, l'utilisation de la bande S permet d'entrer dans le cadre de l'UMTS et des réseaux de mobiles terrestres.

Les fréquences classiquement utilisées pour la transmission par satellite concernent les bandes 4-6 GHz, 11-14 GHz et 20-30 GHz. Les bandes passantes vont jusqu'à 500 MHz et parfois 3 500 MHz. Elles permettent des débits très élevés : jusqu'à plusieurs dizaines de mégabits par seconde. Un satellite comprend des répéteurs, de 5 à 50 actuellement. Chaque répéteur est accordé sur une fréquence différente. Par exemple, pour la bande des 4-6 GHz, il reçoit des signaux modulés dans la bande des 6 GHz, les amplifie et les transpose dans la bande des 4 GHz. S'il existe n stations terrestres à raccorder par le canal satellite, le nombre de liaisons bipoints est égal à $n \times (n - 1)$. Ce

nombre est toujours supérieur à celui des répéteurs. Il faut donc, là aussi, avoir des politiques d'allocation des bandes de fréquences et des répéteurs.

Un exemple fera comprendre le problème posé par l'accès à un canal satellite. Supposons que les stations terrestres n'aient à leur disposition qu'une seule fréquence f_1 , transposée en la fréquence f_2 de retour. Comme les stations terrestres n'ont de relation entre elles que *via* le satellite, si une station veut émettre un signal, elle ne peut le faire qu'indépendamment des autres stations, s'il n'existe pas de politique commune. Si une autre station émet dans le même temps, les signaux entrent en collision et deviennent incompréhensibles, puisque impossibles à décoder. Les deux messages sont alors perdus, et il faut les retransmettre ultérieurement.

Numéro	Bande	Symbole	Fréquence
12		Ondes sous-millimétriques	300-3 000 GHz
		Ondes millimétriques	40-300 GHz
		Bande Ka	27-40 GHz
11	EHF		30-300 GHz
		Bande K	18-27 GHz
		Bande Ku	12-18 GHz
		Bande X	8-12 GHz
		Bande C	4-8 GHz
10	SHF		3-30 GHz
		Bande S	2-4 GHz
		Bande L	1-2 GHz
9	UHF		300 MHz-3 GHz
8	VHF		30-300 MHz
7	HF		3-30 MHz
6	MF		300 KHz-3 MHz
5	LF		30-300 KHz
4	VLF		3-30 KHz

Figure 45. Les fréquences radio.

Les politiques d'accès aux canaux satellite

Les politiques d'accès aux canaux satellite doivent favoriser une utilisation maximale du canal, celui-ci étant la ressource fondamentale du système. Dans les réseaux locaux, le délai de propagation très court permet d'arrêter les transmissions après un temps négligeable. Dans le cas de satellites géostationnaires, les stations terrestres ne découvrent qu'il y a eu chevauchement des signaux que 0,27 s après leur émission — elles peuvent s'écouter grâce à la propriété de diffusion —, ce qui représente une perte importante sur un canal d'une capacité de plusieurs mégabits par seconde.

Les disciplines d'accès sont généralement classées en quatre catégories, mais d'autres possibilités existent, notamment les suivantes :

- les méthodes de réservation fixe, ou FAMA (*Fixed-Assignment Multiple Access*) ;
- les méthodes d'accès aléatoires, ou RA (*Random Access*) ;
- les méthodes de réservation par paquet, ou PR (*Packet Reservation*) ;
- les méthodes de réservation dynamique, ou DAMA (*Demand Assignment Multiple Access*).

Les protocoles de réservation fixe réalisent des accès non dynamiques aux ressources et ne dépendent donc pas de l'activité des stations. Les procédures FDMA, TDMA et CDMA forment les principales techniques de cette catégorie. Ces solutions offrent une qualité de service garantie puisque les ressources sont affectées une fois pour toutes. En revanche, l'utilisation des ressources est mauvaise, comme dans le cas d'un circuit affecté au transport de paquets. Lorsque le flux est variable, les ressources doivent permettre le passage du débit crête.

Les techniques d'accès aléatoires donnent aux utilisateurs la possibilité de transmettre leurs données dans un ordre sans corrélation. En revanche, ces techniques ne se prêtent à aucune qualité de service. Leur point fort réside dans une implémentation simple et un coût de mise en œuvre assez bas.

Les méthodes de réservation par paquet évitent les collisions par l'utilisation d'un schéma de réservation de niveau paquet. Comme les utilisateurs sont distribués dans l'espace, il doit exister un sous-canal de signalisation à même de mettre les utilisateurs en communication pour gérer la réservation.

Les méthodes dynamiques de réservation ont pour but d'optimiser l'utilisation du canal. Ces techniques essaient de multiplexer un maximum d'utilisateurs sur le même canal en demandant aux utilisateurs d'effectuer une réservation pour un

temps relativement court. Une fois la réservation acceptée, l'utilisateur vide ses mémoires tampons jusqu'à la fin de la réservation puis relâche le canal.

Les communications par l'intermédiaire d'un satellite montrent des propriétés légèrement différentes de celles d'un réseau terrestre. Les erreurs, en particulier, se produisent de façon fortement groupée, en raison de phénomènes physiques sur les antennes d'émission ou de réception. Au contraire des réseaux locaux, aucun protocole de niveau liaison n'est normalisé pour les réseaux satellite. Plusieurs procédures ont été proposées, mais aucune ne fait l'unanimité. Le délai d'accès au satellite constitue le problème principal, puisque, pour recevoir un acquittement, un temps égal à deux fois l'aller-retour est nécessaire. À ce délai aller-retour, il faut encore ajouter le passage dans les éléments extrémité, qui est loin d'être négligeable. Ce délai dépend bien sûr de la position de l'orbite sur laquelle se trouve le satellite. Lorsque les capacités des liaisons sont importantes, les techniques de retransmissions automatiques ARQ (*Automatic Repeat reQuest*) classiques ne sont pas efficaces, la quantité d'information à retransmettre devenant très grande. Les techniques sélectives posent aussi des questions de dimensionnement des mémoires permettant d'attendre l'information qui n'est pas parvenue au récepteur. Il faut trouver de nouvelles solutions pour optimiser le débit de la liaison.

Contrairement aux réseaux locaux, les réseaux satellite n'ont pas donné lieu à une normalisation spécifique. Plusieurs protocoles ont été proposés, mais aucun ne s'est vraiment imposé.

Un réseau utilisant un satellite géostationnaire ou un satellite situé sur une orbite moyenne se caractérise par un très long temps de propagation, comparativement au temps d'émission d'une trame. C'est pour cette raison qu'existe un mode étendu dans les procédures classiques, qui permet l'émission de 127 trames sans interruption. Il est aisé de comprendre que si le débit est très élevé et qu'une méthode SREJ soit adoptée, l'anticipation doit être énorme ou bien la longueur des trames très grande. Par exemple, si le débit de la liaison satellite est de 10 Mbit/s, sachant qu'il faut au moins prévoir d'émettre sans interruption pendant un temps égal à deux aller-retour (voir figure 46), la valeur minimale de la trame est de 20 Ko. Cette quantité est très importante, et la qualité de la ligne doit être excellente pour qu'un tel bloc de données (160 000 bits) arrive avec un taux d'erreur bit inférieur à 10^{-10} .

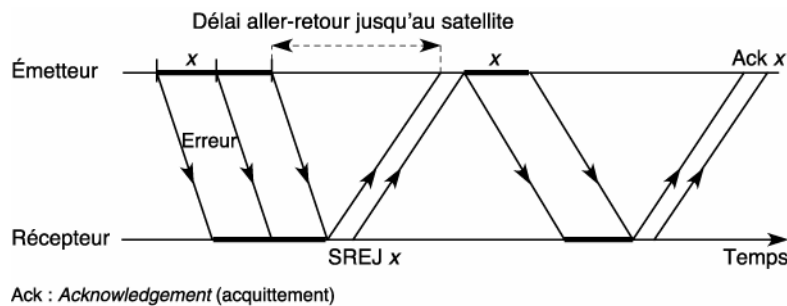


Figure 46. Reprise sur une liaison satellite.

Les systèmes satellite large bande

Les services que nous venons de décrire concernent les services à bande étroite et principalement la téléphonie. L'évolution actuelle pousse à se diriger vers des environnements permettant de transporter des applications multimédias.

Le développement des VSAT (*Very Small Aperture Terminal*) et maintenant des USAT (*Ultra Small Aperture Terminal*) est à l'origine de cette profusion d'antennes que l'on voit fleurir sur les toits et les balcons. L'utilisation du satellite s'étend pour aller des communications bande étroite au transport de canaux vidéo de très bonne qualité, en passant par les systèmes de communication *monovoies*, utilisant une seule direction pour les transmissions et les communications bidirectionnelles.

Le nombre de satellites en orbite pour la diffusion de canaux de télévision ne cesse de croître. De nombreux standards ont été créés, comme DSS (*Direct Satellite System*). La compétition avec les réseaux câblés est de plus en plus forte, ces derniers bénéficiant d'une plus grande bande passante, ce qui leur autorise des services à haut débit, comme la télévision à la demande.

Les constellations s'attaquent au problème réseau des deux façons différentes suivantes :

- Le système n'est qu'un réseau d'accès.
- Le système est un réseau complet en lui-même.

Dans le premier cas, l'utilisateur se connecte au satellite pour entrer sur le réseau d'un opérateur fixe. Il n'est pas possible de passer d'un satellite à un autre, et, à chaque satellite, doit correspondre une station au sol connectée à un

opérateur local. Dans la seconde solution, il peut être possible de passer directement de l'émetteur au récepteur, en allant de satellite en satellite. Quelques stations d'accès à des opérateurs fixes sont également possibles, mais surtout pour atteindre les clients fixes.

La configuration de la constellation est importante pour réaliser l'un ou l'autre type de réseaux. En particulier, il faut essayer de passer au-dessus des zones les plus demandées, et plusieurs choix d'orbites sont en compétition. Les liaisons intersatellite sont nécessaires lorsque le système n'est pas seulement un réseau d'accès. Le nombre de possibilités de liaisons intersatellite à partir d'un même satellite varie de deux à huit. La figure 47 illustre une liaison intersatellite et la figure 48 le délai de retour du satellite dans le processus d'acquiescement. Le fait de posséder des liaisons intersatellite réduit la dépendance de la constellation envers les opérateurs de réseaux fixes. En contrepartie, ils ajoutent au coût et à la complexité du système global. Plusieurs technologies sont disponibles pour la réalisation des liaisons intersatellite, comme la radio, l'infrarouge et le laser.

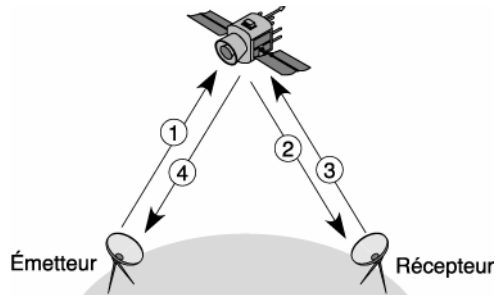


Figure 47. Exemple de liaison intersatellite.

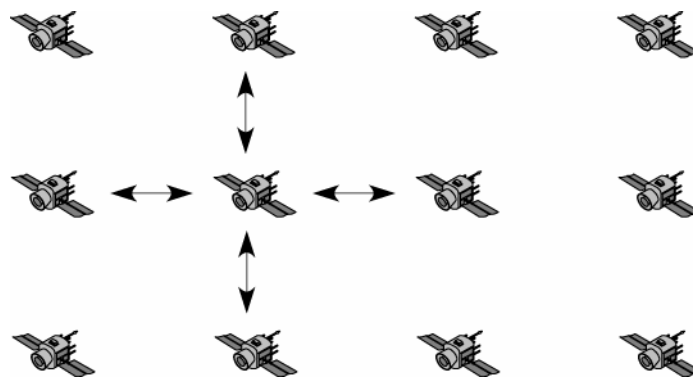


Figure 48. Délai de retour du satellite dans le processus d'acquiescement.

Le handover dans les constellations basse orbite

Le handover correspond à une modification de la cellule qui prend en charge le client, que ce soit à la suite d'un déplacement du client qui change de cellule ou de celui du satellite qui, en tournant, finit par perdre de vue son client. Dans ce dernier cas, un autre satellite de la constellation prend le relais pour que la communication ne soit pas interrompue.

Dans les systèmes qui nous intéressent ici, les constellations basse orbite, le satellite se déplace à la vitesse de 4 à 5 km/s, une vitesse bien supérieure à celle du client. Le nombre de handovers dus à un client qui se déplace peut être considéré comme négligeable et constitue en cela un avantage important. Comme la vitesse de défilement d'un satellite est constante et que le client est assimilé à un point fixe, les handovers sont prévisibles pour autant que le client ne termine pas brutalement sa communication.

En général, on distingue deux catégories de handovers :

- Le handover dur (hard-handover), dans lequel il ne doit y avoir aucune coupure et où le relais sur la nouvelle cellule commence juste au moment où se termine la communication avec la cellule précédente.
- Le handover mou (soft-handover), dans lequel des éléments de communication commencent à transiter par la nouvelle cellule tandis que la cellule précédente est toujours en cours de communication.

Un satellite peut gérer plusieurs cellules — jusqu'à une centaine — grâce à de nombreuses antennes d'émission. Un handover intrasatellite correspond à un handover qui s'effectue entre deux cellules gérées par le même satellite. En revanche, un handover intersatellite correspond à un client qui est pris en charge par une cellule dépendant d'un autre satellite. Le premier type de handover est assez simple, en ce sens qu'un seul et même satellite gère les deux cellules. Un handover intersatellite est nettement plus complexe, car il faut gérer la communication entre les deux satellites sans interruption.

On distingue deux grands types de handovers, en fonction du défilement du satellite :

- Les antennes sont fixes et les cellules se déplacent à la vitesse du satellite.
- Les cellules sont fixes, et le satellite pointe pendant un laps de temps assez long sur la même cellule géographique ; les antennes doivent donc être orientables dans ce cas.

En d'autres termes, dans le premier cas, l'antenne pointe vers la Terre suivant une position fixe. En revanche, dans le second cas, l'antenne pointe vers une zone terrestre fixe qui est suivie par le satellite. Lorsque le satellite disparaît (en général, en dessous d'un angle de 30°), la cellule fixe est prise en charge par le nouveau satellite qui arrive.

Les réseaux de mobiles

Les réseaux cellulaires

L'utilisation de la voie hertzienne pour le transport de l'information a donné naissance à des architectures de réseau assez différentes de celles des réseaux fixes. L'une des raisons à cela est que, dans ces réseaux de mobiles, la communication doit continuer sans interruption, même en cas de déplacement de l'émetteur ou du récepteur. Le réseau est donc constitué de *cellules*, qui recouvrent le territoire que l'opérateur souhaite desservir. Lorsqu'un mobile quitte une cellule, il doit entrer dans une autre cellule pour que la communication puisse continuer. L'un des aspects délicats de ces réseaux concerne le passage du mobile d'une cellule dans une autre. Ce changement de cellule s'appelle un *handover*, ou handoff. Un réseau cellulaire et un handover sont illustrés aux figures 49 et 50.

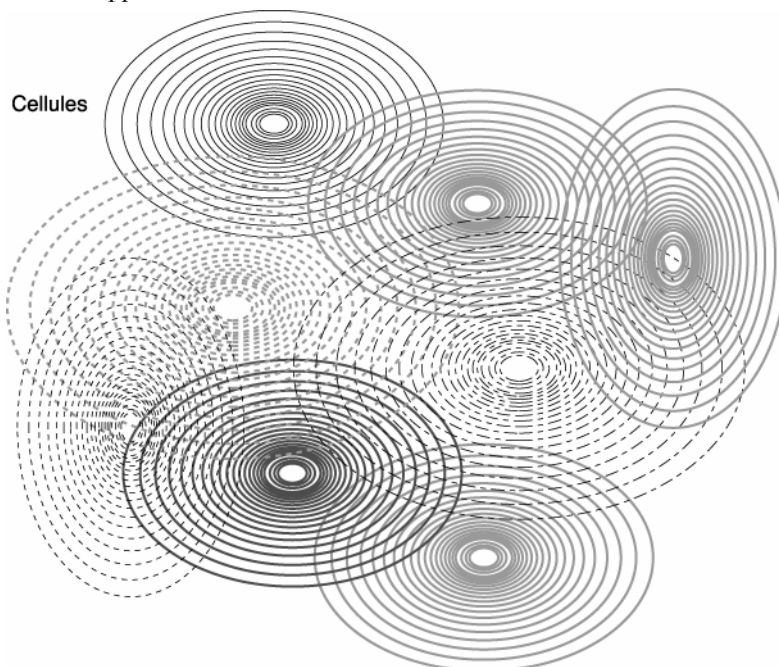


Figure 49. Un réseau cellulaire.

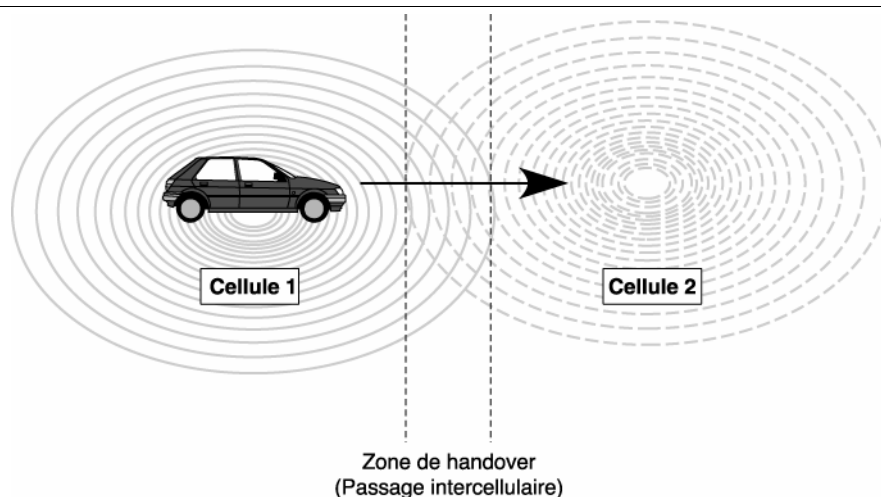


Figure 50. Un handover.

Les cellules se superposent partiellement de façon à assurer une couverture la plus complète possible du territoire cible. Le mobile communique par une *interface radio* avec l'antenne centrale, qui joue le rôle d'émetteur-récepteur de la cellule. Cette interface radio utilise en général des bandes de fréquences spécifiques du pays dans lequel est implanté le réseau. De ce fait, les interfaces radio ne sont pas toujours compatibles entre elles.

La première génération des réseaux de mobiles apparaît à la fin des années 70. Ce sont des réseaux cellulaires, et l'émission des informations sur l'interface radio s'effectue en analogique. Comme il n'existe pas alors de norme prépondérante pour l'interface radio analogique, un fractionnement des marchés se produit, qui empêche cette génération de réseaux de mobiles de rencontrer le succès escompté. Avec la deuxième génération, la transmission sur l'interface radio devient numérique. Enfin, la troisième génération prend en charge les applications multimédias et l'intégration à l'environnement Internet.

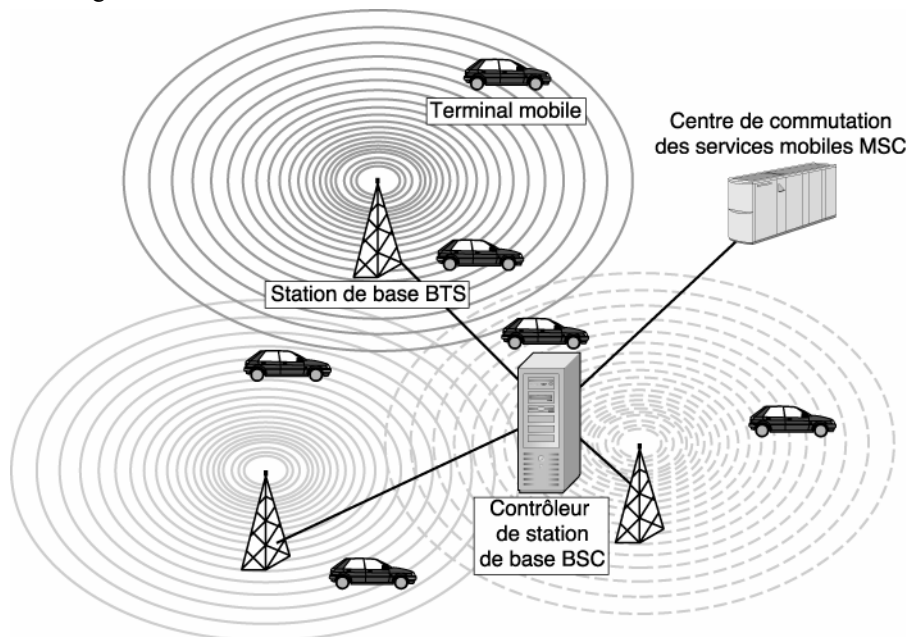


Figure 51. Architecture d'un réseau de mobiles.

Le territoire à desservir est subdivisé en cellules. Une cellule est liée à une station de base, ou *BTS* (*Base Transceiver Station*), qui possède l'antenne permettant d'émettre vers les mobiles, et *vice versa*. Si la densité du trafic est très forte sur une zone donnée, plusieurs stations de base peuvent couvrir cette zone, ce qui augmente la largeur de bande disponible. De même, plus le rayon de la cellule est petit, et plus la bande passante disponible s'adresse proportionnellement à un petit nombre d'utilisateurs. Le rayons d'une cellule peut descendre sous les 500 m.

Un sous-système radio rassemble ces stations de base, auxquelles sont ajoutés des contrôleurs de stations de base, ou *BSC* (*Base Station Controller*). Cet ensemble gère l'interface radio. Le travail des stations de base consiste à

prendre en charge les fonctions de transmission et de signalisation. Le contrôleur de station de base gère les *ressources radioélectriques* des stations de base qui dépendent de lui.

Le sous-système réseau contient les centres de commutation du service mobile, ou *MSC (Mobile service Switching Center)*, qui assurent l'interconnexion des stations de base, à la fois entre elles et avec les autres réseaux de télécommunications. Ces centres n'assurent pas la gestion des abonnés. Leur rôle est essentiellement la commutation, qui permet de relier, directement ou par le biais d'un réseau extérieur, les contrôleurs de stations de base. L'architecture d'un réseau de mobiles est illustrée à la figure 51.

Le sous-système réseau contient aussi deux bases de données, l'enregistreur de localisation nominal, ou *HLR (Home Location Register)* et l'enregistreur de localisation des visiteurs, ou *VLR (Visitor Location Register)*.

L'enregistreur de localisation nominal, HLR, gère les abonnés qui sont rattachés au MSC. L'enregistreur de localisation des visiteurs, VLR, a pour but de localiser les mobiles qui traversent la zone prise en charge par le MSC. Le premier enregistreur n'est pas dynamique, à la différence du second.

L'accès d'un utilisateur ne peut s'effectuer qu'au travers d'une *carte SIM (Subscriber Identity Module)*, qui comporte l'identité de l'abonné. Un code secret permet au réseau de vérifier que l'abonnement est valide. L'architecture de communication d'un mobile vers un autre s'exprime sous la forme d'interfaces à traverser. Les quatre interfaces définies dans les systèmes d'aujourd'hui sont illustrées à la figure 52.

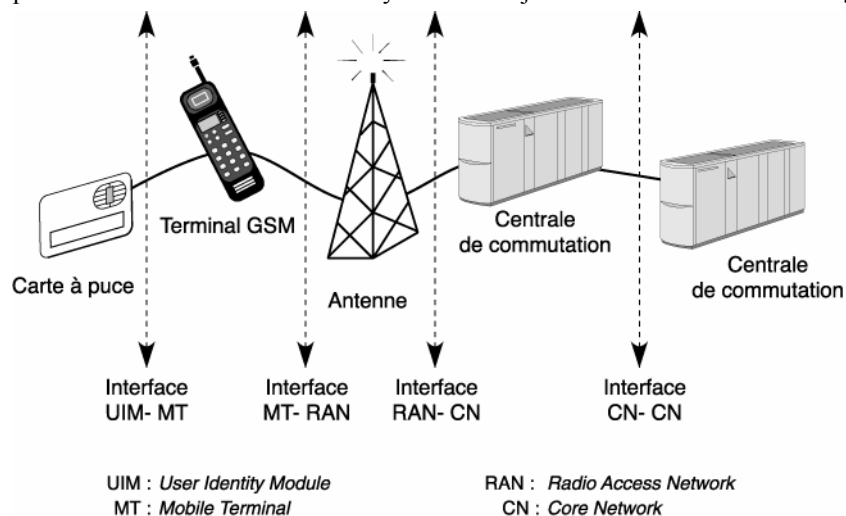


Figure 52. Les quatre interfaces définies dans l'architecture IMT-2000.

Ces interfaces sont les suivantes :

- UIM-MT (*User Identity Module-Mobile Terminal*), qui relie la carte déterminant l'identité de l'utilisateur au terminal mobile.
- MT-RAN (*Radio Access Network*), qui relie le terminal mobile et l'antenne, que l'on appelle aussi *interface air*.
- RAN-CN (*Core Network*, ou réseau central), qui relie l'antenne au réseau fixe du réseau de mobiles.
- CN-CN, qui interconnecte les commutateurs du réseau fixe.
- L'interface air constitue une partie importante de la normalisation d'un réseau de mobiles. Elle est chargée du partage des bandes de fréquences entre les utilisateurs, comme dans un réseau local. Si deux mobiles émettent au même moment sur une même fréquence, les signaux entrent en collision. Il faut donc une politique empêchant ces collisions de se produire. Nous sommes en présence de techniques semblables à celles de la couche MAC (*Medium Access Control*) des réseaux partagés. Les techniques de CSMA/CD ou de jeton ne s'adaptant guère à l'interface air, il a fallu développer d'autres solutions.

Les trois principales politiques de réservation utilisées dans le cadre des systèmes mobiles sont le FDMA (*Frequency Division Multiple Access*), ou accès multiple à répartition en fréquence, le TDMA (*Time Division Multiple Access*), ou accès multiple à répartition dans le temps, et le CDMA (*Code Division Multiple Access*), ou accès multiple à répartition en code.

Utilisé en premier, le FDMA a tendance à disparaître. Dans cette solution, la bande de fréquence f est découpée en n sous-bandes (voir figure 53) permettant à n mobiles distincts d'émettre en parallèle.

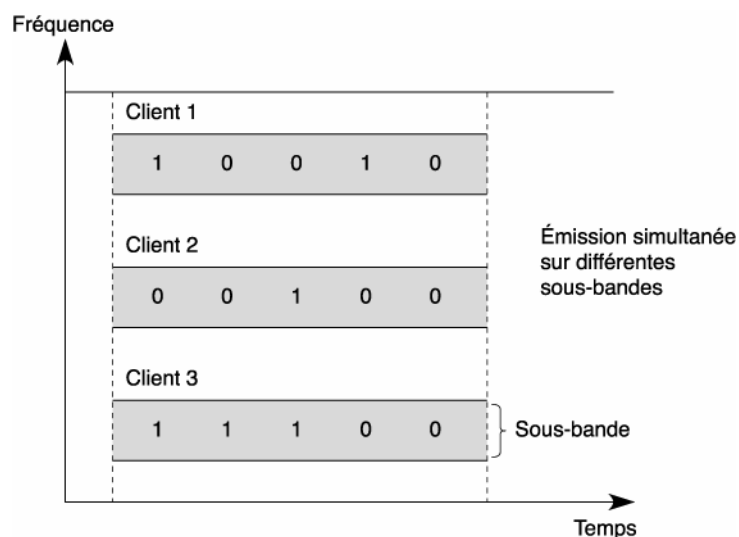


Figure 53. Le FDMA (génération du radiotéléphone analogique).

Chaque mobile comporte de ce fait un *modulateur*, un émetteur, n récepteurs et n démodulateurs. De plus, la station d'émission doit amplifier simultanément n porteuses. Il se crée donc nécessairement des produits d'*intermodulation*, dont l'intensité croît très rapidement en fonction de la puissance utile. Plus de la moitié de la capacité de transmission peut être ainsi perdue.

On évite les collisions en attribuant les fréquences entre les divers mobiles au fur et à mesure qu'ils se présentent dans la cellule. Les limites de cette technique sont évidentes : si une ou plusieurs liaisons sont inutilisées, suite à un manque de données à transmettre de la part de l'utilisateur, ce dernier conservant néanmoins la connexion, il y a perte sèche des bandes correspondantes.

Le TDMA propose une solution totalement différente, dans laquelle le temps est découpé en tranches. Ces tranches sont affectées successivement aux différents mobiles, comme illustré à la figure 54.

Tous les mobiles émettent avec la même fréquence sur l'ensemble de la bande passante, mais à tour de rôle. À l'opposé du fonctionnement en TDMA, chaque mobile est équipé d'un seul récepteur-démodulateur.

Un bloc transmis dans une tranche de temps comporte un en-tête, ou préambule, qui permet d'obtenir les différentes synchronisations nécessaires entre le mobile et la station de base. En particulier, il est nécessaire de synchroniser l'émission en début de tranche de façon qu'il n'y ait pas de chevauchement possible entre deux stations, l'une dépassant légèrement sa tranche ou l'autre commençant un peu tôt à transmettre.

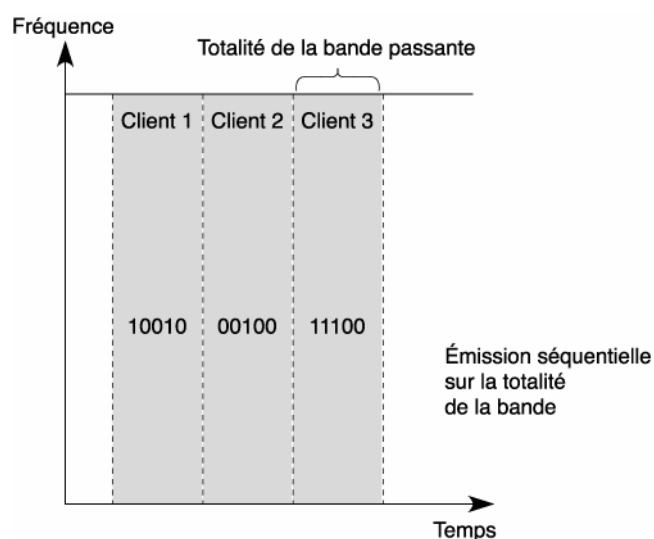


Figure 54. Le TDMA (génération du GSM).

Globalement, le rendement du TDMA est bien meilleur que celui du FDMA. C'est pourquoi la technique la plus employée pour accéder au *canal radio* a longtemps été le TDMA. Cette méthode soulève cependant des problèmes lors de la réutilisation des canaux radio dans des cellules avoisinantes pour cause d'interférences.

La méthode aujourd'hui retenue, aussi bien en Amérique du Nord qu'en Europe, est le CDMA (*Code Division Multiple Access*). Les mobiles d'une même cellule partagent un même canal radio en utilisant une forte partie de la largeur de bande passante attribuée à l'interface air. Le système affecte à chaque client un code unique, déterminant les fréquences et les puissances utilisées. Ce code est employé pour émettre le signal sur la largeur de bande B, à l'intérieur de la bande utile du signal R. Cette technique du CDMA, qui permet de réutiliser les mêmes fréquences dans les cellules adjacentes, est illustrée à la figure 55.

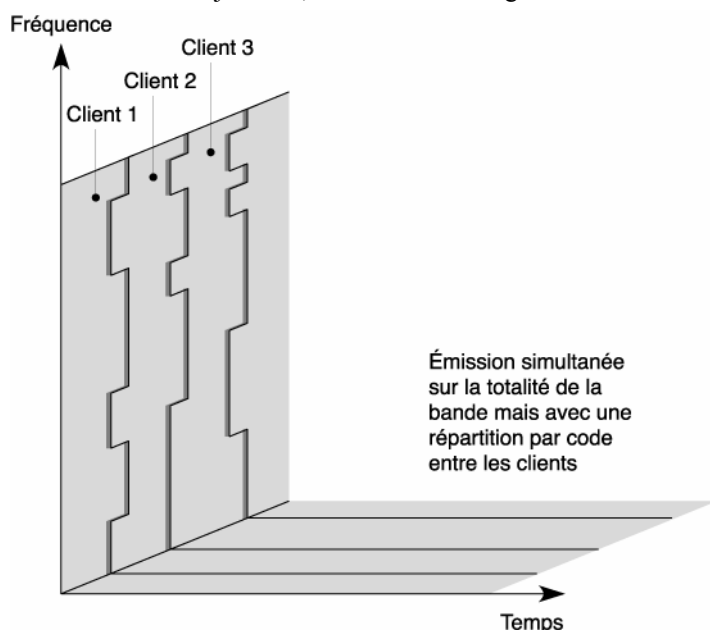


Figure 55. Le CDMA (génération future de l'UMTS).

La taille d'une cellule dépend fortement de la fréquence utilisée. Plus la fréquence est élevée, moins la portée est grande, et plus la fréquence est basse plus la portée est importante. Au départ, les fréquences utilisées allaient de 30 MHz à 300 MHz dans les bandes *UHF* (*Ultra High Frequency*) puis augmentaient, dans la gamme des *VHF* (*Very High Frequency*), de 300 MHz à 3 GHz. On utilise aujourd'hui des gammes de fréquences allant jusqu'à 20 GHz. Des fréquences encore plus élevées, pouvant atteindre 60 GHz, permettent l'utilisation de bande passante importante, les difficultés provenant de la grande directivité des ondes et d'un fort affaiblissement du signal dans les environnements pollués. La portée de ces ondes millimétriques est donc très particulière. En résumé, suivant la fréquence et la puissance utilisées, la taille des cellules varie énormément, allant de grandes cellules, que l'on appelle des cellules parapluie, à de toutes petites cellules, appelées microcellules, voire picocellules. Les différentes tailles de cellules sont illustrées à la figure 56.

L'année 1982 voit le démarrage de la normalisation d'un système de communication mobile dans la gamme des 890-915 MHz et 935-960 MHz, pour l'ensemble de l'Europe. Deux ans plus tard, les premiers grands choix sont faits, avec, en particulier, un système numérique. Le groupe d'étude GSM, (Groupe spécial mobile) finalise en 1987 une première version comportant la définition d'une interface radio et le traitement de la parole téléphonique.

Avec une autre version dans la gamme des 1 800 MHz (le DCS1800, ou *Digital Cellular System*), la norme GSM (*Global System for Mobile Communications*) est finalisée au début de l'année 1990. Cette norme est complète et comprend tous les éléments nécessaires à un système de communication

Le GSM et l'IS-95

1970 marque le début d'un travail entrepris pour établir une norme unique en matière de communication avec les mobiles. Dans le même temps, une bande de 25 MHz dans la bande des 900 MHz est libérée pour réaliser cette norme. En 1987, treize pays européens se mettent d'accord pour développer un réseau GSM (*Global System for Mobile communications*). En 1990, une adaptation dans la bande des 1 800 MHz est mise en place sous le nom de DCS1800 (*Digital Communication System 1 800 MHz*). À cette époque, l'*ETSI* finalise la normalisation du GSM900 et du DCS1800. De leur côté, les Américains reprennent une version du GSM dans la bande des 1 900 MHz, le DCS1900. Les principes généraux du GSM sont les mêmes pour ces trois adaptations.

Les opérateurs américains n'ont pas opté, dans un premier temps, pour la solution GSM européenne, leur propre développement ayant été fait en parallèle. De fait, plusieurs solutions sont disponibles, même si nous ne détaillons dans la suite que la norme IS-95. Aujourd'hui, la norme GSM est de plus en plus utilisée, mais sur une bande de

fréquence de 1 900 MHz, qui n'est pas utilisée en Europe, de telle sorte qu'il est nécessaire d'avoir un téléphone portable spécifique ou bien une option d'accès à la bande des 1 900 MHz sur son portable.

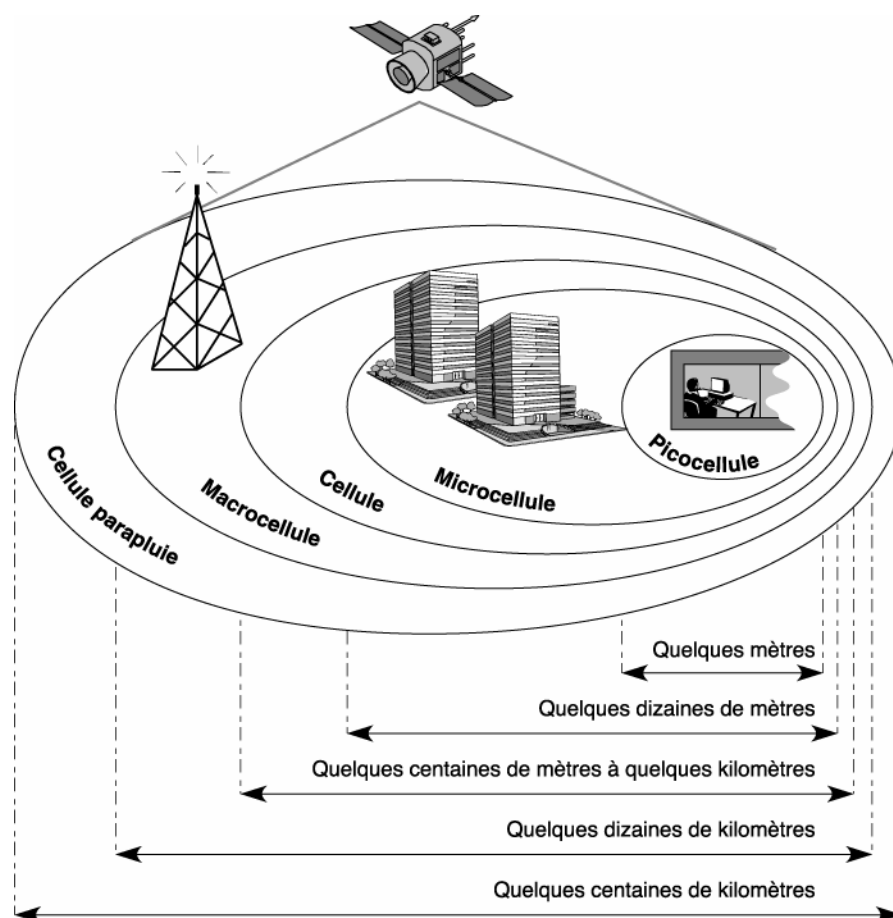


Figure 56. Les différentes tailles de cellules d'un réseau cellulaire.

Le GSM

Le GSM est un environnement complet, rassemblant l'interface radio mais aussi les interfaces entre le système radio et le système de commutation, d'un côté, et l'interface utilisateur, de l'autre. Les appels sont contrôlés par une norme déjà rencontrée dans le RNIS et le relais de trames.

En commençant par la partie la plus proche de l'utilisateur, la station mobile est constituée de deux éléments : le terminal portatif et la carte SIM. Cette carte à puce contient les caractéristiques de l'utilisateur ainsi que les éléments de son abonnement.

L'interface radio travaille dans les bandes 890-915 MHz dans le *sens montant* et 935-960 MHz dans le *sens descendant*. Une version GSM étendue a également été définie, le E-GSM, qui travaille dans les bandes 880-915 MHz dans le sens montant et 925-960 MHz dans le sens descendant. Le réseau DCS1800 utilise un sens montant entre 1 710 et 1 785 MHz et un sens descendant de 1 805 à 1 880 MHz. Enfin, le PCS1900 se place entre 1 850 et 1 910 MHz dans le sens montant et 1 930 à 1 990 MHz dans le sens descendant. Chaque porteuse radio exige 200 KHz, de telle sorte que 124 porteuses sont disponibles en GSM900, 174 en E-GSM, 374 en DCS1800 et 298 en DCS1900.

La solution préconisée dans le GSM est de ne pas multiplexer sur un seul canal tous les flots dont on a besoin pour réaliser une communication. Chaque flot possédant son propre canal, on peut dénombrer dix canaux travaillant en parallèle et ayant chacun leur raison d'être. Par exemple, le canal par lequel passent les données utilisateur est différent des canaux de signalisation, aux buts très divers, allant de la signalisation liée à l'appel à la signalisation correspondant à la gestion des fréquences. Les définitions de ces différents canaux sont regroupées dans l'aparté « Les canaux de l'interface radio GSM ».

Les canaux de l'interface radio GSM.

Le canal plein débit TCH/FS (*Traffic Channel*) offre un débit net de 13 Kbit/s pour la transmission de la parole ou des données. Ce canal peut être remplacé par :

- le canal demi-débit TCH/HS, à 5,6 Kbit/s ;
- le canal plein débit pour les données à 9,6 Kbit/s, pour la transmission de données au débit net de 12 Kbit/s ;
- le canal demi-débit pour les données à 4,8 Kbit/s, pour la transmission de données au débit net de 6 Kbit/s.

Le canal SDCCH (*Standalone Dedicated Control Channel*) offre un débit brut de 0,8 Kbit/s. Il sert à la signalisation (établissement d'appel, mise à jour de localisation, transfert de messages courts, services supplémentaires). Ce canal est associé à un canal de trafic utilisateur.

Le canal SACCH (*Standalone Access Control Channel*), d'un débit brut de 0,4 Kbit/s, est un canal de signalisation lent associé aux canaux de trafic. Son rôle est de transporter des messages de contrôle du handover.

Le canal FACCH (*Fast Access Control Channel*) est obtenu par un vol de trames (utilisation de tranches de temps vide d'un autre canal) sur le canal trafic d'un utilisateur dont il est chargé d'exécuter le handover. Il est donc associé à un canal de trafic. Il peut également servir pour des services supplémentaires, comme l'appel en instance (appel en attente pendant que vous êtes déjà en train de téléphoner).

Le canal CCCH (*Common Control Channel*) est un canal de contrôle commun aux canaux de trafic, où transitent les demandes d'établissement de communication et les contrôles de ressources.

Le canal BCCH (*Broadcast Control Channel*), d'un débit de 0,8 Kbit/s, gère le point à multipoint.

Le canal AGCH, canal d'allocation des accès, s'occupe de la signalisation des appels entrants.

Le canal RACH (*Random Access Channel*) s'occupe de la métasignalisation, correspondant à l'allocation d'un premier canal de signalisation.

Le canal FCCH (*Frequency Control Channel*) prend en charge les informations de correction de fréquence de la station mobile.

Le canal SCH (*Synchronous Channel*) est dédié aux informations de synchronisation des trames pour la station mobile et pour l'identification de la station de base.

Le protocole de niveau trame est chargé de la gestion de la transmission sur l'interface radio. Le protocole choisi provient du standard HDLC, auquel on a apporté quelques modifications de façon à s'adapter à l'interface air. Plus précisément, ce protocole est appelé LAP-D_m (*Link Access Protocol on the D_m channels*). Il transporte des trames avec une fenêtre de taille 1, la reprise éventuelle s'effectuant sur un temporisateur.

Le protocole de niveau paquet est lui-même divisé en trois sous-niveaux :

- La couche RR (*Radio Resource*), qui se préoccupe de l'acheminement de la supervision.
- La couche MM (*Mobility Management*), qui prend en charge la localisation continue des stations mobiles.
- La couche CM (*Connection Management*), qui gère les services supplémentaires, ainsi que le transport des messages courts SMS (*Short Message Service*) et le contrôle d'appel. Ce dernier contrôle reprend en grande partie la normalisation effectuée dans le cadre du réseau numérique à intégration de services (RNIS).

Le GSM définit les relations entre les différents équipements qui constituent le réseau de mobiles. Ces équipements sont les suivants :

- Le sous-système radio.
- Le sous-système réseau, avec ses bases de données pour la localisation des utilisateurs HLR et VLR, décrits à la section précédente.
- Les relations entre les couches de protocoles et les entités du réseau.
- Les interfaces entre sous-système radio et sous-système réseau.
- L'*itinérance* (*roaming*).

IS-95

L'IS-95 est la principale version américaine normalisée pour la seconde génération de réseaux de mobiles. L'interface air utilise la technologie CDMA (*Code Division Multiple Access*), équivalent américain du CDMA. La version IS-95A est celle qui a été déployée pour l'Amérique du Nord. La version 1999, IS-95B, qui augmente les débits numériques, est prise comme référence dans la présente section.

La technique CDMA, même si elle arrive à multiplexer de nombreux utilisateurs par l'utilisation de codes distincts, est assez difficile à maîtriser. En particulier, elle nécessite un contrôle permanent des puissances d'émission de façon à éviter toute ambiguïté entre plusieurs codes à la réception. Sans cela, entre un signal qui aurait perdu la moitié de sa puissance et un signal que l'on enverrait sur la même bande de fréquence avec une puissance divisée par deux, il n'y aurait plus de différence.

D'une façon assez différente de celle du GSM, les canaux de contrôle et les canaux utilisateur sont assez fortement multiplexés par le biais d'une méthode temporelle utilisant des tranches de temps de 20 ms. Il n'y a pas contradiction entre le multiplexage temporel et le *multiplexage en code*, les deux s'effectuant en parallèle.

Les canaux de contrôle de l'IS-95

Le canal de contrôle descendant regroupe le canal pilote, le canal de paging et le canal de synchronisation. Le canal réservé au trafic des utilisateurs est multiplexé avec les canaux de contrôle dans des trames de 20 ms. La trame est ensuite codée pour être transportée sur l'interface air. Le canal de synchronisation travaille à la vitesse de 1,2 Kbit/s. Chaque utilisateur possède un canal en CDMA et jusqu'à sept canaux de trafic supplémentaires. Deux taux de trafic ont été définis : l'ensemble 1, qui contient les débits de 9,6, 4,8, 2,4 et 1,2 Kbit/s, et l'ensemble 2, avec des débits de 14,4, 7,2, 3,6 et 1,8 Kbit/s. Les trames de 20 ms sont divisées en seize groupes de contrôle de puissance d'une durée de 1,25 ms.

La structure du canal montant est différente. Ce canal est en fait subdivisé en deux canaux, le canal de trafic et le canal de gestion de l'accès. Les trames font également 20 ms et prennent en charge l'ensemble du trafic.

L'IS-95 comporte trois mécanismes de contrôle de puissance : un contrôle en boucle ouverte et un contrôle en boucle fermée sur le lien montant et un contrôle en boucle plus lent sur le lien descendant.

L'IS-95 met en œuvre une technologie de codage de la parole à 8 Kbit/s et une autre à 13 Kbit/s. Ce dernier codage utilise le taux de 14,4 Kbit/s du canal de transmission. Le premier codage a recours à un codec EVRC (*Enhanced Variable Rate Codec*) à 8 Kbit/s, s'adaptant aux canaux à 1,2, 2,4, 4,8, et 9,6 Kbit/s pour des décompositions éventuelles du canal dans les débits de 1, 2, 4 et 8 Kbit/s.

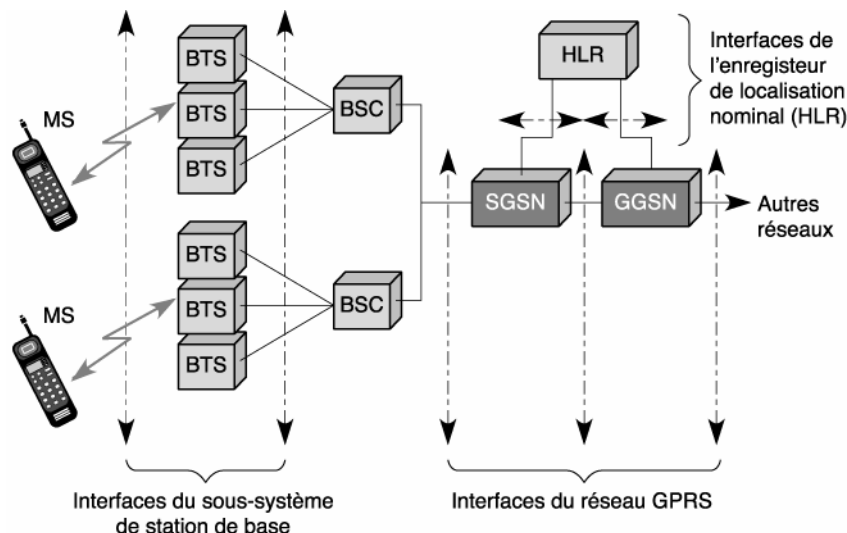
Le GPRS

L'une des activités majeures du développement de la phase 2+ du GSM concerne le GPRS (*General Packet Radio Service*), qui représente une nouvelle génération pour le standard GSM. Le GPRS prend en charge les applications multimédias dans le cadre de la mobilité. Il constitue également une transition vers la troisième génération des réseaux de mobiles par le passage d'un débit de 9,6 Kbit/s ou 14,4 Kbit/s à un débit beaucoup plus important, pouvant atteindre 170 Kbit/s.

Le GPRS peut être considéré comme un réseau de transfert de données avec un accès par interface air. Ce réseau utilise le protocole IP pour le formatage des données. Le transport des paquets IP s'effectue par des réseaux à commutation de trames, notamment le relais de trames.

L'architecture du GPRS est illustrée à la figure 57. Elle est composée des types de nœuds suivants :

- Les SGSN (*Serving GPRS Support Node*), qui sont des routeurs connectés à un ou plusieurs BSS.
- Les GGSN (*Gateway GPRS Support Node*), qui sont des routeurs vers les réseaux de données GPRS ou externes.



SGSN : *Serving GPRS Support Node*
GGSN : *Gateway GPRS Support Node*
HLR : *Home Location Register*

Figure 57. Architecture du GPRS.

Le réseau GPRS possède deux plans, le plan utilisateur et le plan de signalisation. Les couches de protocoles du plan utilisateur sont illustrées à la figure 58.

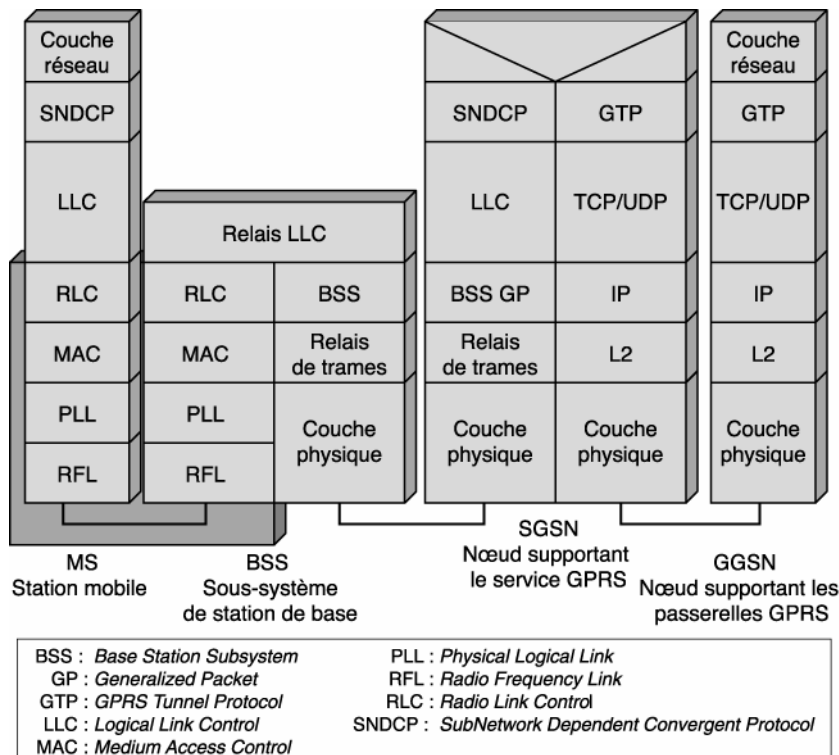


Figure 58. Couches de protocoles du plan utilisateur du réseau GPRS.

Par rapport au GSM, le GPRS requiert de nouveaux éléments pour créer un mode de transfert de paquets de bout en bout. De plus, le HLR est amélioré pour les clients qui demandent à transporter des données. Deux services sont ainsi permis :

- le point à point PTP (*Point-To-Point*) ;
- le point à multipoint PTM (*Point-To-Multipoint*).

Les transferts de paquets et le routage s'effectuent par les nœuds logiques SGSN (*Serving GPRS Support Node*). Ils utilisent également les passerelles GGSN avec les réseaux de transfert de paquets externes. Dans le réseau GPRS, les unités de données sont encapsulées par le SGSN de départ et décapsulées dans le SGSN d'arrivée. Entre les SGSN, c'est le protocole IP qui est utilisé. L'ensemble de ce processus est défini comme le « tunneling » du GPRS. Le GGSN maintient les informations de routage pour réaliser les tunnels et les maintenir. Ces informations sont stockées dans le HLR. Le protocole en charge de ce travail, le GTP (*GPRS Tunnel Protocol*) utilise les protocoles TCP et UDP pour effectuer le transport effectif. Entre le SGSN et les MS (*Mobile station*), le protocole SNDCP (*Subnetwork Dependent Convergence Protocol*) effectue le multiplexage de niveau paquet, ainsi que le chiffrement, la segmentation et la compression. Entre les MS et les BSS, le niveau trame est subdivisé en deux sous-couches, la couche LLC (*Logical Link Control*) et la couche RLC/MAC (*Radio Link Control/Medium Access Control*).

La couche LLC se sert du protocole LAP-D_m, déjà utilisé pour la signalisation dans l'environnement GSM. Le RLC/MAC s'occupe du transfert physique de l'information sur l'interface radio. En outre, ce protocole prend en charge les retransmissions éventuelles sur erreur par une technique BEC (*Backward Error Correction*), consistant en une retransmission sélective des blocs en erreur.

L'UMTS

L'UIT-T travaille sur une nouvelle génération de réseau de mobiles depuis 1985. D'abord connue sous le nom de FPLMTS (*Future Public Land Mobile Telephone System*) puis sous celui d'IMT-2000 (*International Mobile Telecommunications for the year 2000*), sa standardisation a commencé en 1990 en Europe, à l'ETSI. La version européenne s'appelle désormais UMTS (*Universal Mobile Telecommunications System*). Aux États-Unis, plusieurs propositions se sont fait jour, notamment une extension de l'IS-95 et le CDMA 2000. Les propriétés générales de cette génération, appelée 3G, sont les suivantes :

- couverture totale et mobilité complète jusqu'à 144 Kbit/s, voire 384 Kbit/s ;

- couverture plus limitée et mobilité jusqu'à 2 Mbit/s ;
- grande flexibilité pour introduire de nouveaux services.

La recommandation IMT-2000 n'a cependant pas fait l'unanimité, et de nombreux organismes de standardisation locaux ont préféré développer leur propre standard, tout en gardant une certaine compatibilité avec la norme de base. IMT-2000 est devenu le nom générique pour toute une famille de standards. Ce sont les Japonais qui ont démarré le processus en normalisant une version du CDMA (*Code Division Multiple Access*), suivis des Européens et des Américains. La version UMTS a repris en grande partie les spécifications japonaises sous le vocable W-CDMA (*Wideband CDMA*).

L'interface air normalisée par l'ETSI comprend deux types de bande :

- Les bandes de fréquences FDD (*Frequency Domain Duplex*), qui sont les bandes UMTS appariées, correspondant au passage des signaux dans les deux sens sur la même fréquence.
- Les bandes de fréquences TDD (*Time Domain Duplex*), qui sont les bandes UMTS non appariées, correspondant à l'utilisation d'une fréquence par sens de transmission.

Les bandes FDD utilisent le W-CDMA, tandis que, pour les bandes TDD, le choix s'est porté sur une méthode mixte, TD-CDMA (*Time Division CDMA*). Aux États-Unis, le choix s'est porté sur le standard IS-95-B, qui reprend une version du W-CDMA, appelée CDMA 2000. La même année 1998, les organismes de normalisation américains ont proposé une autre norme, UWC-136 (*Universal Wireless Communications*), fondée sur un multiplexage temporel TDMA.

Plusieurs ensembles de fonctionnalités ont été définis. Les fonctions retenues pour la première partie des années 2000 correspondent aux caractéristiques suivantes (*Capacity Set 1*, ou ensemble de capacité 1) :

- mobilité terminale ;
- mobilité personnelle ;
- environnement personnel virtuel ;
- fonctions permettant de recevoir et d'émettre des applications multimédias : 144 Kbit/s avec une mobilité forte et 2 Mbit/s avec une mobilité faible ;
- handovers devant pouvoir être effectués entre plusieurs membres de la famille IMT-2000 et avec les systèmes de deuxième génération ;
- réseau fixe devant suivre les technologies paquet et circuit ;
- interconnexion possible avec les réseaux RNIS, X.25 et Internet.

L'ensemble de capacité 2 (*Capacity Set 2*), mis en place dans les années 2002-2003, doit fournir les fonctionnalités suivantes :

- débits de 2 Mbit/s par utilisateur ;
- équipements de réseaux d'accès pouvant être mobiles, par exemple, une base postée dans un avion ;
- handovers entre tous les membres de la famille IMT-2000 ;
- handovers entre l'IMT-2000 et des systèmes non-IMT-2000.

De nombreuses investigations ont été menées depuis quelques années sur la technique CDMA (*Code Division Multiple Access*) dans le but de réaliser une interface air pour la troisième génération IMT-2000/UMTS/IS-95. La technique CDMA semble être la mieux placée pour les réseaux hertziens de troisième génération. Comme les débits devant transiter par ces interfaces atteignent 2 Mbit/s, il a fallu développer des versions spécifiques du CDMA pour le large-bande, le W-CDMA (*Wideband CDMA*) ou le CDMA 2000, comme expliqué précédemment. La figure 59 illustre le réseau UMTS.

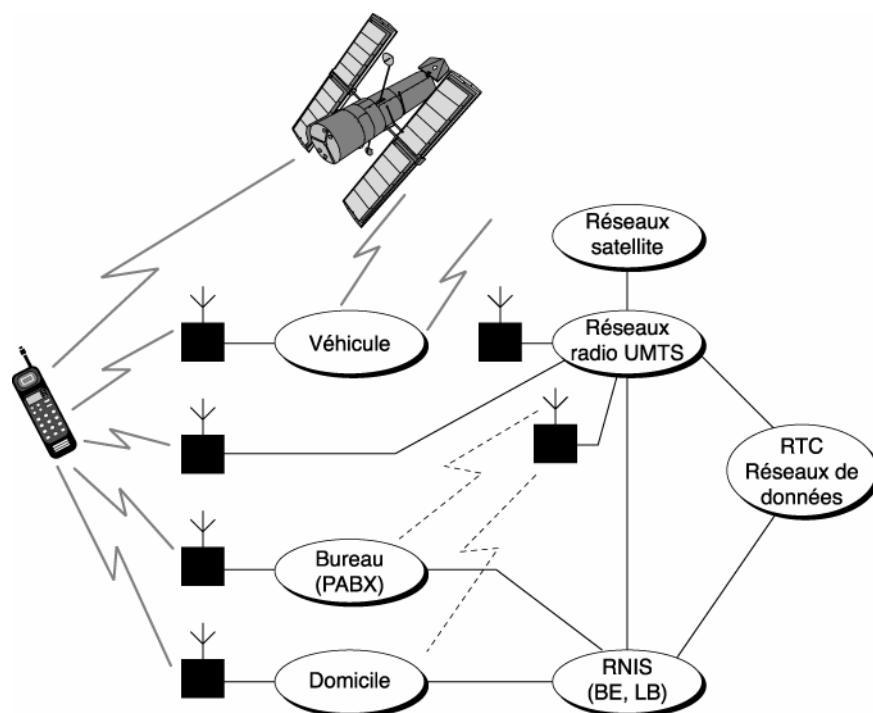


Figure 59. Le réseau UMTS global.

Les figures 60 et 61 donnent une idée du développement de l'architecture UMTS.

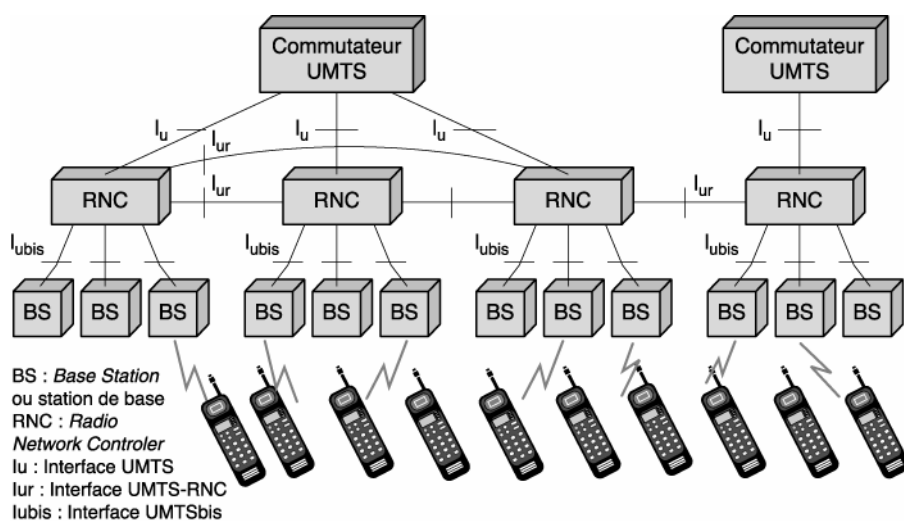


Figure 60. Équipements et interfaces de l'UMTS

Le radiologique (Software Radio)

Le radiologique (*Software Radio*) définit un émetteur-récepteur de fréquences radio, comme un téléphone portable ou un pager, reconfigurable ou personnalisable. Un radiologique idéal serait un système multifréquence et multimode capable de redéfinir dynamiquement par logiciel toutes les couches, y compris la couche physique. Le concept d'un tel système se fonde à la fois sur les standards radio que le terminal peut utiliser (GSM, W-CDMA, etc.) et sur l'architecture et la conception du produit final. L'architecture du produit comporte deux parties bien distinctes : la partie matérielle et la partie logicielle.

L'architecture matérielle utilisée dans un radiologique se fonde généralement sur un processeur de signaux numériques (DSP, pour *Digital Signal Processor*) totalement reprogrammable ou sur un circuit de type ASIC (*Application Specific Integrated Circuit*), où des fonctionnalités matérielles sont déjà prédéfinies mais où l'on peut tout de même reprogrammer au moins une de ces fonctions.

Pour la partie logicielle, qui inclut l'interprétation des signaux envoyés ainsi que les différents traitements tels que le contrôle des flux d'informations, il n'existe pas vraiment de spécification. On parle, entre autres choses, d'utiliser un moteur Java, qui permettrait de créer des interfaces utilisateur personnalisées à l'aide d'applets ou de scripts.

Pour arriver à réaliser un tel système, plusieurs problèmes doivent être surmontés, en particulier celui de l'énergie. Les processeurs de signaux numériques et les codecs sont très gourmands en énergie. De plus, la gestion logicielle de nombreux standards radio demande un traitement puissant et une capacité mémoire importante. Ces radiologiciels doivent gérer de nombreuses bandes de fréquences, tout en évitant les interférences. Il reste beaucoup à faire pour que le prix d'un portable radiologiciel devienne abordable pour le grand public.

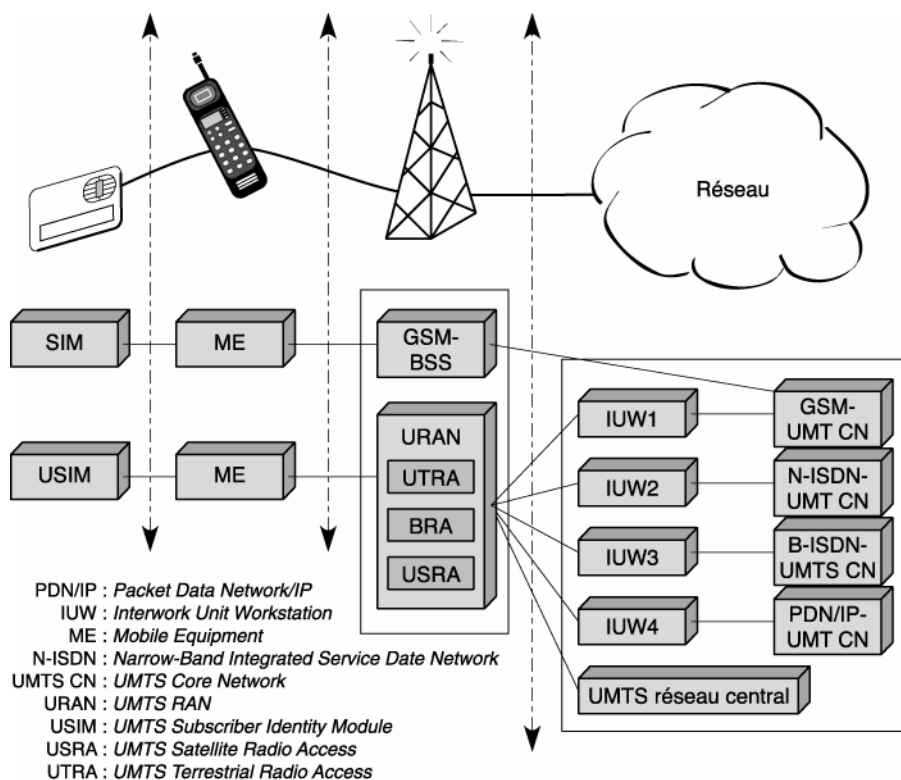


Figure 61. Architecture de l'UMTS et du GSM.

La mobilité locale

Nous avons principalement étudié les réseaux de mobiles au sens classique, dans lesquels les stations peuvent se déplacer sur le réseau d'un opérateur. Cette dernière section s'intéresse aux réseaux géographiquement limités, dans lesquels les terminaux se trouvent tous dans une même cellule. On peut appeler cela une mobilité restreinte, au contraire de la mobilité forte. Les exemples que nous examinons concernent les réseaux locaux sans fil, ainsi que les réseaux personnels et les réseaux *ad hoc*.

Les réseaux locaux sans fil

Les réseaux locaux sans fil sont en plein développement du fait de la flexibilité de leur interface, qui permet à un utilisateur de changer de place dans l'entreprise tout en restant connecté. Plusieurs produits sont actuellement commercialisés, mais ils sont souvent incompatibles entre eux en raison d'une normalisation relativement récente. Ces réseaux atteignent des débits de plusieurs mégabits par seconde, voire de plusieurs dizaines de mégabits par seconde.

Plusieurs solutions peuvent être envisagées : soit la communication hertzienne s'effectue sur l'ensemble du site, et tous les terminaux sont alors connectés directement entre eux ou par une seule borne, soit les communications s'effectuent à l'intérieur de microcellules, déterminées en général par les murs, et utilisent l'*infrarouge*. Les communications entre les équipements terminaux peuvent s'effectuer directement ou par le biais d'une borne intermédiaire. Quant aux communications entre bornes de concentration, elles peuvent s'effectuer de façon hertzienne ou par câbles.

La normalisation devrait avoir un fort impact sur les réseaux locaux sans fil. Aux États-Unis, c'est le groupe de travail IEEE 802.11 qui est en charge de cette normalisation, tandis que le groupe HiperLAN (*High Performance Radio LAN*) en est chargé sur le Vieux Continent.

Pour HiperLAN, les bandes de fréquences retenues se situent entre 5 150 et 5 300 MHz, auxquelles il faut ajouter une bande de 200 MHz dans les fréquences autour de 17 GHz. Les vitesses de transfert, de l'ordre de 19 à 25 Mbit/s, ne devraient pas atteindre les capacités des réseaux locaux filaires les plus rapides du marché, c'est-à-dire une centaine de mégabits par seconde, voire davantage. La distance entre les postes de travail les plus éloignés est de 5 km, mais une restriction des distances permet de garantir plus facilement la qualité du service demandé par l'utilisateur.

La communication peut se faire directement de station à station ou par l'intermédiaire d'un nœud central.

Du côté de l'IEEE 802.11, les fréquences choisies se situent dans la gamme des 2,4 GHz. Dans cette solution de réseau local par voie hertzienne, les communications peuvent se faire soit directement de station à station, mais sans qu'une station puisse relayer les paquets vers une autre station terminale, soit en passant par une borne de concentration. Les débits sont de 1 ou 2 Mbit/s, suivant la technique de codage utilisée.

La technique d'accès au support physique, le protocole MAC (*Medium Access Control*), est complexe mais unique et s'adapte à tous les supports physiques. De nombreuses options sont disponibles et rendent sa mise en œuvre assez complexe. La base provient de la technique CSMA/CD, mais comme la détection de collision n'est pas possible, on utilise un algorithme CSMA/CA (*Collision Avoidance*).

Pour éviter les collisions, chaque station possède un temporisateur avec une valeur spécifique. Lorsqu'une station écoute la porteuse et que le canal est vide, elle transmet. Le risque qu'une collision se produise est extrêmement faible, puisque la probabilité que deux stations démarrent leur émission dans une même microseconde est quasiment nulle. En revanche, lorsqu'une transmission a lieu et que d'autres stations se mettent à l'écoute et persistent à écouter, la collision devient inévitable. Pour empêcher la collision, il faut que les stations attendent, avant de transmettre, un temps permettant de séparer leurs instants d'émission respectifs. On ajoute pour cela un premier temporisateur très petit, qui permet au récepteur d'envoyer immédiatement un acquittement. Un deuxième temporisateur permet de donner une forte priorité à une application temps réel. Enfin, le temporisateur le plus long, dévolu aux paquets asynchrones, détermine l'instant d'émission pour les trames asynchrones.

Les réseaux sans fil

Les catégories de réseaux sans fil

Les réseaux sans fil sont en plein développement du fait de la flexibilité de leur interface, qui permet à un utilisateur de changer facilement de place dans son entreprise. Les communications entre équipements terminaux peuvent s'effectuer directement ou par le biais de stations de base. Les communications entre points d'accès s'effectuent de façon hertzienne ou par câble. Ces réseaux atteignent des débits de plusieurs mégabits par seconde, voire de plusieurs dizaines de mégabits par seconde.

Plusieurs gammes de produits sont actuellement commercialisées, mais la normalisation en cours devrait introduire de nouveaux environnements. Les groupes de travail qui se chargent de cette normalisation sont l'IEEE 802.15, pour les petits réseaux personnels d'une dizaine de mètres de portée, l'IEEE 802.11, pour les réseaux LAN (*Local Area Network*), ainsi que le l'IEEE 802.20, nouveau groupe de travail créé en 2003 pour le développement de réseaux un peu plus étendus.

Dans le groupe IEEE 802.15, trois sous-groupes normalisent des gammes de produits en parallèle :

- IEEE 802.15.1, le plus connu, en charge de la norme Bluetooth, aujourd'hui largement commercialisée.
- IEEE 802.15.3, en charge de la norme UWB (*Ultra-Wide Band*), qui met en œuvre une technologie très spéciale : l'émission à une puissance extrêmement faible, sous le bruit ambiant, mais sur pratiquement l'ensemble du spectre radio (entre 3,1 et 10,6 GHz). Les débits atteints sont de l'ordre du Gbit/s sur une distance de 10 mètres.
- IEEE 802.15.4, en charge de la norme ZigBee, qui a pour objectif de promouvoir une puce offrant un débit relativement faible mais à un coût très bas.

Du côté de la norme IEEE 802.11, dont les produits sont nommés Wi-Fi (*Wireless-Fidelity*), il existe aujourd'hui trois propositions, dont les débits sont situés entre 11 et 54 Mbit/s. Une quatrième proposition devrait bientôt proposer un débit de 320 Mbit/s.

La très grande majorité des produits sans fil utilise les fréquences de la bande 2,4-2,483 5 MHz et de la bande 5,15 à 5,3 MHz. Ces deux bandes de fréquences sont libres et peuvent être utilisées par tout le monde, à condition de respecter la réglementation en cours.

Les réseaux IEEE 802.11

La norme IEEE 802.11 a donné lieu à deux générations de réseaux sans fil, les réseaux Wi-Fi qui travaillent à la vitesse de 11 Mbit/s et ceux qui montent à 54 Mbit/s. Les premiers se fondent sur la norme IEEE 802.11b et les seconds sur les normes IEEE 802.11a et IEEE 802.11g. La troisième génération atteindra 320 Mbit/s avec la norme IEEE 802.11n.

Les fréquences du réseau Wi-Fi de base se situent dans la gamme des 2,4 GHz. Dans cette solution de réseau local par voie hertzienne, les communications peuvent se faire soit directement de station à station, mais sans qu'une station puisse relayer automatiquement les paquets vers une autre station terminale, à la différence des *réseaux ad-hoc*, soit en passant par un point d'accès, ou AP (*Access Point*).

Le point d'accès est partagé par tous les utilisateurs qui se situent dans la même cellule. On a donc un système partagé, dans lequel les utilisateurs entrent en compétition pour accéder au point d'accès. Pour sérialiser les accès, il faut définir une technique d'accès au support physique. Cette dernière est effectuée par le biais d'un protocole de niveau *MAC (Medium Access Control)* comparable à celui d'Ethernet. Ce protocole d'accès est le même pour tous les réseaux Wi-Fi.

De nombreuses options rendent toutefois sa mise en œuvre assez complexe. La différence entre le protocole hertzien et le protocole terrestre *CSMA/CD* d'Ethernet provient de la façon de gérer les collisions potentielles. Dans le second cas, l'émetteur continue à écouter le support physique et détecte si une collision se produit, ce qui est impossible dans une émission hertzienne, un émetteur ne pouvant à la fois émettre et écouter.

Le CSMA/CA (*Carrier Sense Multiple Access/Collision Avoidance*)

Dans le protocole terrestre CSMA/CD, on détecte les collisions en écoutant la porteuse, mais lorsque deux stations veulent émettre pendant qu'une troisième est en train de transmettre sa trame, cela mène automatiquement à une collision. Dans le cas hertzien, le protocole d'accès permet d'éviter la collision en obligeant les deux stations à attendre un temps différent avant d'avoir le droit de transmettre. Comme la différence entre les deux temps d'attente est supérieure au temps de propagation sur le support de transmission, la station qui a le temps d'attente le plus long trouve le support physique déjà occupé et évite ainsi la collision, d'où son suffixe *CA (Collision Avoidance)*.

Pour éviter les collisions, chaque station possède un temporisateur avec une valeur spécifique. Lorsqu'une station écoute la porteuse et que le canal est vide, elle transmet. Le risque qu'une collision se produise est extrêmement faible, puisque la probabilité que deux stations démarrent leur émission dans une même microseconde est quasiment nul. En revanche, lorsqu'une transmission a lieu et que deux stations ou plus se mettent à l'écoute et persistent à écouter, la collision devient inévitable. Pour empêcher la collision, il faut que les stations attendent avant de transmettre un temps suffisant pour permettre de séparer leurs instants d'émission respectifs. On ajoute également un petit temporisateur à la fin de la transmission afin d'empêcher les autres stations de transmettre et de permettre au récepteur d'envoyer immédiatement un acquittement.

L'architecture d'un réseau Wi-Fi est cellulaire. Un groupe de terminaux munis d'une carte d'interface réseau 802.11, s'associent pour établir des communications directes et forment un BSS (*Basic Set Service*).

Comme illustré à la figure 62, le standard 802.11 offre deux modes de fonctionnement, le mode infrastructure et le mode *ad hoc*. Le mode infrastructure est défini pour fournir aux différentes stations des services spécifiques sur une zone de couverture déterminée par la taille du réseau. Les réseaux d'infrastructure sont établis en utilisant des points d'accès, ou AP (*Access Point*), qui jouent le rôle de station de base pour une BSS.

Lorsque le réseau est composé de plusieurs BSS, chacun d'eux est relié à un système de distribution, ou DS (*Distribution System*), par l'intermédiaire de leur point d'accès (AP) respectif. Un système de distribution correspond en règle générale à un réseau Ethernet utilisant du câble métallique. Un groupe de BSS interconnectés par un système de distribution (DS) forment un ESS (*Extended Set Service*), qui n'est pas très différent d'un sous-système radio de réseau de mobiles.

Le système de distribution (DS) est responsable du transfert des paquets entre les différentes stations de base. Dans les spécifications du standard, le DS est implémenté de manière indépendante de la structure hertzienne et utilise un réseau Ethernet métallique. Il pourrait tout aussi bien utiliser des connexions hertziennes entre les points d'accès.

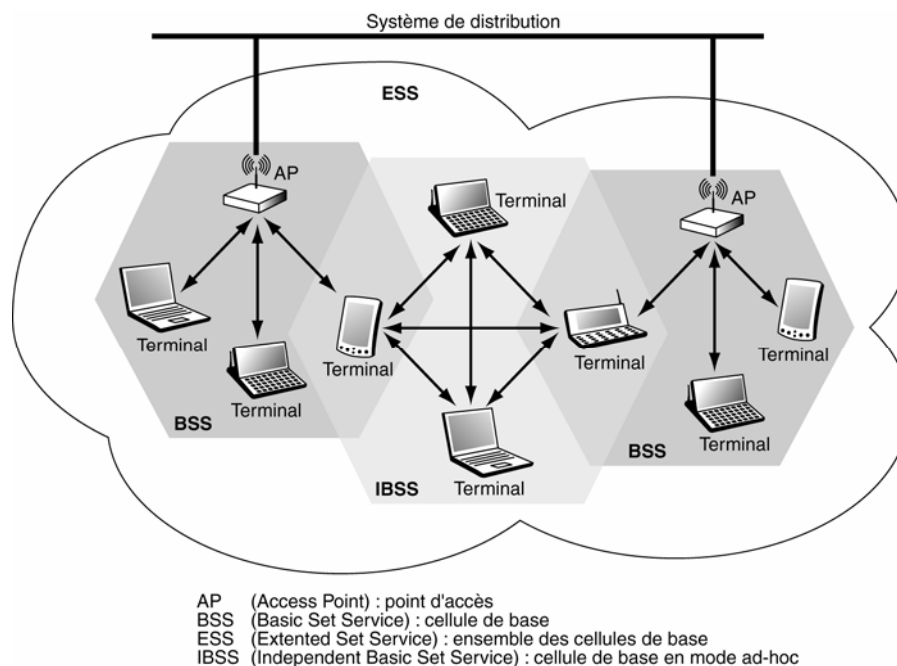


Figure 62. Architecture d'un réseau Wi-Fi.

Sur le système de distribution qui interconnecte les points d'accès auxquels sont connectées les stations mobiles, il est possible de placer une passerelle d'accès vers un réseau fixe, tel qu'Internet. Cette passerelle permet de connecter le réseau 802.11 à un autre réseau. Si ce réseau est de type IEEE 802.x, la passerelle incorpore des fonctions similaires à celles d'un pont.

Un réseau en mode *ad hoc* est un groupe de terminaux formant un IBSS (*Independent Basic Set Service*), dont le rôle consiste à permettre aux stations de communiquer sans l'aide d'une quelconque infrastructure, telle qu'un point d'accès ou une connexion au système de distribution. Chaque station peut établir une communication avec n'importe quelle autre station dans l'IBSS, sans être obligée de passer par un point d'accès. Comme il n'y a pas de point d'accès, les stations n'intègrent qu'un certain nombre de fonctionnalités, telles les trames utilisées pour la synchronisation.

Ce mode de fonctionnement se révèle très utile pour mettre en place facilement un réseau sans fil lorsqu'une infrastructure sans fil ou fixe fait défaut.

Les réseaux Wi-Fi

Pour qu'un signal soit reçu correctement, sa portée ne peut dépasser 50 m dans un environnement de bureau, 500 m sans obstacle et plusieurs kilomètres avec une antenne directive. En règle générale, les stations ont une portée maximale d'une vingtaine de mètres. Lorsqu'il y a traversée de murs porteurs, cette distance est plus faible.

La couche liaison de données

La couche liaison de données du protocole 802.11 est composée essentiellement de deux sous-couches, LLC (*Logical Link Control*) et MAC. La couche LLC utilise les mêmes propriétés que la couche LLC 802.2. Il est de ce fait possible de relier un WLAN à tout autre réseau local appartenant à un standard de l'IEEE. La couche MAC, quant à elle, est spécifique de l'IEEE 802.11.

Le rôle de la couche MAC 802.11 est assez similaire à celui de la couche MAC 802.3 du réseau Ethernet terrestre, puisque les terminaux écoutent la porteuse avant d'émettre. Si la porteuse est libre, le terminal émet, sinon il se met en attente. Cependant, la couche MAC 802.11 intègre un grand nombre de fonctionnalités que l'on ne trouve pas dans la version terrestre.

La méthode d'accès utilisée dans Wi-Fi est appelé DCF (*Distributed Coordination Function*). Elle est assez similaire à celle des réseaux traditionnels supportant le best-effort. Le DCF a été conçu pour prendre en charge le

transport de données asynchrones, transport dans lequel tous les utilisateurs qui veulent transmettre des données ont une chance égale d'accéder au support.

La sécurité

Dans les réseaux sans fil, le support est partagé. Tout ce qui est transmis et envoyé sur le support peut donc être intercepté. Pour permettre aux réseaux sans fil d'avoir un trafic aussi sécurisé que dans les réseaux fixes, le groupe de travail IEEE 802.11 a mis en place le protocole WEP (*Wired Equivalent Privacy*), dont les mécanismes s'appuient sur le chiffrement des données et l'authentification des stations. D'après le standard, le protocole WEP est défini de manière optionnelle, et les terminaux ainsi que les points d'accès ne sont pas obligés de l'implémenter.

Pour empêcher l'écoute clandestine sur le support, le standard fournit un algorithme de chiffrement des données. Chaque terminal possède une clé secrète partagée sur 40 ou 104 bits. Cette clé est concaténée avec un code de 24 bits, l'IV (*Initialisation Vector*), qui est réinitialisé à chaque transmission. La nouvelle clé de 64 ou 128 bits est placée dans un générateur de nombre aléatoire, appelé PRNG (RS4), venant de l'algorithme de chiffrement RSA (*Rivest Shamir Adelman*). Ce générateur détermine une séquence de clés pseudo-aléatoires, qui permet de chiffrer les données. Une fois chiffrée, la trame peut être envoyée avec son IV. Pour le déchiffrement, l'IV sert à retrouver la séquence de clés qui permet de déchiffrer les données.

Le chiffrement des données ne protège que les données de la trame MAC et non l'en-tête de la trame de la couche physique. Les autres stations ont donc toujours la possibilité d'écouter les trames qui ont été chiffrées.

Associés au WEP, deux systèmes d'authentification peuvent être utilisés :

- Open System Authentication ;
- Shared Key Authentication.

Le premier définit un système d'authentification par défaut. Il n'y a aucune authentification explicite, et un terminal peut s'associer avec n'importe quel point d'accès et écouter toutes les données qui transitent au sein du BSS. Le second fournit un meilleur système d'authentification puisqu'il utilise un mécanisme de clé secrète partagée.

Le mécanisme standard d'authentification de Wi-Fi

Ce mécanisme fonctionne en quatre étapes :

1. Une station voulant s'associer avec un point d'accès lui envoie une trame d'authentification.
2. Lorsque le point d'accès reçoit cette trame, il envoie à la station une trame contenant 128 bits d'un texte aléatoire généré par l'algorithme WEP.
3. Après avoir reçu la trame contenant le texte, la station la copie dans une trame d'authentification et la chiffre avec la clé secrète partagée avant d'envoyer le tout au point d'accès.
4. Le point d'accès déchiffre le texte chiffré à l'aide de la même clé secrète partagée et le compare à celui qui a été envoyé plus tôt. Si le texte est identique, le point d'accès lui confirme son authentification, sinon il envoie une trame d'authentification négative.

La figure 63 décrit le processus d'authentification d'une station, reprenant les quatre étapes que nous venons de détailler.

Pour restreindre encore plus la possibilité d'accéder à un point d'accès, ce dernier possède une liste d'adresses MAC, appelée ACL (*Access Control List*), qui ne permet de fournir l'accès qu'aux stations dont l'adresse MAC est spécifiée dans la liste.

Économie d'énergie

Les réseaux sans fil peuvent posséder des terminaux fixes ou mobiles. Le problème principal des terminaux mobiles concerne leur batterie, qui n'a généralement que peu d'autonomie. Pour augmenter le temps d'activité de ces terminaux mobiles, le standard prévoit un mode d'économie d'énergie.

Il existe deux modes de travail pour le terminal :

- Continuous Aware Mode ;
- Power Save Polling Mode.

Le premier correspond au fonctionnement par défaut : la station est tout le temps allumée et écoute constamment le support. Le second permet une économie d'énergie. Dans ce cas, le point d'accès tient à jour un enregistrement de toutes les stations qui sont en mode d'économie d'énergie et stocke les données qui leur sont adressées. Les stations qui sont en veille s'activent à des périodes de temps régulières pour recevoir une trame particulière, la trame TIM (*Traffic Information Map*), envoyée par le point d'accès.

Entre les trames TIM, les terminaux retournent en mode veille. Toutes les stations partagent le même intervalle de temps pour recevoir les trames TIM, de sorte à toutes s'activer au même moment pour les recevoir. Les trames TIM font savoir aux terminaux mobiles si elles ont ou non des données stockées dans le point d'accès. Lorsqu'un terminal s'active pour recevoir une trame TIM et s'aperçoit que le point d'accès contient des données qui lui sont destinées, il envoie au point d'accès une requête, appelée Polling Request Frame, pour mettre en place le transfert des données. Une fois le transfert terminé, il retourne en mode veille jusqu'à réception de la prochaine trame TIM.

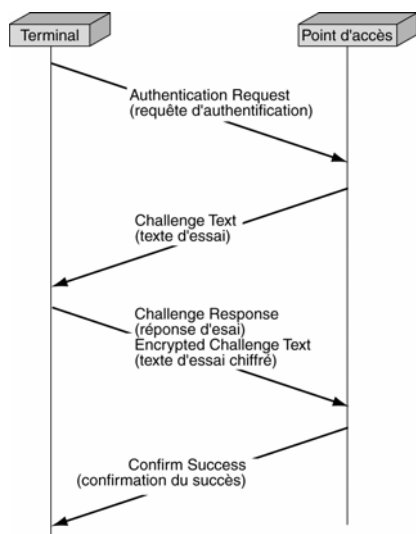


Figure 63. Le mécanisme d'authentification d'une station.

Pour des trafics de type broadcast ou multicast, le point d'accès envoie aux terminaux une trame DTIM (*Delivery Traffic Information Map*), qui réveille l'ensemble des points concernés.

IEEE 802.11b

Le réseau IEEE 802.11b provient de la normalisation effectuée sur la bande des 2,4 GHz. Cette norme a pour origine des études effectuées dans le cadre général du groupe IEEE 802.11.

En ce début des années 2000, la norme IEEE 802.11b s'est imposée comme standard, et plusieurs millions de cartes d'accès réseau Wi-Fi ont été vendues. Wi-Fi a d'abord été déployé dans les campus universitaires, les aéroports, les gares et les grandes administrations publiques et privées, avant de s'imposer dans les réseaux des entreprises pour permettre la connexion des PC portables et des équipements de type PDA.

Wi-Fi travaille avec des stations de base dont la vitesse de transmission est de 11 Mbit/s et la portée de quelques dizaines de mètres. Pour obtenir cette valeur maximale de la porteuse, il faut que le terminal soit assez près de la station de base, à moins d'une vingtaine de mètres. Il faut donc bien calculer, au moment de l'ingénierie du réseau, le positionnement des différents points d'accès. Si la station est trop loin, elle peut certes se connecter mais à une vitesse inférieure.

Aux États-Unis, treize fréquences sont disponibles sur la bande des 83,5 MHz. En Europe, lorsque la bande sera entièrement libérée, quatorze fréquences seront disponibles. Un point d'accès ne peut utiliser que trois fréquences au maximum, car l'émission demande une bande passante qui recouvre quatre fréquences.

Les fréquences peuvent être réutilisées régulièrement. De la sorte, dans une entreprise, le nombre de machines que l'on peut raccorder est très important et permet à chaque station terminale de se raccorder à haut débit à son serveur ou à un client distant.

IEEE 802.11a et g

Les produits Wi-Fi provenant des normes IEEE 802.11a et g utilisent la bande des 5 GHz. Cette norme a pour origine des études effectuées dans le cadre de la normalisation *HiperLAN* de l'ETSI au niveau européen en ce qui concerne la couche physique. La couche MAC de l'IEEE 802.11b est en revanche conservée.

Les produits Wi-Fi provenant de la norme 802.11a ne sont pas compatibles avec ceux de la norme Wi-Fi 802.11b, les fréquences utilisées étant totalement différentes. Les fréquences peuvent toutefois se superposer si l'équipement

qui souhaite accéder aux deux réseaux comporte deux cartes d'accès. En revanche, les produits Wi-Fi 802.11g travaillant dans la bande des 2,4 GHz sont compatibles et se dégradent en 802.11b si un point d'accès 802.11b peut être accroché. Il y a donc compatibilité avec la norme 802.11b.

La distance maximale entre la carte d'accès et la station de base peut dépasser les 100 m, mais la chute du débit de la communication est fortement liée à la distance. Pour le débit de 54 Mbit/s, la station mobile contenant la carte d'accès ne peut s'éloigner que de quelques mètres du point d'accès. Au-delà, le débit chute très vite pour être approximativement équivalent à celui qui serait obtenu avec la norme 802.11b à 100 m de distance.

En réalisant de petites cellules, de façon que les fréquences soient fortement réutilisables, et compte tenu du nombre important de fréquences disponibles en parallèle (jusqu'à 8), le réseau 802.11a permet à plusieurs dizaines de clients sur 100 m² de se partager plusieurs dizaines de mégabits par seconde. De ce fait, le réseau 802.11a est capable de prendre en charge des flux vidéo de bonne qualité.

La norme IEEE 802.11g a une tout autre ambition, puisqu'elle vise à remplacer la norme IEEE 802.11b sur la fréquence des 2,4 GHz, mais avec un débit supérieur à celui du 802.11b, atteignant théoriquement 54 Mbit/s mais pratiquement nettement moins, plutôt de l'ordre d'une vingtaine de mégabits par seconde.

Qualité de service et sécurité

La qualité de service est toujours un élément essentiel dans un réseau. Les réseaux 802.11 posent de nombreux problèmes pour obtenir de la qualité de service. Tout d'abord, le débit réel du réseau n'est pas stable et peut varier dans le temps. Ensuite, le réseau étant partagé, les ressources sont partagées entre tous les utilisateurs se trouvant dans la même cellule.

En ce qui concerne la première difficulté, les points d'accès Wi-Fi ont la particularité assez astucieuse de s'adapter à la vitesse des terminaux. Lorsqu'une station n'a plus la qualité suffisante pour émettre à 11 Mbit/s, elle dégrade sa vitesse à 5,5 puis 2 puis 1 Mbit/s. Cette dégradation provient soit d'un éloignement, soit d'interférences. Cette solution permet de conserver des cellules assez grandes, puisque le point d'accès s'adapte. L'inconvénient est bien sûr qu'il est impossible de prédire le débit d'un point d'accès. On voit bien que si une station travaille à 1 Mbit/s et une autre à 11 Mbit/s, le débit réel du point d'accès est plus proche de 1 Mbit/s que de 11 Mbit/s. De plus, comme l'accès est partagé, il faut diviser le débit disponible entre les différents utilisateurs.

Le groupe de travail IEEE 802.11 a défini deux normes, 802.11e et 802.11i, dans l'objectif d'améliorer les diverses normes 802.11 en introduisant de la qualité de service et des fonctionnalités de sécurité et d'authentification.

Ces ajouts ont pour fonction de faire transiter des applications possédant des contraintes temporelles, comme la parole téléphonique ou les applications multimédias. Pour cela, il a fallu définir des classes de service et permettre aux terminaux de choisir la bonne priorité en fonction de la nature de l'application transportée.

La gestion des priorités s'effectue au niveau du terminal par l'intermédiaire d'une technique d'accès au support physique modifiée par rapport à celle utilisée dans la norme de base IEEE 802.11. Les stations prioritaires ont des temporisateurs d'émission beaucoup plus courts que ceux des stations non prioritaires, ce qui leur permet de toujours prendre l'avantage lorsque deux stations de niveaux différents essayent d'accéder au support.

Le protocole IEEE 802.11i devrait apporter une sécurité bien meilleure que celle proposée par le WEP. Il devrait être mis en œuvre à partir du début de 2005. Il utilise un algorithme de chiffrement plus performant, avec l'adoption de l'AES (*Advanced Encryption Standard*). Déjà utilisé par la Défense américaine, ce standard a toutefois le défaut d'être incompatible avec la génération actuelle, de même qu'avec les extensions de sécurité en cours (*voir l'encadré sur la sécurité WPA*). Le protocole IEEE 802.11i devrait être mis en œuvre avec la génération IEEE 802.11n à un débit de 320 Mbit/s.

La sécurité WPA (*Wi-Fi Protected Access*)

Les mécanismes de sécurité ont fortement progressé depuis le début des années 2000. Le WPA a été introduit en 2003. Il propose deux processus de sécurité : un WEP amélioré, appelé TKIP (*Temporal Key Integrity Protocol*), et une procédure d'authentification des utilisateurs avec la norme IEEE 802.1x.

TKIP apporte une modification régulière des clés secrètes, de telle sorte que même un attaquant n'a pas le temps d'acquiescer un nombre suffisant de trames pour avoir un espoir de casser les clés secrètes. La norme IEEE 802.1x apporte une authentification, qui déborde du strict cadre de l'environnement Wi-Fi.