

Chapitre 3

Transmission des données couche physique

1. Rôle d'une interface réseau

Dans un premier temps, nous allons examiner les paramètres qui permettent de configurer les périphériques d'un PC et plus particulièrement une carte réseau.

1.1 Principes

L'interface réseau fait office d'intermédiaire entre l'ordinateur et le support de transmission. Au départ la carte réseau était un périphérique dédié (NIC - *Network Interface Card*), enfiché dans un connecteur d'extension (*slot*) de la carte mère ; depuis bien longtemps maintenant, elle est un simple composant soudé à la carte mère. Son rôle est de préparer les données à transmettre avant de les envoyer et d'interpréter celles reçues. Pour cela, elle contient un émetteur-récepteur.

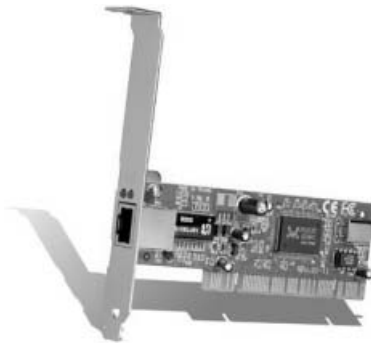


Carte réseau Ethernet intégrée sur une carte mère d'un PC

Le lien entre la carte et le système d'exploitation réseau est assuré par le pilote (*driver*) périphérique. Ce composant logiciel correspond à la couche Liaison de données du modèle OSI.

1.2 Préparation des données

La couche physique met en forme les données (bits) à transmettre sous forme de signaux. Les échanges entre l'ordinateur et la carte s'effectuent via le bus de la machine en parallèle. La carte réseau va donc sérialiser les informations avant de transmettre les signaux sur le support physique.



Ancienne carte réseau Ethernet

2. Options et paramètres de configuration

Tout point d'entrée/sortie sur un réseau doit être identifié afin que la trame soit reçue (acceptée) par le bon périphérique. Une carte réseau ou un port série doivent avoir un numéro qui doit permettre de les repérer au plus bas niveau (du modèle OSI).

2.1 Adresse physique

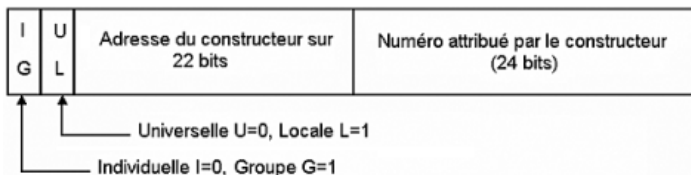
Sur un réseau local de type Ethernet (le plus courant, que nous aborderons plus tard), c'est une adresse physique sur six octets, qui permet d'identifier l'interface réseau. Les trois premiers octets de cette adresse sont attribués par l'IEEE pour identifier le constructeur du matériel (par exemple, HP : 64-4E-D7, Dell : 28-F1-0E, VMware : 00-50-56). Les trois octets restants sont laissés à la disposition du constructeur, qui doit faire en sorte de vendre des cartes, de telle manière qu'aucune n'ait la même adresse physique, sur le même réseau de niveau 2.

■ Remarque

Attention, un ordinateur peut disposer de plusieurs adresses physiques ou adresses MAC (Medium Access Control) : par exemple, un PC portable disposera d'une adresse pour sa carte Ethernet, d'une adresse pour sa carte Wi-Fi mais aussi d'une autre adresse lorsque celui-ci sera connecté à la station d'accueil (Dock USB-C faisant transiter également le réseau).

Un serveur (physique) disposera quant-à-lui généralement d'au moins trois adresses physiques : deux adresses pour les cartes réseau (pour faire un agrégat ou teaming), ainsi qu'une adresse pour la carte de gestion à distance.

Une adresse MAC (présente dans une trame réseau) va soit identifier une carte réseau unique ($I=0$), soit être associée à un ensemble de cartes ($G=1$). Cette adresse pourra être unique globalement ($U=0$) ou simplement unique sur un périmètre limité ($L=1$).



Remarque

Théoriquement, rien n'empêche le système d'exploitation réseau de travailler avec des adresses physiques différentes de celles du constructeur. Par exemple, sous Windows, en accédant aux propriétés de la carte réseau, il est possible d'imposer une nouvelle adresse physique différente de celle proposée par défaut. Il suffit alors de valider, et la nouvelle adresse MAC devient effective immédiatement !



Propriétés avancées d'une carte réseau sous Windows 11 (Oracle Virtual Box)

Remarque

Pour accéder directement aux propriétés réseau sous Windows vous pouvez exécuter **ncpa.cpl**.

Remarque

La commande **ipconfig /all** sous Windows ou **ip address** sous Linux permet d'obtenir l'adresse MAC.

Cette adresse est utilisée chaque fois qu'une station, ou plutôt sa carte réseau, a besoin d'émettre une trame vers une autre carte réseau. Il est néanmoins possible d'envoyer un paquet non pas à une, mais à plusieurs cartes en remplaçant l'adresse unique du destinataire par une adresse multiple (souvent une adresse de diffusion, soit **FFFFFFFFFFFF**, c'est-à-dire tous les bits des six octets mis à 1).

Ainsi, toute adresse référençant plusieurs hôtes aura son bit de poids fort (le plus à gauche) à '1' (ex. : **FFFFFF.FFFFFFFF**), à '0' dans le cas contraire (ex. : **00A024.B6132D**).

Par exemple, lorsqu'une carte réseau effectue une requête *Address Resolution Protocol* (ARP), elle envoie une diffusion sur son réseau de niveau 2, c'est-à-dire que le destinataire physique de la trame émise est « Tout le monde », **FF-FF-FF-FF-FF-FF**, comme ci-dessous :

```

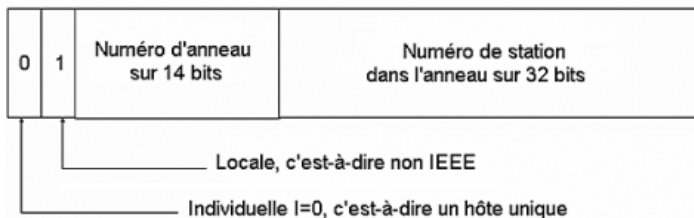
+Frame: Base frame properties
+ETHERNET: ETYPE = 0x0806 : Protocol = ARP: Address Resolution Protocol
+ETHERNET: Destination address : FFFFFFFFFFFFFF
+ETHERNET: Source address : 00A024B6132D
  ETHERNET: Frame Length : 60 (0x003C)
  ETHERNET: Ethernet Type : 0x0806 (ARP: Address Resolution Protocol)
  ETHERNET: Ethernet Data: Number of data bytes remaining = 46 (0x002E)
+ARP_RARP: ARP: Request, Target IP: 172.17.0.3
  ARP_RARP: Hardware Type = Ethernet (10Mb)
  ARP_RARP: Protocol Type = 2048 (0x800)
  ARP_RARP: Hardware Address Length = 6 (0x6)
  ARP_RARP: Protocol Address Length = 4 (0x4)
  ARP_RARP: Opcode = Request
  ARP_RARP: Sender's Hardware Address = 00A024B6132D
  ARP_RARP: Sender's Protocol Address = 172.17.0.92
  ARP_RARP: Target's Hardware Address = 000000000000
  ARP_RARP: Target's Protocol Address = 172.17.0.3
  ARP_RARP: Frame Padding

```

Identification d'une adresse de diffusion (niv. 2)

Une adresse attribuée par l'IEEE aura son deuxième bit de poids fort à '0', tandis qu'une valeur '1' précisera que l'adresse correspond à une adresse non normalisée.

Par exemple, en Token Ring (anciens réseaux locaux proposés par IBM et concurrents d'Ethernet), l'adresse d'un hôte était constituée comme suit :



Adressage physique Token Ring

■ Remarque

Historiquement, il était possible de créer des groupes en Token Ring (G=1).

■ Remarque

La liste exhaustive des préfixes d'adresses MAC attribués aux constructeurs (OUI - Organizationally Unique Identifiers) peut être consultée à partir de l'URL suivante : <http://standards-oui.ieee.org/oui.txt>

Par exemple, sur un PC portable Dell Latitude 5470, on obtient les informations ci-dessous :

- Carte Ethernet
 - Description : Intel(R) Ethernet Connection I219-LM
 - Adresse physique : 28-F1-0E-44-43-37
- Carte Wi-Fi
 - Description : Qualcomm QCA61x4A 802.11ac Wireless Adapter
 - Adresse physique : 94-53-30-19-15-0F

Sur le site de l'IEEE, on fait correspondre les préfixes des adresses MAC avec les informations dont on dispose :

28-F1-0E	(hex)	Dell Inc. One Dell Way Round Rock TX 78682 US
94-53-30	(hex)	Hon Hai Precision Ind. Co.,Ltd. Building D21,No.1, East Zone 1st Road Chongqing Chongqing 401332 CN

2.2 Interruption

Tout périphérique du PC est relié au microprocesseur par une ligne dédiée, ou ligne d'interruption (IRQ - *Interrupt ReQuest*). Lorsque le périphérique a besoin du microprocesseur pour travailler, il lui envoie un signal par cette ligne (tension électrique qui passe à l'état bas). Historiquement, les premiers PC comportaient 2 fois 8 lignes en cascade. Aujourd'hui, les systèmes d'exploitation intègrent 256 interruptions gérées de manière logicielle (Plug and Play). Certaines lignes sont attribuées par défaut et d'autres sont disponibles pour recevoir les périphériques supplémentaires. Le microprocesseur gère ces lignes par ordre de priorité : plus le numéro de l'interruption est faible, plus la priorité est élevée.

■ Remarque

Grâce à la technique du Plug and Play, qui permet la détection de la carte et l'affectation automatique de ses paramètres, il n'est plus vraiment utile aujourd'hui de connaître ces informations. La période où les IRQ pouvaient générer des conflits entre périphériques est désormais révolue.

2.3 Adresse d'entrée/sortie

Un périphérique interrompt le microprocesseur chaque fois que des informations ont besoin d'être échangées. Ces informations sont reçues ou envoyées par une porte d'entrée/sortie localisée à une adresse particulière : l'adresse d'entrée/sortie. Cette adresse pointe sur une plage d'au plus 32 octets, qui va permettre de stocker des données, mais aussi des informations indiquant ce qu'il faut faire de ces données.

2.4 Adresse de mémoire de base

Il s'agit d'une adresse de mémoire volatile dont le rôle est de faire un tampon (buffer), lors de la réception ou l'émission de trame sur le réseau.

Cette adresse doit être un multiple de 16, elle est donc souvent écrite en hexadécimal sans le '0' final qui est sous-entendu.

2.5 Canal DMA (Direct Memory Access)

Dans la plupart des cas, les périphériques dépendent du microprocesseur pour transférer des informations de leur tampon vers la mémoire vive ou en sens inverse. Ainsi, il existe des périphériques qui disposent d'un canal particulier pour pouvoir échanger directement des informations avec la mémoire vive du PC, sans avoir recours au microprocesseur (dans un deuxième temps).

Certains périphériques, notamment des cartes réseau disposent d'un canal DMA, de 1 à 7.

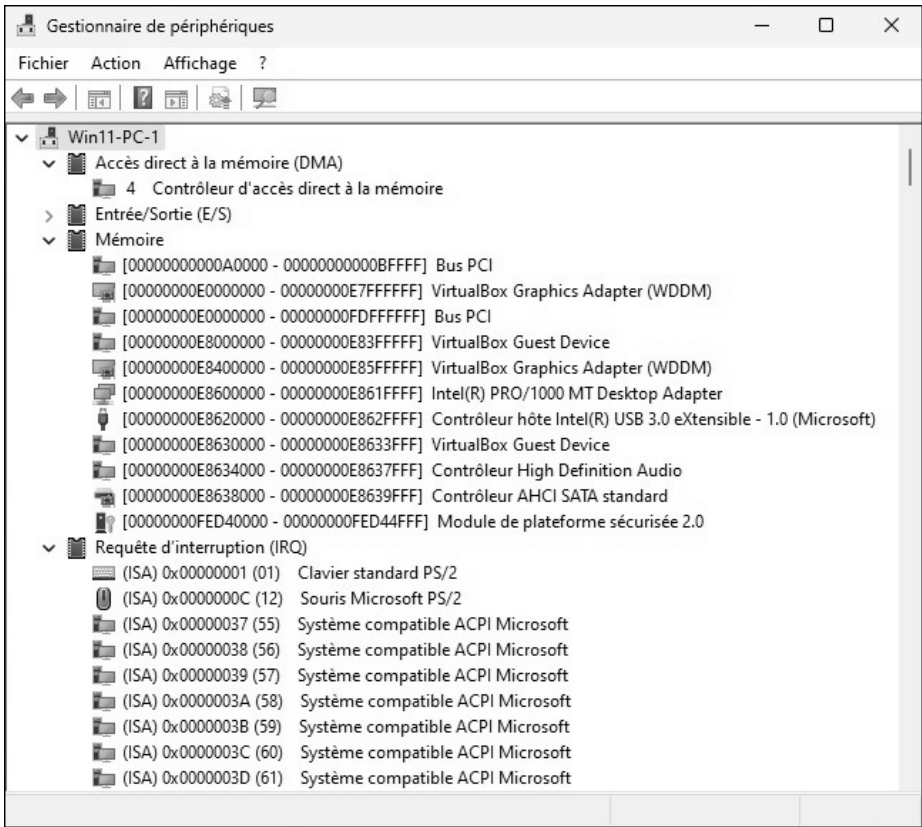
2.6 Bus

Toutes les données échangées entre les périphériques et l'ordinateur passent par des bus de données. Pendant longtemps, cet échange était surtout effectué à travers des voies parallèles, et la vitesse de transmission dépendait beaucoup de sa largeur, par exemple 16, 32 ou 64 bits. Les nouvelles technologies de bus privilégient des solutions de transferts en série, dans lesquels les bits sont envoyés les uns après les autres.

Avec l'évolution des moyens, les débits sont désormais bien supérieurs et les connecteurs sont plus petits.

Les bus historiques, *Industry Standard Architecture* (ISA), *Extended Industry Standard Architecture* (EISA) et *Micro Channel Architecture* (MCA) ont laissé leur place à d'autres, plus modernes.

Le **Gestionnaire de périphériques** sous Windows permet d'obtenir les informations précises sur l'allocation des ressources matérielles (menu **Affichage**, **Ressources par Type**) :



Affichage des adresses mémoire

Chapitre 3

Évolution technique en 5G

1. Cas d'usage

La cinquième génération des réseaux mobiles a apporté des évolutions techniques importantes par rapport aux standards précédents. Parmi ces évolutions figure le très haut débit, beaucoup plus important que la 4G, la réduction de la latence jusqu'à 1 milliseconde, la fiabilité et la grande capacité du réseau en nombre d'objets connectés.

Ces évolutions techniques reposent sur des innovations technologiques en rapport avec l'architecture du réseau, qui font appel au cloud, à la virtualisation et au Network Slicing et qui font évoluer la partie radio pour plus d'efficacité tout en réduisant la consommation d'énergie et le coût d'exploitation. Les plus importantes de ces innovations font l'objet de ce chapitre.

Les utilisations de la 5G sont réparties dans trois volets majeurs :

- eMBB (*Enhanced Mobile Broadband*) : le haut débit mobile amélioré ;
- CC&URLLC (*Critical Communications and Ultra Reliable and Low Latency Communications*) : les communications critiques, à ultra haute fiabilité et faible latence ;
- mMTC (*massive Internet of Things*) : l'Internet des objets massif.

1.1 eMBB

Il s'agit du service classique de connexion à Internet sur des smartphones ou tablettes. L'un des objectifs de la 5G est d'atteindre des débits de transferts de données cent fois plus importants que ceux réalisés en 4G. Cela rend plus fluides la navigation web, le streaming vidéo, les appels vidéo et la réalité virtuelle. En plus, cela ouvre la possibilité pour d'autres cas d'usage sous le volet eMBB. Par exemple, en aviation, eMBB peut fournir à un avion en plein vol un débit de 1,2 Gb/s.

L'eMBB repose sur les innovations suivantes :

- l'usage d'un spectre de fréquence plus large avec des fréquences millimétriques (*mm Wave* ou *mmW*) qui permettent d'atteindre les 20 Gb/s théoriques avec, bien entendu, une portée beaucoup plus limitée ;
- l'évolution des antennes utilisées avec l'usage du mMIMO (*massive Multi Input Multi Output*) et du *beamforming*. En 5G, il est possible de cibler une direction particulière, un client particulier. C'est le *beamforming* ou formation de faisceaux. Tandis qu'en 4G, le signal est diffusé par l'antenne dans toutes les directions de son champ de propagation.

1.2 CC&URLLC

Une communication critique est une communication dont l'échec n'est pas envisageable car il peut y avoir des conséquences très graves, dont la perte de vies. C'est pourquoi ces communications doivent avoir une fiabilité et une sécurité maximales, ainsi qu'une latence très faible et le plus proche possible de 0 ms. Ainsi, elles sont dites communications à ultra haute fiabilité et faible latence URLLC.

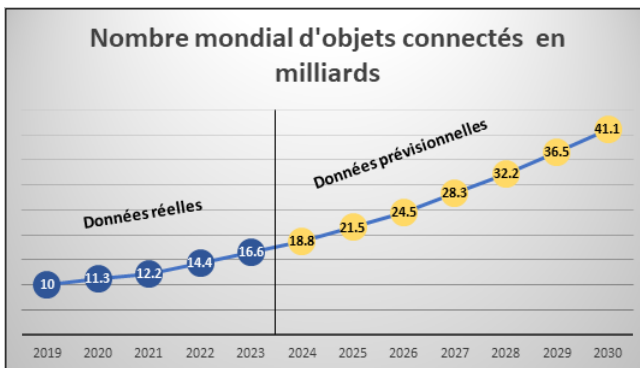
Les domaines où l'on trouve des communications critiques sont les services d'urgence (protection civile, ambulanciers, police), la santé (télémédecine et surveillance à distance), le transport (véhicules autonomes et systèmes de gestion de trafic), l'infrastructure critique (réseaux électriques et systèmes dont l'interruption même brève n'est pas tolérée) et l'industrie intelligente (automatisation et cycles de production dont l'arrêt peut avoir des conséquences graves pour la sécurité des employés et pour le résultat de l'entreprise).

Les communications critiques reposent sur les points suivants :

- Edge Computing : le traitement de l'information à la périphérie réduit considérablement la latence. Il permet aussi d'améliorer la sécurité des données puisque celles-ci restent sur place et le risque d'exposition aux attaques est minimisé.
- mMIMO : la multiplication du nombre de connexions grâce au mMIMO apporte plus de fiabilité et améliore la qualité du signal.
- Le Network Slicing : la création de réseaux spécifiques permet d'optimiser les ressources et de séparer les flux, ce qui assure la fiabilité et la sécurité.
- Les mécanismes de redondance intégrés permettent d'assurer la transmission des données.
- L'optimisation des protocoles de transport et de signalisation permet davantage de fiabilité et une réduction maximale de la latence.
- Le renforcement des protocoles de sécurité.

1.3 mIoT

Ci-dessous l'évolution du nombre mondial en milliards d'objets connectés, selon l'entreprise allemande IoT Analytics. Les données de 2024 sont prévisionnelles puisque le chiffre 18,8 milliards est l'estimation du nombre d'appareils IoT connectés dans le monde en fin d'année 2024.



Évolution du nombre mondial d'objets connectés de 2019 à 2030

Entreprise fondée et implantée en Allemagne, IoT Analytics est l'un des principaux fournisseurs mondiaux de la veille stratégique et des informations du marché de l'Industrie 4.0 (IoT, intelligence artificielle, cloud, etc.). Leurs principaux axes de travail dans l'ensemble de la pile technologique comprennent les applications IoT, les plateformes et logiciels IoT, la connectivité et le matériel IoT et l'IoT industriel. Des entreprises comme Qualcomm, Huawei, Oracle et Sony leur font confiance pour leurs informations sur le marché. Ainsi, les données de IoT Analytics sont considérées comme source fiable pour les statistiques relatives à l'IoT.

L'annexe 1 détaille en anglais et en français la méthodologie d'IoT Analytics dans le comptage du nombre d'objets connectés.

Le mIoT est aussi connu par mMTC (*massive Machine Type Communications*). Dans des lieux comme les usines, les gares, les stades, les moyens de transport ou les centres de congrès, le mIoT doit supporter un grand nombre de machines, pour des usages professionnels ou personnels, d'une manière transparente, en générant une quantité immense de données et avec des débits de transfert très élevés.

Parmi les cas d'usage du mIoT :

- le développement des villes intelligentes : gestion intelligente du trafic, gestion des déchets et la surveillance environnementale ;
- l'agriculture de précision : surveillance des sols, évaluation de la santé des cultures et contrôle de l'irrigation ;
- dans le domaine de la santé : appareils de santé portables, surveillance à distance des patients et suivi des actifs hospitaliers ;
- l'automatisation industrielle : maintenance prédictive dans les usines, suivi des actifs et optimisation des processus.

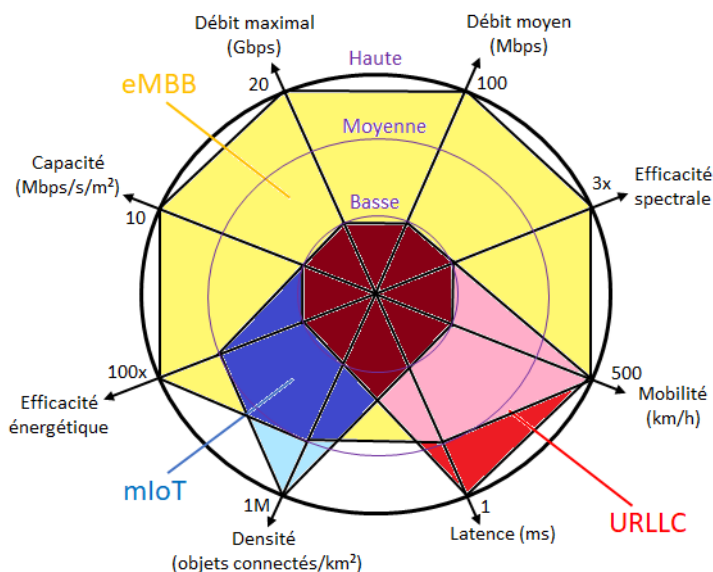
Ci-dessous les défis techniques du mIoT :

- L'évolutivité : pour la conception de réseaux capables d'accueillir plusieurs dizaines de milliers d'appareils connectés au km² (source : le site d'Orange), il faut relever un défi de taille. Il faut reposer sur des protocoles évolutifs, des schémas d'adressage et des mécanismes de routage avec une efficacité optimale.
- L'efficacité énergétique : les protocoles de communication doivent prendre en compte l'usage de batteries dans l'alimentation des objets connectés. Celles-ci doivent fonctionner le plus longtemps possible et ne doivent pas être remplacées fréquemment. Ainsi, les protocoles doivent être économes en énergie. En outre, des conceptions matérielles à faible consommation sont exigées.
- La connectivité : une connectivité fiable pour tous les appareils doit être garantie, quelles que soient les conditions. Des protocoles de communication robustes sont essentiels.
- La gestion des données : le mIoT génère des quantités énormes de données. Cela implique des techniques efficaces d'agrégation, de compression et de stockage des données.
- La sécurité et la confidentialité : la sécurité est peut-être le défi le plus difficile à relever de nos jours. Plus il y a de données, plus les réseaux sont vulnérables. Les métiers de la cybersécurité seront de plus en plus créés avec l'augmentation des volumes de données tous les jours. Sécuriser les réseaux mIoT contre les cyberattaques et garantir la confidentialité des données est plus qu'essentiel.
- La qualité de service (QoS) : c'est très complexe d'équilibrer les exigences de qualité de service de diverses applications (par exemple, surveillance en temps réel, rapports périodiques) au sein d'un réseau partagé. Cependant, c'est indispensable.

Le mIoT reposera sur les techniques qui assurent une connectivité à très grande fiabilité, une très grande capacité du réseau, une efficacité optimale de la consommation d'énergie et de l'exploitation des ressources et une sécurité de pointe. Citons le mMIMO, l'usage des basses fréquences, le beamforming, le Network Slicing, l'Edge Computing et la virtualisation.

1.4 Importance des fonctionnalités clés de la 5G

L'évolution technologique des réseaux 5G doit prendre en charge huit fonctionnalités clés synthétisées par l'illustration ci-dessous. L'importance de ces *features* varie selon le cas d'usage. Par exemple, en eMBB ou en mIoT, la latence peut être supérieure à 1 ms, voire atteindre des valeurs observées en 4G, ce qui n'est pas toléré en URLLC. Un débit moyen de 100 Mb/s est exigé en eMBB, mais en mIoT et en URLLC le très haut débit n'est pas primordial.



Importance des fonctionnalités clés de la 5G

Afin de prendre en charge les huit fonctionnalités et les trois cas d'usage et étant donné la croissance accélérée des réseaux 5G, toutes les bandes de fréquences doivent être exploitées. C'est l'objet de la prochaine section.

2. Spectre de fréquences

Les porteuses des télécoms vont de 400 MHz à 100 GHz, mais les bandes sous licence vont de 600 MHz à 39 GHz. Elles sont réparties en deux catégories :

FR1	450 MHz - 6 GHz
FR2	24.25 GHz – 52.6 GHz

FR2 correspond aux hautes fréquences, également appelées fréquences mmW. FR1 est divisée en fréquences basses (en dessous de 1 GHz) et médianes (1 GHz – 6 GHz).

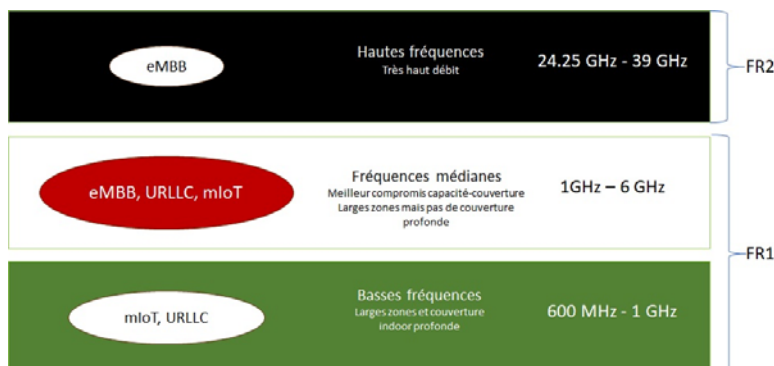
Ainsi, FR1 et FR2 donnent trois plages de fréquences à utiliser en 5G :

- **Basses fréquences** : généralement utilisée pour un déploiement rural, cette plage de basses fréquences, avec ses meilleures caractéristiques de propagation, est destinée à couvrir de grandes zones avec une bande passante maximale de 100 MHz pour une porteuse. Une couverture *indoor* profonde est garantie avec les basses fréquences.
- **Fréquences médianes** : destinée au déploiement de la 5G dans un contexte urbain ou suburbain, cette plage intermédiaire est le meilleur compromis entre la couverture et la capacité. Ici aussi, la bande passante maximale est de 100 MHz.
- **Hautes fréquences (FR2)** : la bande passante est plus élevée vers l'utilisateur (bande passante maximale de 400 MHz), mais la propagation est plus faible avec les hautes fréquences. Cette gamme est destinée aux environnements urbains denses (couverture de type « hotspot »).

L'illustration suivante synthétise le mapping des cas d'usage majeurs de la 5G sur ces trois plages de fréquences.

70 Les réseaux 5G

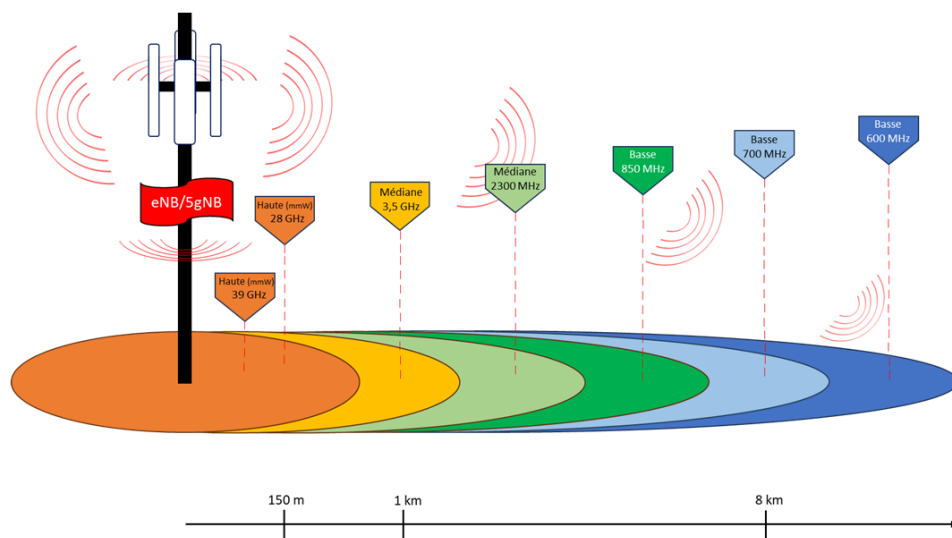
Enjeux et usages d'un nouveau monde numérique



Mapping cas d'usage 5G vs spectre de fréquences

Chez Orange en France, c'est la bande de fréquences 3,5 GHz (3,4 - 3,8 GHz) qui sera utilisée en priorité pour le réseau mobile 5G. Cette bande médiane ou bande C est considérée comme la bande cœur de la 5G.

La figure ci-dessous schématise la couverture radio en fonction de la fréquence. Il s'agit d'un site 4G/5G mais les fréquences millimétriques sont uniquement utilisées en 5G.



Portée des fréquences utilisables en 5G

En se référant aux données de l’UIT, la portée du signal ne dépend pas que de la bande de fréquences utilisée ; elle est aussi impactée par l’environnement.

Le tableau suivant synthétise la portée en zone urbaine et rurale de trois bandes de fréquences : basse (700 MHz), médiane (3.5 GHz) et haute fréquence (26 GHz) :

Catégorie	Bande	Portée en zone urbaine	Portée en zone rurale
FR 1	700 Mhz	2 km	8 km
	3.5 GHz	400 m	1.2 km
FR 2	26 GHz	150 m	-

3. Évolution des antennes

Le déploiement de la 5G avec un usage des fréquences traditionnelles, déjà utilisées en 2G, 3G et 4G, ne nécessite pas forcément l’ajout d’antennes 5G. Il est possible de se limiter à une intégration de la 5G par de simples mises à jour sur les sites existants.

Cependant, pour la fréquence 3.5 GHz par exemple, ou pour les fréquences mmW, il est nécessaire de réaliser des opérations sur l’infrastructure du réseau. Au moins l’ajout d’une antenne 5G est requis pour un déploiement qui, dans un premier temps, se fait par extension de l’architecture 4G (5G NSA ou *Non-Standalone 5G*) sans l’ajout de nouveaux sites 5G. Selon la demande des opérateurs, la deuxième phase de déploiement consiste à ajouter de nouveaux sites 5G SA (*Standalone*).

Chapitre 3

Le Multi-Link Operation

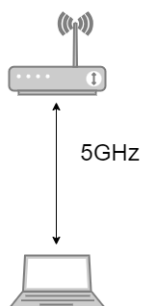
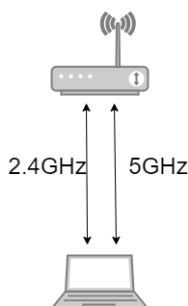
1. Présentation du MLO

Dans un monde où les besoins en connectivité rapide et fiable ne cessent de croître, les réseaux Wi-Fi sont soumis à des exigences sans précédent en termes de débit, de latence et de robustesse. La norme Wi-Fi 7 répond à ces nouveaux défis en introduisant une avancée majeure : le *Multi-Link Operation* (MLO). Cette technologie marque une rupture avec les générations précédentes, en permettant aux appareils d'exploiter plusieurs bandes de fréquences simultanément pour maximiser les performances.

Pour comprendre l'impact du MLO, imaginons les bandes de fréquences comme des modes de transport :

- La bande 2,4 GHz représente une route : elle est largement accessible, capable de transporter des données sur de longues distances, mais est fortement utilisée et peut être sujette à la congestion.
- La bande 5 GHz est un réseau ferroviaire : plus rapide, moins encombré, mais dont l'accès peut être limité par la densité des appareils et des obstacles.
- La bande 6 GHz est l'aviation : rapide et capable de transmettre de grands volumes de données sur des canaux larges, mais souvent réservée aux dispositifs les plus modernes.

Dans les générations précédentes de Wi-Fi, un appareil devait choisir une seule de ces « voies » à la fois, limitant sa capacité de transmission aux caractéristiques d'une bande unique. Avec le MLO, un dispositif compatible peut désormais basculer dynamiquement entre ces voies, ou même les utiliser simultanément, en fonction de la congestion et des besoins en bande passante. Cette flexibilité permet d'optimiser le flux de données en exploitant toujours le chemin le plus performant, assurant une connectivité stable et rapide, même dans des environnements denses.

Wi-Fi 6**Wi-Fi 7
MLO**

L'implémentation du MLO dans la norme Wi-Fi 7 poursuit trois objectifs majeurs :

- L'amélioration du Débit : en autorisant l'utilisation simultanée de plusieurs canaux sur différentes bandes de fréquences, le MLO permet d'agréger les débits, atteignant ainsi des performances bien supérieures aux générations précédentes. Les appareils peuvent ainsi atteindre des débits plus élevés en tirant parti de la capacité totale du spectre disponible.
- La réduction de la Latence : le MLO permet à un appareil de transmettre et de recevoir des données sur plusieurs bandes en parallèle, ce qui diminue considérablement les temps d'attente en cas de congestion sur une bande. Pour les applications en temps réel (comme les jeux en ligne ou les environnements industriels), cette réduction de la latence garantit une réactivité accrue et améliore l'expérience utilisateur.

- L'amélioration de la fiabilité : en permettant un basculement automatique des liens en fonction de l'état du réseau, le MLO assure une meilleure continuité des connexions. Par exemple, si une bande devient temporairement encombrée, un appareil peut basculer vers une autre bande sans interruption de service, garantissant ainsi une connectivité stable dans des environnements à forte densité de dispositifs, tels que les stades, les aéroports ou les conférences.

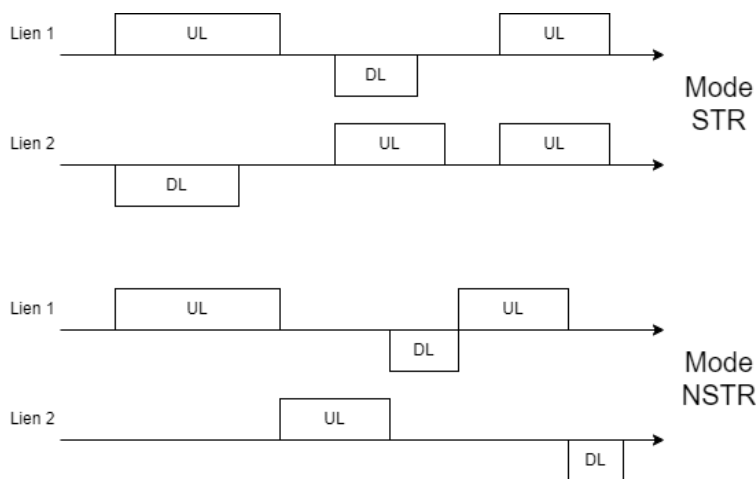
Afin de gérer efficacement cette utilisation multi-bande, le MLO introduit également le concept de *Multi-Link Device* (MLD), une architecture regroupant plusieurs interfaces radio sous une seule entité avec une adresse MAC unique. Cette structure simplifie la gestion des communications, permettant aux dispositifs de maintenir une connexion unique tout en exploitant plusieurs liens en parallèle.

Le MLO dans la norme Wi-Fi 7 propose deux modes principaux pour adapter son fonctionnement aux caractéristiques des appareils et aux contraintes du réseau :

- *Simultaneous Transmit and Receive* (STR) : ce mode permet la transmission et la réception simultanées de données sur plusieurs bandes de fréquences. Il offre un potentiel de débit maximal en exploitant toutes les bandes disponibles en parallèle, idéal pour les applications exigeant une bande passante élevée et une faible latence. Cependant, le STR introduit un défi technique important : les interférences croisées. Lorsqu'un dispositif utilise des canaux proches en fréquence, comme les bandes 5 GHz et 6 GHz, des interférences peuvent se produire, impactant la qualité des transmissions. Des solutions techniques supplémentaires telles que l'isolation des antennes et l'utilisation de filtres spécifiques sont nécessaires pour limiter ces interférences et tirer pleinement parti du mode STR.

- *Non-Simultaneous Transmit and Receive* (NSTR) : pour réduire le risque d'interférences, le mode NSTR impose une synchronisation des transmissions et des réceptions entre les différents liens, de sorte que les bandes ne puissent qu'émettre ou recevoir des données de façons simultanées. Bien que ce mode limite la flexibilité de l'appareil en termes de débit total, il est plus adapté aux appareils à faible consommation énergétique et dans des environnements où la gestion des interférences est cruciale. Le NSTR permet de garantir la qualité des transmissions même en présence de nombreux appareils connectés, en synchronisant les transmissions pour éviter les collisions et interférences.

Le schéma ci-dessous montre la différence entre les modes STR et NSTR :



L'adoption du MLO dans la norme Wi-Fi 7 apporte une flexibilité dans la gestion des transmissions sur plusieurs liens de fréquence simultanément. Cependant, tous les dispositifs ne disposent pas des mêmes capacités matérielles et leur implémentation du MLO dépend de leur architecture radio. En complément des modes de transmission STR et NSTR, les équipements peuvent être classés selon deux grandes catégories :

- les dispositifs à radio unique (*Single Radio*) ;
- les dispositifs multi-radio (*Multiple Radio*).

En fonction de cette classification, plusieurs variantes de MLO existent pour optimiser l'utilisation des ressources réseau en tenant compte des contraintes matérielles et énergétiques.

1.1 MLSR (Multi-Link Single Radio) - un seul lien actif

Le MLSR est une approche du MLO adaptée aux dispositifs équipés d'une seule radio. Un appareil MLSR ne peut activer qu'un seul lien à la fois et doit commuter dynamiquement entre les différentes bandes de fréquences en fonction des besoins.

Cette architecture présente l'avantage d'être énergétiquement efficace, car elle ne nécessite pas l'activation simultanée de plusieurs radios, ce qui est essentiel pour les appareils mobiles tels que les smartphones et les objets connectés. De plus, elle réduit la complexité matérielle, permettant une adoption plus large du MLO sans imposer des coûts excessifs aux fabricants d'équipements. Toutefois, le principal inconvénient du MLSR est le temps de transition entre les liens, qui peut impacter la fluidité des communications, notamment pour les applications nécessitant une latence minimale, comme les jeux en ligne ou la visioconférence.

1.2 EMLSR (Enhanced Multi-Link Single Radio) - un MLSR amélioré

Le EMLSR est une amélioration du MLSR qui optimise la gestion des transitions entre les liens pour minimiser la latence et améliorer la réactivité. Contrairement au MLSR « classique », un dispositif EMLSR peut écouter simultanément sur plusieurs liens tout en maintenant une seule transmission active. Cette capacité d'écoute parallèle lui permet d'anticiper les transitions et d'accélérer le basculement d'un lien à un autre en cas de congestion ou de dégradation du signal.

En réduisant le temps de latence lors des changements de lien, l'EMLSR devient particulièrement intéressant pour les applications où la continuité du signal est critique, comme le streaming vidéo, le cloud gaming ou les environnements de travail collaboratifs en temps réel. Néanmoins, cette amélioration s'accompagne d'une complexité logicielle accrue, nécessitant une coordination plus avancée entre le matériel et les algorithmes de gestion des liens.

1.3 MLMR (Multi-Link Multiple Radio) - l'exploitation simultanée de plusieurs liens

Contrairement aux dispositifs MLSR et EMLSR, qui doivent alterner entre plusieurs liens, les appareils basés sur une architecture Multi-Link Multiple Radio (MLMR) sont dotés de plusieurs radios physiques, leur permettant d'exploiter plusieurs liens simultanément. Cette architecture est particulièrement avantageuse pour maximiser les performances et réduire la latence en évitant les interruptions liées aux commutations de lien.

Les appareils MLMR peuvent être classés en deux sous-catégories selon leur mode de gestion des transmissions :

- MLMR STR, les radios peuvent transmettre et recevoir en parallèle sur plusieurs liens.
- MLMR NSTR, plusieurs radios sont présentes, mais leur utilisation est alternée pour éviter les interférences.

1.4 MLMR STR - la performance maximale

Le MLMR STR (*Multi-Link Multiple Radio Simultaneous Transmit and Receive*) est l'implémentation la plus avancée du MLO en termes de performances. Grâce à la présence de plusieurs radios indépendantes, un équipement de ce type peut émettre et recevoir sur plusieurs liens en simultané, permettant ainsi d'agréger les débits et de réduire considérablement la latence.

Cette architecture est idéale pour les applications nécessitant un débit élevé et une latence ultra-faible, telles que la réalité virtuelle (VR), les jeux en ligne compétitifs, ou encore les applications industrielles en temps réel. En exploitant plusieurs bandes de fréquences simultanément, un terminal peut contourner les congestions réseau, garantissant ainsi une connectivité stable et rapide. Toutefois, cette approche consomme davantage d'énergie, ce qui peut limiter son adoption sur des appareils mobiles fonctionnant sur batterie. De plus, la présence de plusieurs radios pose des défis en termes de gestion des interférences et de coordination des bandes ce qui la rend très difficile à mettre réellement en œuvre.

1.5 MLMR NSTR - une alternative plus accessible

Le MLMR NSTR (*Multi-Link Multiple Radio Non-Simultaneous Transmit and Receive*) représente une alternative au STR, permettant d'exploiter plusieurs radios tout en évitant de les faire fonctionner en parallèle. Dans cette configuration, un appareil peut transmettre sur un lien et recevoir sur un autre, mais il ne peut pas effectuer ces opérations simultanément.

L'avantage principal de cette approche est qu'elle limite les interférences intra-appareil, qui peuvent survenir lorsque plusieurs radios émettent sur des bandes proches. Elle est particulièrement utile dans les environnements Wi-Fi denses, comme les bureaux, les centres commerciaux ou les événements où de nombreux appareils sont connectés en simultané. Par ailleurs, en évitant l'activation simultanée de plusieurs radios, cette architecture réduit la consommation énergétique par rapport au MLMR STR, tout en offrant une meilleure gestion des ressources spectrales. Cependant, le principal compromis du MLMR NSTR est que les performances en débit et en latence sont inférieures à celles du STR, car les transmissions doivent être alternées, ce qui peut créer des délais dans l'acheminement des données.

L'adoption du MLO dans la norme Wi-Fi 7 ouvre des possibilités étendues pour des applications modernes, particulièrement celles nécessitant une réactivité extrême et une bande passante élevée :

- La réalité augmentée et virtuelle : ces technologies, utilisées dans les domaines du divertissement, de la formation et de l'industrie, nécessitent une faible latence pour éviter tout décalage entre les actions de l'utilisateur et la réponse visuelle. Le MLO, en réduisant la latence et en assurant une connectivité ininterrompue, améliore grandement l'expérience utilisateur.
- Les jeux en ligne et cloud gaming : les jeux en ligne, en particulier ceux de type cloud gaming, dépendent de connexions rapides et stables pour une expérience fluide. Le MLO permet de maintenir des taux de transmission élevés tout en réduisant les interruptions dues à la congestion du réseau.
- Les applications industrielles : dans les environnements industriels, où la fiabilité de la connexion est cruciale pour la sécurité et l'efficacité, le MLO garantit une continuité de service en basculant automatiquement entre les liens en cas de congestion ou d'interférence.

En offrant la possibilité d'exploiter plusieurs bandes de manière simultanée ou adaptative, le *Multi-Link Operation* transforme l'expérience réseau, que ce soit pour les utilisateurs quotidiens cherchant une meilleure qualité de service ou pour les applications industrielles nécessitant des performances extrêmes, le MLO ouvre une nouvelle ère de connectivité sans fil, assurant flexibilité, robustesse et efficacité dans les environnements les plus exigeants. Tout au long de ce chapitre, nous allons explorer en profondeur l'implémentation du Multi-Link Operation. Nous analyserons les différents modes de fonctionnement, les mécanismes de gestion des interférences et les méthodes d'allocation des ressources qui permettent au MLO de répondre aux besoins spécifiques des environnements modernes. Nous examinerons également les défis techniques à surmonter pour en tirer pleinement parti, afin de fournir une compréhension complète de ces capacités.

2. Architecture des dispositifs Multi-Lien (MLD)

Dans les générations précédentes des normes Wi-Fi, chaque client connecté devait se limiter à l'utilisation d'une seule bande de fréquence à la fois, entraînant des limitations en termes de débit, de latence et de continuité de service, particulièrement en cas de congestion ou de perturbation réseau. Cette contrainte est devenue un obstacle pour répondre aux besoins croissants des applications modernes, qui exigent une connectivité rapide et stable, même dans des environnements densément peuplés. Pour surmonter ces limites, la norme Wi-Fi 7 introduit le Multi-Link Device, une nouvelle architecture conçue pour exploiter simultanément plusieurs bandes de fréquences, offrant ainsi une connectivité optimisée.

Dans cette section, nous examinerons comment les MLD permettent aux appareils de bénéficier d'une gestion des liens plus intelligente et dynamique, assurant des performances élevées et une meilleure résilience. Nous détaillerons également la structure des couches MAC et PHY de cette architecture, qui centralise les connexions multi-lien sous une adresse MAC unique, simplifiant ainsi la gestion et réduisant les interruptions de service. Cette architecture marque une étape cruciale dans l'évolution des réseaux Wi-Fi, ouvrant la voie à des connexions plus fiables et performantes dans les environnements réseau modernes.

2.1 Présentation des dispositifs Multi-Lien (MLD)

Traditionnellement, les appareils Wi-Fi se connectent via un seul lien sur une bande de fréquence donnée (par exemple, 2,4 GHz, 5 GHz ou 6 GHz), chaque bande ayant ses avantages et ses inconvénients en termes de portée, de débit et de sensibilité aux interférences. Cependant, cette approche à lien unique présente des limitations qui deviennent de plus en plus problématiques pour répondre aux besoins actuels. Par exemple, si un canal dans la bande 5 GHz devient encombré en raison d'un grand nombre de connexions, l'appareil ne peut rien faire pour contourner cette congestion autrement qu'en basculant manuellement vers un autre canal. Cependant, ce changement peut entraîner des interruptions, de la latence et des déconnexions temporaires, rendant l'expérience utilisateur frustrante pour des applications sensibles à la continuité de service.