

N° d'ordre : 08/2005-M/IN

REPUBLIQUE ALGERINNE DEMOCRATIQUE ET POPULAIRE
MINISTERE D'ENSEIGNEMENT SUPERIEUR ET DE LA RECHERCHE SCIENTIFIQUE
UNIVERSITE DES SCIENCES ET DE LA TECHNOLOGIE
« HOUARI BOUMEDIENNE »
FACULTE DE L'ELECTRONIQUE ET INFORMATIQUE
Département d'Informatique



Mémoire

Présenté pour l'obtention du diplôme de Magister

En : INFORMATIQUE

Spécialité : Programmation et

Par : ZERAOULIA Khaled

Sujet :

Transmission vidéo sur des liaisons sans fil

Soutenue le 03/07/2005, devant le Jury composé de :

Mme. A. AISSANI, Professeur, U.S.T.H.B

Président

Mr. N. BADACHE, Professeur, U.S.T.H.B.

Directeur de thèse

Mr. A. BELKHIR, Maître de conférence, U.S.T.H.B

Examineur

Mme. M. BOUKALA, Maître de conférence, U.S.T.H.B

Examineur

Remerciements

Je tiens à remercier le tout puissant de m'avoir donné le courage et la patience jusqu'à l'achèvement de travail.

J'exprime ma profonde reconnaissance et mes vifs remerciements à mon directeur de thèse Mr. N. BADACHE pour m'avoir encadrer. Je le remercie également pour les lectures attentives de mes rapports et dont les critiques et suggestions ont été d'un grand apport pour la finalité de ce travail.

Je remercie Mme. A. ISSANI, Mr. A. BELKHIR, Mme. M. BOUKALA d'avoir accepté de juger ce travail.

Mes remerciements s'adressent aussi à ma famille, mes amis et tous ceux qui m'ont aidé et contribué de près ou de loin par une quelconque façon, que ce soit par leur amitié, leurs conseils ou leur soutien moral ou matériel.

Table des matières

I. INTRODUCTION	6
1. LES RÉSEAUX SANS FIL IEEE 802.11	8
1.1 Introduction	8
1.2. Problèmes des transmissions radios et des réseaux sans fil	9
1.3. Architecture IEEE 802.11	10
1.3.1. Description des couches physiques, MAC, Réseau et transport.....	10
1.3.2. Physique.....	11
1.3.3. Liaison MAC.....	13
1.3.4. Synchronisation.....	19
1.3.5. L'économie d'énergie.....	19
1.4. Avantages	20
1.5. Conclusion	20
2. LA COMPRESSION VIDÉO	21
2.1. Introduction	21
2.2. La compression MPEG	21
2.2.1. Décomposition d'un fichier MPEG	23
2.2.2. Compression temporelle.....	23
2.2.3. Compression spatiale.....	26
2.3. Evolution des standards MPEG	26
a. MPEG1.....	26
b. MPEG2.....	27
c. MPEG3.....	28
d. MPEG4.....	28
e. H264/ MPEG4 AVC.....	31
2.4. Conclusion	33
3. MOBILITÉ IP	34
3.1. Introduction	34
3.2. Problème de la mobilité IP	34
3.3. Mobile IP	35
3.3.1. Définitions.....	35

3.3.2. Fonctionnement du protocole.....	36
3.4. MIPv6	42
3.4.1. Fonctionnalités requises.....	42
3.4.2. Découverte des routeurs d'accès.....	42
3.4.3. Enregistrement.....	42
3.5. Limites de Mobile IP.....	42
3.6. Conclusion	46
4. Transmission du trafic vidéo MPEG sur un réseau WLAN 802.11b.....	47
4.1. Introduction	47
4.2. La diffusion de la vidéo sur un Wlan	47
4.2.1. La préparation des flux MPEG	47
4.2.2. RTP et RTCP	48
4.3. Caractéristique du trafic vidéo encodé	49
4.4. Mesure de la qualité de la vidéo	50
4.4.1. Introduction	50
4.4.2. Indicateurs objectifs	50
4.5. Les caractéristiques des canaux sans fil	51
4.5.1. Impact des erreurs des canaux sans fil sur la qualité vidéo	51
4.6. Mécanisme vidéo streaming.....	51
A. Entrelacement des données codage robuste	52
B. Contrôle de pertes.....	52
C. Contrôle de congestion et de débit.....	53
4.7. Le modèle de simulation	53
4.7.1. Génération de la trace d'une source vidéo.....	56
4.7.2. Topologie utilisée.....	56
4.7.3. Analyse des résultats de simulation	57
4.8. Conclusion	61
5. Etude de performance et l'impact du handover 802.11 sur le trafic vidéo.....	62
5.1. Introduction	62
5.2. Principe de base de cette partie	63
5.3. Related work	63
5.4. Procédure handover couche liaison.....	64
5.5. Conception des PAs	69
5.6. Envoi de Flux ascendant et descendant durant le handover L2	70
5.6. 1. Envoi de flux ascendant.....	71
5.6.2 Envoi de flux descendant	73

5.6.3. Evaluation de la portée des PA.....	74
5.7. Conclusion	76
6. La gestion du handover couche IP (macro mobilité).....	77
6.1. Introduction	77
6.2. Les diverses techniques de handover IP.....	79
6.2.1. Hierarchical mobile IP	80
6.2.2. Smooth-handover.....	81
6.2.3. Fast-handover	83
6.3. Analyse générale des diverses techniques de handover	84
6.4. Conclusion.....	86
7. Modèle générique pour le Contrôle du handover IP.....	87
7.1. Introduction	87
7.2. Les besoins du modèle générique pour le contrôle de mobilité	87
7.3. Latence du handover pour le rapport des applications temps réel	87
7.4. Modèle générique du handover IP	89
7.4.1. Introduction.....	89
7.4.2 Les étapes de contrôle du handover dans le modèle générique	89
7.4.3. Les buts de la conception du modèle générique	90
7.4.4. Définition des différents composants du modèle générique	90
7.5. Conclusion	96
8. GHMM (Generic Hierarchical Mobility Managment).....	97
8.1. Introduction.....	97
8.2. Macro-mobilité IP.....	98
8.3. Protocoles avec une architecture basée sur une hiérarchie d'agents de mobilité.....	98
8.4. Limites des protocoles de macro-mobilité.....	98
8.5. GHMM (Genric Hierarchical Mangement Mobility).....	99
8.5.1. Architecture considérée.....	99
8.5.2. Descriptions des techniques et des composants du GHMM.....	101
8.5.3. Comportement du GHMM	110
8.5.4. Gestion du handover dans GHMM.....	112
8.6. Conclusion	113

CONCLUSION115

Bibliographie117

ANNEX A123

ANNEX B129

Table des figures

Figure 1.1- Architecture.....	10
Figure 1.2- Architecture de la couche 1 et 2.....	10
Figure 1.3- Architecture détaillée de la couche 1 et 2	11
Figure 1.4- Méthodes DSSS et FHSS	12
Figure 1.5- Explication du Backoff exponentiel (sans prendre en compte les ACK).....	15
Figure 1.6- Backoff exponentiel + ACK	15
Figure 1.7- La station émettrice (A) envoie un RTS	17
Figure 1.8- Réponse de la station réceptrice- Envoi d'un CTS	17
Figure 1.9- Station cachée	18
Figure 1.10- NAV.....	18
Figure 2.1- Luminance – exemple	22
Figure 2.2- Schéma IBP utilisé dans les normes Vidéo MPEG	24
Figure 2.3- Estimation et compensation de mouvement	25
Figure 2.4- Exemple d'une scène MPEG4.....	30
Figure 2.5- Structure d'une scène MPEG-4.....	31
Figure 2.6- Evolution de la compression vidéo	32
Figure 3.1 - Encapsulation IP dans IP	37
Figure 3.2- Enregistrement	39
Figure 3.3- Dialogue avec un Mobile.....	40
Figure 3.4- Exemple de routage	41
Figure 3.5 - Routage Optimisé	42
Figure 3.6- Communication V 6.....	44
Figure 4.1- Illustration de processus d'entrelacement	52
Figure 4.2- Modèle de simulation	54
Figure 4.3- Diagramme de flux de contrôle et de données.....	55
Figure 4.4- La topologie de simulation	57
Figure 4.5- Retard des paquets en fonction des séquences vidéo (retard 90%).....	58
Figure 4.6- Retard des paquets en fonctions des séquences vidéo.....	58
Figure 4.7- Qualité vidéo PSNR en fonction des séquences vidéo	58
Figure 4.8- PSNR en fonction du retard des paquet.....	59
Figure 4.9 -PSNR en fonction des tailles des paquets.....	59
Figure 4.10-PSNR en fonction des probabilités d'erreurs	59
Figure 5.1-La topologie proposée pour l'exécution du handover L2	65

Figure 5.2-La procédure de roaming	66
Figure 5.3 Handover L2 classique	67
Figure 5.4 Handover L2 dans notre étude	68
Figure 5.5-La procédure du handover pour un trafic upstream	72
Figure 5.6-La procédure du handover un trafic downstream	73
Figure 6.1- Interaction entre un terminal mobile et le cœur de réseau	78
Figure 6.2- La latence du handover	78
Figure 6.3- Diagramme temporel pour le smooth handover	82
Figure 6.4- Diagramme temporel pour le fast handover	84
Figure 6.5- Stratégie du Handover.....	85
Figure 7.1- Définition de la latence handover mobile IP	87
Figure 7.2- Modèle générique du handover IP	89
Figure 7.3- Différentes scénarios du registration	92
Figure 8.1- Architecture proposée	100
Figure 8.2- Architecture de HMIPv6.....	101
Figure 8.3- Procédure d'anticipation	103
Figure 8.4- Schéma buffering optimal	108
Figure 8.5- Le Bicating avec pré- registration	110
Figure 8.6- Comportement d'un nœud mobile dans un domaine GHMM	111
Figure 8.7- Gestion du handover dans le GHMM	112
Tab 6.1 Notations générales.	80
Tab 8.1 Procédure du pré- registration	113

1. Introduction

Durant les dernières années, nous avons été témoins d'une croissance spectaculaire dans l'intérêt mondial pour les communications mobiles. La demande croissante en services et le développement de la technologie sans fil continueront leur progression dans les années à venir. Avec l'augmentation du nombre d'utilisateurs mobiles et l'évolution rapide des réseaux mobiles sans fil ainsi que les applications multimédia en temps réel, les demandes des utilisateurs en terme de qualité de service **QOS** deviennent de plus en plus exigeantes. Ceci nécessite l'introduction des méthodes et techniques pour l'amélioration et la gestion des communications mobiles dans les **WLANs (Wireless Local Area network) (802.11b)**.

Dans cette optique nous allons décrire nos travaux qui portent sur la transmission de la vidéo sur des réseaux sans fil et mobiles.

En réalité, La transmission vidéo sur un WLAN 802.11b est un problème critique dans les réseaux locaux sans fil. De nombreux travaux [Fen04] [Wan99][You03] ont déjà été menés sur ce problème, cependant il est loin d'être résolu convenablement dans toutes ses formes.

La vidéo est caractérisée par le volume important des données qu'elle occupe. Or les liens constituant le réseau WLAN présentent des capacités de transmission disparates et parfois très limitées (plus de 10Mbit/s pour des réseaux locaux, 2 Mbit/s pour un lien ADSL, 64Mbit/s pour un lien sans fil GSM, 56 ou 33 Kbit/s pour un modem). Par conséquent, il est nécessaire de réduire fortement le débit de données pour pouvoir les transmettre. Ceci est l'objectif du codage de source, chargé de supprimer les redondances et impertinences contenues dans la source vidéo, au prix d'une baisse de la qualité de représentation à des taux de compression élevés. Or, un canal de communication introduit généralement des erreurs dans les données transmises, avec une certaine probabilité. Pour faire face à ces erreurs, les systèmes de transmission classiques consacrent une partie du débit transmis à des données de redondance. L'introduction de ces données de redondances, qui permettent de réduire la probabilité d'erreur résiduelle, constitue l'étape d'amélioration de transmission sur le canal sans fil. A cet effet, nous avons mené des études portant sur la qualité vidéo perçue lors d'une communication mobile dans les réseaux Wlan 802.11. Une étude a été faite sur des améliorations possibles des protocoles et des techniques de la couche liaison et la couche réseau (IP) en vue de minimiser la perte et le retard des paquets durant le mécanisme du processus de changement de cellule (handover) déclenché lors du mouvement d'un mobile. Plus précisément le handover au niveau de la couche MAC et le handover au niveau de la couche IP.

D'autre part, la gestion de la mobilité constitue un défi technique à relever. En effet, un protocole de mobilité efficace doit pouvoir permettre l'exécution des applications d'une manière transparente au mouvement de l'utilisateur.

Les travaux réalisés portent d'une part sur la transmission vidéo sur un réseau local sans fil (WLAN802.11) (définir les modèles et mécanismes nécessaires pour une telle transmission) et d'autre part sur l'impact de la mobilité (handoff) et la nature sans fil sur cette transmission

Cette thèse est composée de huit chapitres:

Le chapitre 1 introduit les réseaux sans fil et les principaux concepts liés à ces nouveaux réseaux, ce chapitre présente une description technique des réseaux sans fil 802.11(la couche physique , la couche Mac ainsi que les différents schémas d'accès sans fil).

Le chapitre 2 présente la notion de la compression vidéo. Nous étudions la norme de compression MPEG en détail. En premier lieu nous illustrons la décomposition d'un fichier vidéo MPEG, ainsi que les deux techniques de base de la compression MPEG : la compression temporelle et la compression spatiale. Ensuite nous présentons l'évolution du standard MPEG durant ces dernières années. En dernier lieu nous détaillons la version la plus récente du standard MPEG4.

Le chapitre 3 consacré aux réseaux mobiles, nous définissons d'abord le fonctionnement et les caractéristiques du protocole mobile IP, tel que le routage et l'encapsulation. Ensuite nous étudions le protocole MIPv6 et ces différents mécanismes de routage, d'enregistrement, de découvertes des routeur d'accès.

Le chapitre 4 présente notre première contribution pour l'étude de la transmission vidéo sur un réseau sans fil 802.11 (le cas où le nœud mobile ne change pas de points d'accès), le but de cette partie est d'étudier l'impact de la nature sans fil sur une transmission vidéo. Nous commençons d'abord par la définition de la transmission vidéo sur un Wlan ainsi que les caractéristiques de la vidéo encodée et les canaux sans fil. Après nous définissons un modèle de simulation pour la transmission vidéo sur un Wlan 802.11b. Ensuite nous nous basons sur ce modèle pour étudier l'impact de la nature sans fil sur la qualité vidéo perçue.

Le chapitre 5 présente l'étude de l'impact du handover couche liaison sur la transmission vidéo (le cas où un nœud mobile change de point d'accès sans changer de routeur d'accès). Nous étudions en détail dans ce cas le processus du handover couche liaison (procédure de changement de point d'accès) durant la transmission vidéo et nous proposons d'optimiser cette procédure du handover en ajoutons les techniques de smoothing et la notion du soft handover.

Le chapitre 6 consacré à l'étude des différentes extensions du protocole MIPv6. Le but de ce chapitre est de cadrer les différents axes que nous pouvons les exploiter pour proposer une extension optimale de MIPv6 qui nous permet d'exécuter un handover couche IP (c'est le processus de changement de routeur d'accès) d'une façon efficace

Chapitre 7 consacré pour notre contribution pour la proposition d'un modèle générique pour la gestion du handover couche IP. C'est un modèle à base de composants ou chaque composant permet d'améliorer la gestion de la mobilité (handover couche IP). Dans ce chapitre nous avons présenté en premier lieu les besoins d'un tel modèle. Ensuite nous étudions les différents composants du modèle générique.

Chapitre 8 est consacré à notre proposition d'un protocole de gestion du handover couche IP GHMM((*Generic Hierarchical Mobility Management*)) , ce protocole **GHMM** (basée sur ce modèle et le contrôle combiné du handover par **NM** (Nœud Mobile) et le **MAP** (Agent de Mobilité) ; ie : le nœud mobile initie le handover et les différents composants du modèle générique contrôlent toutes les étapes du handover. Le protocole que nous avons proposé se caractérise par l'optimisation de la perte et le retard ainsi que la charge de signalisation causés par la transmission vidéo sur un Wlan 802.11

Chapitre 2

La Compression Vidéo

2.1. Introduction

Ce chapitre présente un état de l'art sur le codage vidéo pour la transmission vidéo sur un réseau Wlan. Plusieurs standards de compression existent. nous étudions dans ce qui suit le standard MPEG, qui est un des standards les plus connus dans le domaine de la compression [Hur03].

MPEG est l'acronyme de **M**oving **P**icture **E**xpert **G**roup, un groupe formé sous les auspices de l'Organisation Internationale de Normalisation (**ISO**) et de la Commission Internationale d'Electrotechnique (**IEC**). Il s'agit d'un groupe d'industriels qui se réunit afin d'élaborer des standards dans le domaine des compressions de la vidéo et de l'audio numériques. Ils définissent en particulier des flux binaires compressés desquels découlent implicitement des décodeurs. Cependant les algorithmes utilisés afin d'obtenir ce format de flux binaire compressé est laissé au bon vouloir des implanteurs de la norme [Hur03].

2.2. La compression MPEG

L'intérêt de la compression des données informatiques est immédiat : la taille des données s'en trouve réduite, ce qui implique en particulier moins de place de stockage utilisée et moins de transfert réseau. Dans une application vidéo, on peut se permettre d'utiliser une compression avec pertes. En effet il n'est pas nécessaire d'avoir une qualité de vidéo parfaite pour que l'information soit exploitable. L'information peut donc être dégradée dans une certaine mesure tout en restant acceptable [Hur03].

Un des formats de compression les plus utilisés en vidéo est le format **MPEG** (**M**otion **P**ictures **E**xpert **G**roup) qui s'appuie sur une compression spatiale des images et une compression temporelle inter-images. Ce type de compression a fait preuve de son efficacité par rapport à d'autres mécanismes. Par exemple, une compression sans perte comme la compression de type Huffman [Try96] n'est pas suffisante. Elle est basée sur le nombre d'occurrences d'un caractère dans une série de données. Plus le caractère en question est fréquent, plus la séquence l'encodant sera petite. La décompression se fait par l'application de la grille d'encodage inverse et se fait donc sans perte. Il est logique qu'une telle compression soit moins efficace qu'une compression présentant une perte

d'informations acceptable. D'autres compressions se basent sur la compression image par image, indépendamment, de la vidéo [Hur03].

C'est le cas par exemple du format MJPEG. Cette compression n'est pas optimale non plus car la redondance entre des images d'une même scène vidéo est forte, et ce fait peut être exploité, ce qui est fait avec la compression MPEG. Voyons en détail les mécanismes utilisés dans la compression MPEG.

2.2.1. Décomposition d'un fichier MPEG

Les données dans une séquence vidéo sont hiérarchisées. La séquence vidéo commence par une entête contenant des données d'identification et d'encodage. Elle contient un ou plusieurs groupes d'images, et se termine par un code permettant de détecter la fin de la séquence. Un groupe d'images (Group of Pictures - GOP) contient une entête et plusieurs images. C'est l'entité qui permet l'accès aléatoire à une partie de la vidéo.

Une image est l'entité élémentaire encodée. Elle est composée de trois matrices Y, Cr et Cb comme l'explique [Try96] (ou elles sont notées Y, U, V). Y correspond à la luminance de l'image : c'est en quelque sorte la quantité de lumière de la couleur, une « distance » de la couleur par rapport au noir. Ainsi, sur la **Fig 2.1** on peut voir le même rouge à des luminances croissantes. Cr et Cb sont des composantes chromatiques.

On peut estimer qu'une couleur peut être repérée suivant les deux axes rouge/vert et bleu/jaune. Cr représente la distance de la couleur à la couleur rouge. Cb représente la distance de la couleur à la couleur bleue. L'intérêt de cette représentation chromatique pour la compression est que l'oeil humain est moins sensible aux légères différences au niveau de la chrominance qu'à des différences équivalentes au niveau de la luminance. Il est donc possible de représenter les matrices Cr et Cb moins finement que la matrice Y. Ceci est un avantage par rapport au système RGB (Voir aussi [Try96] pour ces différents systèmes).

(red green blue - rouge vert bleu) qui lui réclame la même quantité de données pour les trois composantes de la couleur.

Ces images sont elles-mêmes découpées en tranches (slice), qui sont principalement Utilisées dans la gestion d'erreurs (de compression, de décompression, de transmission).

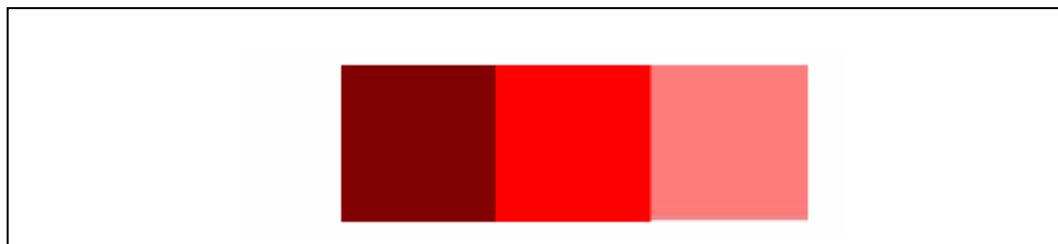


Fig 2.1. Luminance : exemple

Les tranches sont composées de macroblocs qui regroupent eux-mêmes des blocs de pixels représentés par leurs valeurs Y, Cr et Cv. Les blocs sont de dimension 8x8 pixels, les macroblocs sont de dimension 2x2 blocs.

2.2.2 Compression temporelle

La compression temporelle repose sur le fait que la différence entre deux images consécutives d'une vidéo est faible. Ceci est exploité en considérant des groupes d'images, les images étant de types différents. Il existe trois types d'images dans le format MPEG :

- les images I (I-frames pour intracoded frames -images intracodées).
- les images P (P-frames pour predicted frames - images prédites).
- les images B (B-frames pour bidirectional frames - images bidirectionnelles).

Dans ce qui suit nous présentons les différentes techniques normalisées de compression vidéo :

A- Le schéma IPB

Pour atteindre des taux de compression élevés, on exploite la redondance spatiale et temporelle contenue dans le signal. Pour cela, les schémas de compression vidéo standard définissent les trois types d'images I, P, B (**Fig 2.2**).

Les images I sont encodées indépendamment du reste de la vidéo. L'image compressée contient de l'information absolue, plus ou moins dégradée en fonction de la qualité finale voulue de la vidéo.

Les images P sont encodées par rapport à des images de référence : on calcule leur différence par rapport à cette référence. L'encodage se fait ici au niveau des macroblocs.

Si un macrobloc correspond à un élément nouveau par rapport à l'image de référence (arrivée d'un élément nouveau dans une scène, changement de scène), le macrobloc inconnu est encodé comme un macrobloc d'une image I. Sinon, le mouvement entre l'image de référence et l'image actuelle est prévisible.

Les données sont alors stockées sous la forme d'un vecteur de mouvement (motion vector) qui détermine la position du macrobloc par rapport à son équivalent dans l'image de référence. Il est à noter que cette image P peut être utilisée comme référence pour un autre calcul relatif. L'inconvénient est donc la propagation inévitable des erreurs.

Les images de type B sont encodées avec une double référence : une image précédente et une image suivante, selon le même principe que les images P. Les images B ne sont pas utilisées comme référence pour un autre calcul relatif.

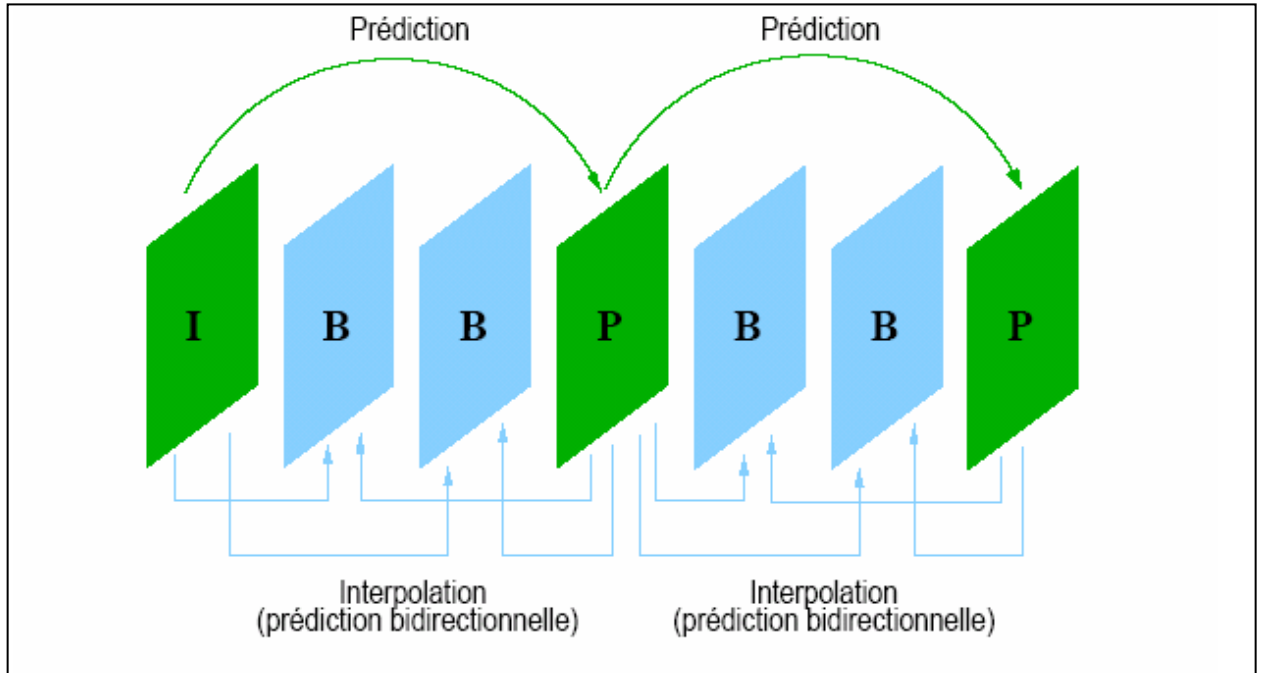


Fig 2.2. Schéma IBP utilisé dans les normes Vidéo MPEG

A.1 Codage des images intra (I)

Le codage des images intra est similaire à celui de la norme de compression d'images fixes avec pertes JPEG **[Jpeg]**. Une Transformée en Cosinus Discrète (DCT) est appliquée, tout d'abord, à chacun des blocs 8x8 de pixels et les convertit en blocs de coefficients de fréquences spatiales. C'est l'étape de corrélation spatiale **[Fab01]** Une quantification scalaire des coefficients DCT est ensuite effectuée. Pour chaque composante fréquentielle, le pas de quantification est fonction de la sensibilité de l'œil à la fréquence considérée. Les pas de quantification ont été définis afin de prendre en compte les caractéristiques du système visuel humain, plus sensible aux basses qu'aux hautes fréquences (les pas de quantification sont donc plus élevés pour les coefficients de hautes fréquences).

A.2. Codage des images prédites (P)

Le codage des images P consiste en un codage prédictif temporel utilisant une estimation/compensation de mouvement basé bloc.

- **Estimation et compensation de mouvement**

Les standards vidéo opèrent une compensation de mouvement par bloc, ou blockmatching. Celle-ci est illustrée sur la Figure suivante. A chaque image P est attribuée une image de référence, le plus souvent égale à l'image I ou P qui précède, reconstruite après codage et décodage.

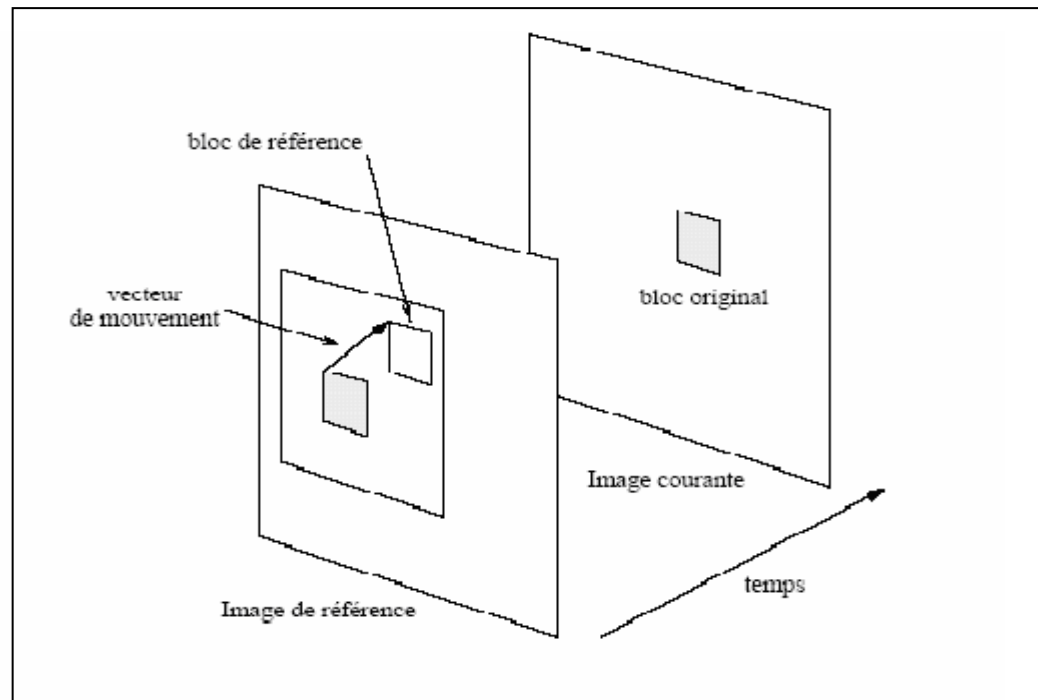


Fig 2.3. Estimation et compensation de mouvement

Pour chaque bloc de l'image courante, on cherche, dans l'image de référence, le bloc qui lui ressemble le plus au sens de la somme des différences absolues, calculée sur les pixels de luminance. La recherche est limitée à une fenêtre (voir **Fig. 2.3**) de taille variable, et le bloc trouvé est appelé le bloc de référence. Notons que la taille de la fenêtre influe sensiblement sur la rapidité du codeur vidéo. Un vecteur de mouvement, résultant de cette recherche, est transmis au décodeur, ainsi que l'erreur résiduelle, différence entre le bloc perdu et le bloc de référence. Le codage de l'erreur se fait par transformation par DCT, suivie d'une quantification scalaire et du codage statistique, selon une approche identique à celle employée pour les images I. L'image codée est reconstruite et mémorisée pour servir de référence pour la prédiction d'images futures. [Fab01].

A.3. Codage des images bidirectionnelles (B)

Le codage des images B est similaire à celui des images P. Néanmoins, les blocs des images B sont prédits à partir de deux images de référence respectivement antérieures et postérieures à l'image courante.

Chaque bloc peut donc avoir 3 blocs de référence possibles:

- bloc appartenant à l'image de référence antérieure.
- bloc appartenant à l'image de référence postérieure.
- interpolation des deux premières possibilités.

2.2.3 Compression spatiale

Les images de type I sont les plus volumineuses car elles contiennent plus d'information, c'est pourquoi il est nécessaire de les compresser beaucoup. Le principe général de cette compression est que l'oeil humain est plus sensible au bruit dans les fréquences spatiales basses que dans les fréquences spatiales hautes. Pour tenir compte de ce fait, l'image doit être décomposée fréquentiellement, ce qui se fait par le moyen d'une transformée discrète en cosinus bidimensionnel (DCT – Discrete Cosine Transform : Transformée en Cosinus Discrète). Pour chaque bloc (i; j) on calcule

$$D(i, j) = \frac{1}{4} C_i * C_j * \sum_{x=0}^{N-1} \sum_{y=0}^{N-1} S_{xy} \cos\left(i\pi \frac{2x-1}{2N}\right) \cos\left(j\pi \frac{2y-1}{2N}\right)$$

$$C_n = \begin{cases} \frac{1}{\sqrt{2}} & \text{si } n = 0 \\ 0 & \text{sinon} \end{cases}$$

Avec N la taille du bloc (soit i) et S x y la valeur (Y, Cb ou Cr) de la matrice considérée au point (x; y) du bloc. En pratique, ce n'est pas cette formule qui est directement utilisée, mais une version optimisée (algorithmiquement parlant) pour accélérer la compression.

Ceci en soit n'est pas une compression : on passe d'une valeur entière à une valeur décimale, ce qui est a priori plus volumineux. Il faut donc quantifier les valeurs de DCT obtenues suivant une grille de quantification. Cette grille sera plus fine en hautes fréquences, plus large en basses fréquences. La quantité d'information est donc réduite [Hur03].

Une compression standard (de type Huffman par exemple) permet de compresser sans perte cette information. Globalement, la perte d'information de la compression est uniquement liée à la quantification des composantes fréquentielles de l'image.

Pour décompresser l'image, il suffit de décompresser les données quantifiées, d'appliquer la grille définie (et présente dans les informations du flux vidéo), et d'appliquer une transformée discrète en cosinus inverse. Les images de type P sont ensuite reconstruites à partir des images décompressées (ou des images P précédentes) et des informations contenues dans les macroblocs. Finalement, les images B sont reconstruites à partir des images existantes.

Les sections suivantes présentent les spécificités de chacun de standards MPEG-4.

2.3. Evolution des standards MPEG

a. MPEG1

MPEG1 est le premier essai du groupe MPEG afin de définir un format efficace de compression vidéo. Le standard MPEG1 est un ensemble de 5 publications de l'Organisation Internationale de Normalisations, sous la dénomination ISO/IEC 11172 Part 1 à 5. Les 3 premières publications ont été écrites en 1993, ils définissent les flux binaires vidéo, audio ainsi que leur

multiplexage pour aboutir à un flux multimédia complet. La 4ème publication en 1995, standardise une plate-forme de tests afin de s'assurer de l'interopérabilité des matériels et logiciels se voulant être compatibles MPEG1.

En 1998, l'ISO publie la part 5 du standard, mais il ne s'agit pas d'un texte, c'est en fait un logiciel de référence publié sous forme de code source. C'est de lui que dérivent tous les logiciels les plus utilisés par le public.

Le standard MPEG1 est prévu pour un vitesse de flux autour de 1,5MB pour une image n'excédent pas les 352x240. Cela correspond approximativement à une qualité équivalente à du VHS. On pourra noter que le MPEG1 a surtout été utilisé dans les pays d'Asie où il a su s'imposer à travers le format VCD (Vidéo CD) très employés entre autre dans les karaokés.

Cependant MPEG1 a très vite montré ses limites. Dès lors que le matériel a permis de traiter des images de taille plus conséquente, MPEG1 n'offrait plus les moyens techniques de garantir une qualité suffisante [Her02].

b. MPEG2

Voyant que le MPEG1 ne saurait satisfaire les besoins de la vidéo numérique (haute définition, haute qualité), le groupe MPEG a entamé une autre collaboration afin d'établir les bases d'un standard capable de supporter des hauts débits ainsi qu'une image de plus haute résolution.

Ce standard a été progressivement publié de 1998 à 2000, un peu à la façon de MPEG1 en publiant divers écrit traitant chacun d'une partie du standard. Il définit une nouvelle génération de compression "MPEG" qui délivre une qualité supérieure et se montre plus efficace pour les hauts débits et grandes images.

MPEG2 intègre aussi le support de formats d'image professionnels tel que le format YUV422. Il introduit aussi de nouveaux concepts de compression tels que les références temporelles bidirectionnelles et le support natif des séquences entrelacées (très utilisées par/pour la télévision). MPEG2 apporte aussi énormément au niveau de la couche transport/multiplexage et en fait un candidat idéal pour des flux multimédias complexes avec une interactivité réduite. Sa fourchette de débit idéal se situe entre 1.5MB/set 15MB/s. L'un des principaux problèmes lié à la diffusion de MPEG1 est sa totale absence de contrôle de flux. C'est à dire qu'il se contente d'être émit sans se soucier des capacités de l'équipement terminal. Le MPEG2 permet un contrôle de flux en proposant au niveau de son codage différentes couches

- une couche de base avec une qualité restreinte,
- une ou plusieurs couches supplémentaires d'améliorations.

Ceci permet d'avoir plusieurs films vidéos transmis dans les mêmes trames (suite de bits envoyés sur le réseau). Quatre modes de codages hiérarchiques sont définis par MPEG2 (temporel, spatial, avec variation du pas de quantification, avec partitionnement des données prioritaires). Pour chacun d'eux il y a une couche de base et une ou plusieurs couches d'améliorations [Her02].

c. MPEG3

Cette norme n'a jamais vu le jour, elle était censée couvrir les très hauts débits. Mais MPEG2 se montrait étonnamment performant dans les très hauts débits et MPEG3 fut directement intégré dans MPEG2.

d. MPEG4

- **Description**

MPEG4 est plus qu'une simple amélioration des générations précédentes, il définit une véritable philosophie de la vidéo en dépassant les concepts de codage et de compression. Le MPEG4, en effet, définit à la fois une syntaxe du flux vidéo codé et un ensemble d'outils permettant la diffusion de cette vidéo sur des réseaux hétérogènes. L'aspect le plus visible du MPEG4 est l'introduction de la notion d'objets médias. La scène vidéo est décomposée en objets visuels indépendants, une personne, un meuble, une carte, ... et en objets audio comme la voix d'une personne, le bruit de moteur d'une voiture, ... De plus ces objets peuvent être naturels ou synthétiques : ils peuvent provenir d'une caméra ou d'un ordinateur. En fait un mixage est fait entre des sources réelles et des sources synthétiques. Ces sources " rapportées " peuvent être aussi bien en 2D, qu'en 3D. En plus de définir des objets, MPEG4 définit une arborescence structurant ces objets en les décomposant en d'autres objets. Une personne pourra être décomposée en sa voix, son visage, ses mains, et le reste du corps. La scène vidéo sera donc ainsi décomposée selon une hiérarchie de ces médias objets et selon leur disposition spatiale. MPEG4 permet également de les manipuler de différentes façons et ainsi de rendre interactif la scène vidéo transmise.

Les objets média pourront être identifiés et recevoir un code qui leur permettra de ne pas être réutilisé à outrance. On pourra ainsi protéger les droits d'un auteur sur son objet média.

- **Les applications visées**

Le standard MPEG4 va fournir un ensemble de technologies satisfaisant le besoin des auteurs, des fournisseurs et au final des utilisateurs.

Pour les auteurs: MPEG4 permettra la production de séquences réutilisables. Il leur permettra une grande flexibilité, autorisant l'amalgame de la télévision numérique, des animations graphiques, et des pages web. En outre, ils auront la possibilité de protéger leurs œuvres.

Pour les fournisseurs d'accès internet : MPEG4 offrira des informations transparentes qu'ils pourront aisément adapter à la demande de l'utilisateur (par exemple: l'adaptation en fonction de la langue de l'utilisateur), ainsi que le contrôle des transferts (gestion des pertes de données).

Pour les utilisateurs : MPEG4 aura de nombreuses possibilités qui pourront être accessible à partir d'un simple terminal. Voici un large éventail de toutes les applications concernées par les apports d'une telle standardisation :

1. La communication temps réel (vidéophone,...)
2. La surveillance
3. Le multimédia mobile (mini portable faisant office de téléphone, fax, agenda,... par liaison GSM ou satellite)
4. Le stockage et la recherche d'informations basés sur le contenu
5. La lecture de vidéo sur internet/intranet sans avoir à télécharger toute la source.
6. La visualisation de scène simultanément à plusieurs endroits (téléconférence,...)
7. La transmission (tout types de données : vidéo, audio,...)
8. La post production (cinéma et télé)
9. Le DVD
10. Les applications de l'animation de visage : réunions virtuelles,...
11. La hiérarchisation et la gestion des objets audio dans une scène.

- **Description Technique**

Elle concerne ici essentiellement l'aspect visuel de la norme.

Structure générale :

La norme MPEG4 propose une solution radicalement différente pour le codage des vidéos afin de satisfaire à tous ses besoins dans les différentes applications qu'elle propose. Les scènes audiovisuelles sont ainsi composées de plusieurs objets médias hiérarchisés.

Ainsi, dans l'arborescence de cette hiérarchie, on trouve:

- des images fixes (background)
- des objets vidéo (objets en mouvement sans background)
- des objets audio (la voix associée à l'objet en mouvement)

MPEG4 définit donc précisément la manière de décrire d'une scène. La description d'une scène codée par MPEG4 peut être comparée au langage VRML dans sa structure et ses fonctionnalités.

Le schéma suivant donne l'exemple d'une scène décomposé suivant l'idée de MPEG4:

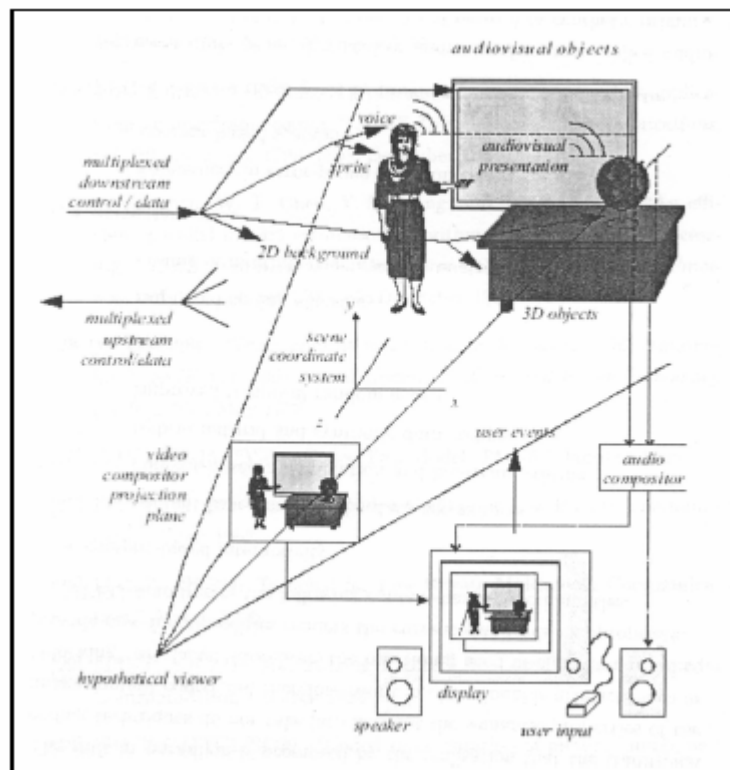


Fig 2.4. Exemple d'une scène MPEG4

- **Description d'une scène**

Une scène audiovisuelle, codée par MPEG4, est décrite comme un ensemble d'éléments individualisés. Elle contient des composants "média" simples regroupés par type. Ces groupes correspondent aux branches d'un arbre de découpage où chaque feuille représente un élément simple. Voici un exemple de branche de codage. Où on distingue les feuilles :

Textes, graphiques, mouvements de la bouche et son texte associé, plus l'animation de la tête son synthétique. Par exemple, si cette branche correspondait à une personne qui parle, elle serait divisée en feuilles contenant le fond, la parole et les divers composants graphiques représentant la personne en train de parler. Une telle construction permet ainsi la construction de scènes complexes tout en autorisant l'utilisateur à manipuler qu'une partie des objets. Un objet média peut donc être associé à une information, comme on associe ici la parole à la tête d'un personnage.

Voici le schéma de structure d'une scène MPEG4 (Fig2.4) :

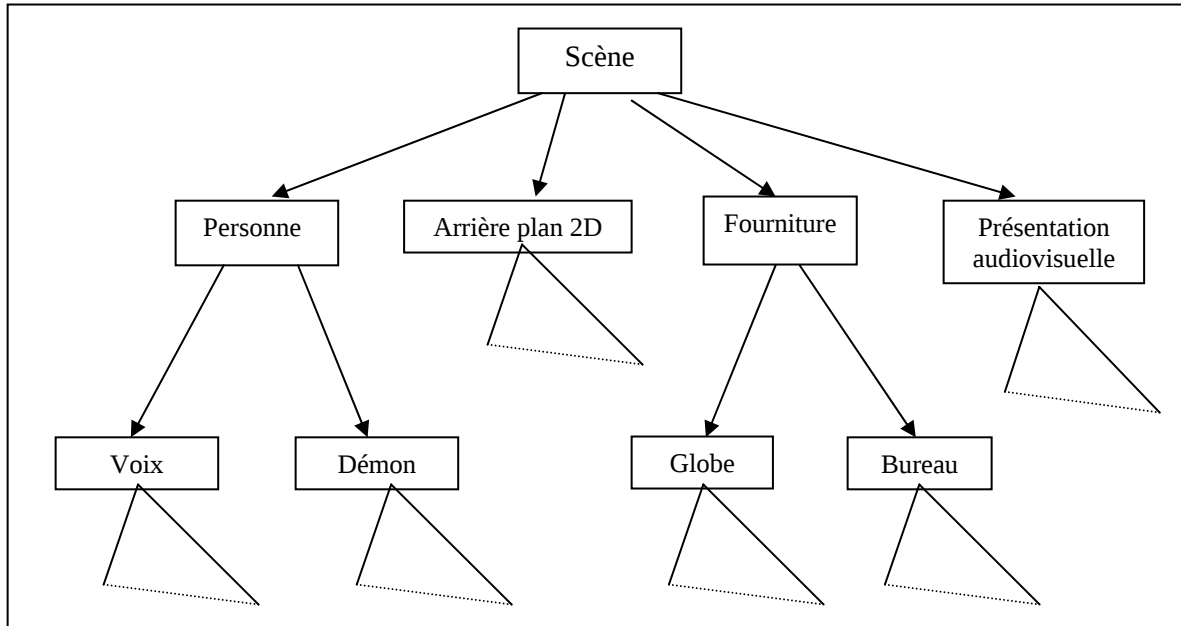


Fig 2.5. Structure d'une scène MPEG-4

MPEG4 fournit donc des méthodes de codage pour les objets individuels (comme nous venons de le voir). La norme permet également d'optimiser le codage de plusieurs objets dans une scène. L'information nécessaire à la composition d'une scène est contenue dans la description de la scène. Celle-ci est codée et transmise avec les objets média. Ainsi, pour faciliter l'interactivité, la description de la scène est codée indépendamment des primitives "Objets média". Une grande attention est portée sur l'identification des paramètres relatifs à la scène. Ces paramètres sont donnés par différents algorithmes qui codent de façon optimale les objets (cf " motion vectors in video coding algorithms " pour les objets fixes, et " position of the object in the scene " pour les mouvements d'objets). MPEG4 autorise la modification de ces paramètres sans avoir à décoder les objets média. Pour cela, ils sont placés dans la partie description de la scène et non avec les objets média [Hur03].

2.3.5 H264/ MPEG4 AVC

a- D'autres normes de codage visuelles: H.261, H.263 et H.264 :

Les normes H.261 et H.263 sont principalement employées pour le système de téléconférences. Bien que le MPEG soit employé pour principalement des films cinématographiques, ses nouveaux perfectionnements permettent son utilisation dans des applications de vidéoconférence, (Fig 2.7)

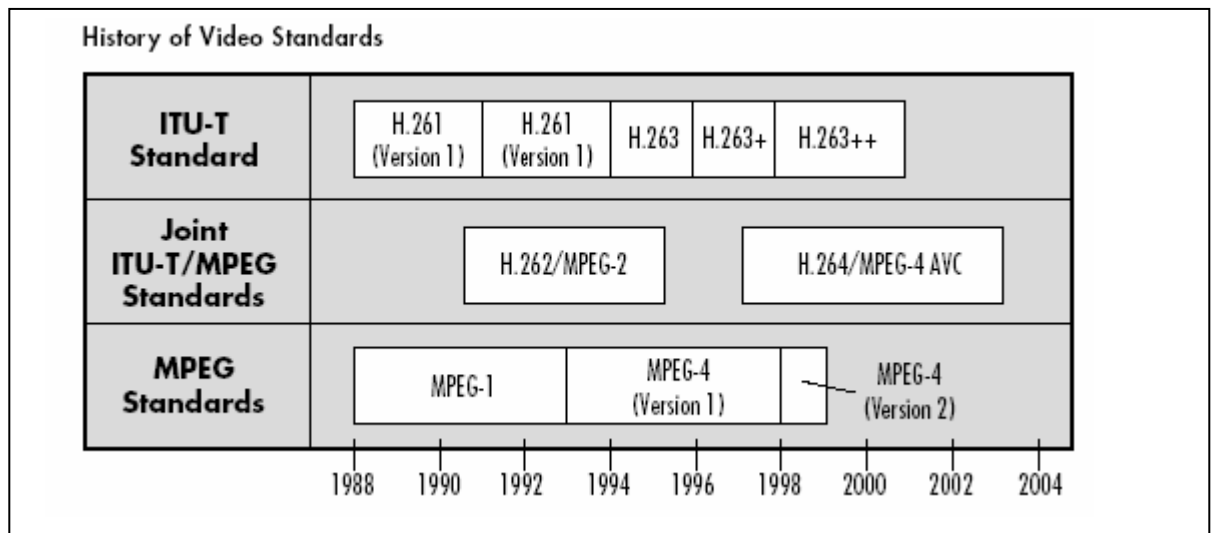


Fig 2.6. Evolution de la compression vidéo

b- H264/MPEG4 AVC

La nouvelle norme de compression vidéo H.264/mpeg-4 AVC apporte une amélioration significative au-dessus de toutes les normes visuelles précédentes de compression. En termes d'efficacité de codage, on s'attend à ce que la nouvelle norme fournisse au moins l'amélioration de la compression 2x au- meilleures que les normes précédentes et des améliorations de la qualité perceptuelles substantielles au-dessus de MPEG-2 et de MPEG-4. La norme, étant conjointement développé par ITU-T et ISO/IEC, adressera la gamme complète des applications visuelles comprenant de applications sans fil de basses débit binaire, vidéo streaming sur internet, livraison DVD, et la de la plus haute qualité vidéo pour des applications numérique de cinéma.

Le nom pour la norme est H.264 (précédemment appelé H.261) pour ITU/T, alors que le nom d'ISO/IEC est par MPEG-4 (AVC) qui deviendra la norme MPEG-4 PARTIE 10 Puisque AVC est une prolongation à la norme MPEG-4 courante, il tirera bénéfice des outils bien développés d'infrastructure de MPEG-4 . On s'attend à ce que MPEG-4 AVC soit choisi au-dessus de la norme visuelle courante de la compression MPEG-4, pour la majorité d'applications qui exigent les niveaux les plus élevés de compression et de qualité [Hur03].

Comme nous avons vu dans le diagramme " histoire des normes visuelles", ITU-T et ISO/IEC sont responsables de toutes les normes visuelles internationales précédentes de compression. Jusqu'ici, le plus réussi de ces normes a été MPEG-2, qui a continué pour réaliser l'acceptation de marché grand-public dans les secteurs tels que DVD, émission numérique de télévision (câble et satellite d'excédent), et boîte numérique de placer- dessus. La nouvelle norme de H.264/mpeg-4 AVC représente la plus grande amélioration simple de l'efficacité et de la qualité de codage depuis l'introduction de MPEG-2.

En conséquence, avec le temps, on s'attend à ce que H.264/mpeg-4 AVC déplace MPEG-2 et MPEG-4 dans beaucoup d'applications existantes.

2.4. Conclusion

MPEG4 cible la vidéo à très faible débit (**VLBV** : Very Low Bitrate Video). Pour cela, elle reprend les fonctionnalités de MPEG2 et élargit son champ d'action. MPEG4 peut ainsi se concevoir comme une boîte à outils permettant à la fois de décomposer une vidéo en différents types d'objets et de manipuler ces objets avec des outils adéquats. En plus de cela, MPEG4 est une norme portable qui lui permettra de s'affranchir des contraintes liées aux réseaux. Sa façon de définir des interfaces vers certains outils plutôt que ces outils eux-mêmes la rend très malléable et n'impose pas ainsi à un distributeur de services une seule façon d'implémenter un décodeur ou un codeur.

L'extension de la notion de codage hiérarchique introduite par MPEG2 est aussi une révolution en matière de flexibilité et de réduction de trafic dans les réseaux.

Le nouveau H.264/mpeg-4 AVC offre des avantages significatifs de débit binaire et de qualité par rapport à toutes les normes précédentes. La norme H.264/mpeg-4 AVC est ratifiée par ISO en 2003. Les produits initiaux basés sur la norme sont disponibles maintenant.

Chapitre 3

MOBILITÉ IP

3.1. Introduction

L'objectif de ce chapitre est de présenter l'une des solutions proposées par l'IETF pour résoudre le problème de la mobilité dans les environnements IP : Mobile IP **[Mon02]**. Dans ce chapitre, après une brève introduction aux problèmes liés à la mobilité dans les environnements IP, nous décrivons le protocole Mobile IP, qui est actuellement le plus largement utilisé dans l'Internet pour gérer la mobilité. Ce protocole est adapté aux problèmes de macro mobilité (c'est à dire entre différents domaines de mobilité). Nous présentons ensuite des solutions de micro mobilité qui permettent de traiter les déplacements à l'intérieur d'un même domaine.

3.2. Problème de la mobilité IP

IPv4 identifie le point d'accès d'un noeud sur l'Internet de manière unique grâce à son adresse IP; Celle-ci se décompose en deux parties :

- le préfixe qui détermine le sous-réseau sur lequel la machine se trouve,
- l'identifiant de la machine sur son sous-réseau.

L'Internet est un réseau de très grande taille pour que chaque routeur puisse mémoriser une route vers toutes les machines qui y sont attachées. En fait, les routeurs ne stockent que des entrées correspondant à des sous-réseaux, considérant que des datagrammes destinés à des machines ayant le même préfixe seront tous routés de manière identique.

La mobilité introduit un nouveau problème de routage : les mobiles se déplacent d'un sous réseau IP vers un autre sous-réseau IP, mais ont un mauvais préfixe sur le réseau destination.

Par conséquent, un noeud doit être situé sur le réseau indiqué par son adresse IP afin de pouvoir recevoir les paquets qui lui sont destinés. Pour qu'un noeud puisse changer de point d'accès sans perdre la possibilité de communiquer, deux mécanismes peuvent être employés :

- le noeud doit changer d'adresse IP à chaque fois qu'il change de point d'accès
- des chemins spécifiques à l'hôte doivent être propagés dans presque toute la structure de routage de l'Internet.

Ces deux alternatives sont souvent inacceptables. La première ne permet pas à un noeud de conserver des connexions au niveau de la couche transport ou des couches supérieures lorsqu'il change de position. La seconde pose des problèmes de passage à l'échelle.

Un nouveau mécanisme flexible est nécessaire afin de s'adapter à la mobilité des noeuds sur Internet. Le protocole Mobile IP permet aux noeuds de changer de point d'accès à l'Internet sans changer d'adresse IP. Les spécifications minimales de la solution recherchée sont les suivantes :

- le déplacement d'un mobile ne doit pas provoquer de coupure des connexions ouvertes ;
- l'opération doit être simple à mettre en oeuvre et d'un coût raisonnable;
- l'accès aux ressources doit être transparent ;
- la solution doit être compatible avec le protocole IP et en particulier avec les algorithmes de routage. Le support de la mobilité ne doit pas nécessiter la modification de tous les routeurs.

3.3. Mobile IP

L'IETF (Internet Engineering Task Force) est l'organisme chargé de développer et de publier les protocoles reconnus comme les protocoles standard sur l'Internet. Cet organisme est composé de groupes de travail qui s'occupent chacun d'un domaine bien précis. L'un d'entre eux (Mobile IP) est chargé de proposer un protocole pour le support des mobiles sur l'Internet.

3.3.1. Définitions

Noeud mobile :

Hôte ou routeur qui change de point d'accès d'un réseau (ou sous-réseau) à un autre. Comme toute machine fixe, un mobile appartient initialement à un réseau sur Internet. Ce réseau, appelé réseau mère (Home Network), affecte son adresse IP au mobile. Un noeud mobile peut changer de position sans changer d'adresse IP.

Home Agent (HA) :

Routeur sur le réseau mère d'un mobile, qui envoie les datagrammes dans un tunnel pour les remettre au mobile lorsqu'il visite un autre réseau. Le Home Agent met à jour les informations concernant la position du mobile.

Foreign Agent (FA) :

Routeur sur un réseau visité par le noeud mobile, qui fournit des services de routage au mobile lorsqu'il est enregistré auprès de lui.

Adresse permanente / temporaire (COA) :

Un noeud mobile possède une adresse IP permanente sur son réseau mère. Lorsqu'il visite un autre réseau, une adresse temporaire (care-of address) est affectée au mobile. Cette adresse reflète le point d'accès du mobile. En général, le mobile utilise son adresse mère comme adresse source dans tous les datagrammes IP qu'il envoie.

Correspondant (CN) :

Machine (mobile ou non) qui dialogue avec un mobile.

Handoff :

Changement de point d'attachement d'un mobile à l'Internet. C'est le concept central de la mobilité, qui consiste à établir un lien au niveau de chaque nouveau point d'attachement au réseau.

3.3.2. Fonctionnement du protocole

3.3.2.1. Principe de base

Un noeud mobile, noté MN, se déplace de réseau en réseau et s'attache à des bases. Une base implémente des fonctionnalités de niveau 2 du modèle OSI. Elle assure une connectivité de niveau liaison de données et permet l'échange d'informations avec un mobile par un canal sans fil. Quand un correspondant veut envoyer des données à un mobile, il crée un paquet IP ordinaire. L'adresse IP source est celle du correspondant, l'adresse IP destination, celle du mobile sur son réseau mère. Les paquets arrivent donc sur le réseau mère du mobile où ils sont interceptés par le Home Agent. Celui-ci transmet les paquets vers la position courante du mobile.

3.3.2.2. Adresse temporaire

Lorsqu'un mobile quitte son réseau mère, Mobile IP utilise le "tunnelage" pour cacher l'adresse mère du mobile aux routeurs situés entre le réseau mère et le mobile. La fin du tunnel correspond à l'adresse temporaire du mobile. Cette adresse temporaire doit être une adresse à laquelle les datagrammes peuvent être remis par des mécanismes de routage IP classiques. A l'adresse temporaire, le datagramme d'origine est enlevé du tunnel et remis au noeud mobile.

L'adresse temporaire peut être obtenue de deux manières différentes :

1. Une "Foreign Agent care-of address" est une adresse temporaire fournie par un Foreign Agent grâce aux messages Agent Advertisement. Dans ce cas, l'adresse temporaire est l'adresse IP du Foreign Agent. C'est donc le Foreign Agent qui est à l'extrémité du tunnel ; lorsqu'il reçoit les datagrammes tunnelés, il les décapsule et remet le datagramme d'origine au mobile. Ce mode d'acquisition d'adresse temporaire est préférable car il permet à de nombreux noeuds mobiles de partager une même adresse temporaire.

2. Une "co-located care-of address" est une adresse IP locale, acquise par des moyens externes, que le mobile associe à l'une des ses interfaces de réseau. Cette adresse peut être acquise dynamiquement par des mécanismes tels que DHCP, ou peut être possédée par le mobile comme adresse à utiliser dans certains réseaux visités. Lorsque le noeud mobile utilise une co-located care-of address, il se trouve lui-même à l'extrémité du tunnel et décapsule les datagrammes tunnelés jusqu'à lui. Ce mode permet à un mobile de fonctionner sans Foreign Agent, mais il pose un problème au niveau de l'espace d'adressage d'IPv4.

3.3.2.3. Encapsulation IP dans IP

Le Home Agent ne fait qu'encapsuler les paquets, c'est-à-dire qu'il ne les modifie pas de telle façon qu'il n'a pas besoin de recalculer les sommes de contrôle (*checksum*) des couches supérieures. Il se contente donc d'ajouter un nouvel en-tête IP devant l'en-tête du datagramme reçu [Per96], comme le montre la Figure 3-1. Ainsi les datagrammes sont facilement redirigés et routés dans l'Internet. L'adresse IP vers laquelle le Home Agent fait suivre les paquets destinés au mobile est la care-of address.

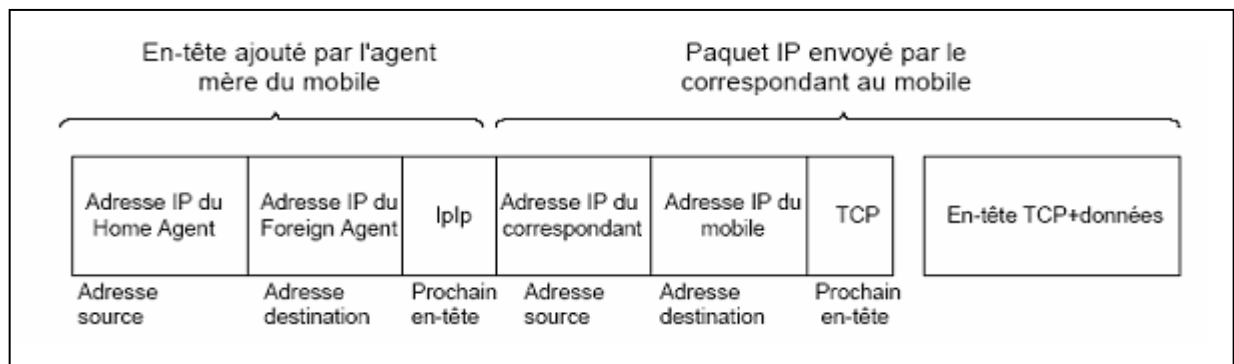


Fig 3.1 Encapsulation IP dans IP.

Mobile IP définit un ensemble de messages de contrôle, envoyés avec UDP en utilisant le numéro de port 434. Nous nous intéressons à deux types de messages : demande d'enregistrement (Registration Request) et réponse d'enregistrement (Registration Reply).

Pour la découverte des agents, Mobile IP utilise les messages existants Router Advertisement et Router Solicitation définis pour la découverte de routeurs ICMP [Dee91].

3.3.2.4. Découverte des Foreign Agents et enregistrement

Des mécanismes de communication entre les mobiles et leurs Foreign Agents sont nécessaires. En particulier, il faut qu'un mobile puisse apprendre l'adresse IP des agents relais. Le mobile doit ensuite s'enregistrer auprès d'un Foreign Agent et obtenir son accord pour qu'il relaye les paquets. Puis, le Home Agent doit être informé de l'adresse IP du nouveau Foreign Agent, qui devient l'adresse

temporaire du mobile (COA). On peut également transmettre cette adresse à l'ancien Foreign Agent (sur le réseau que le mobile a quitté) pour que les paquets en transit au moment du déplacement du mobile arrivent bien à destination. En effet, il se peut que le Home Agent continue à envoyer des paquets lors de la procédure de changement de Foreign Agent.

La gestion des agents relais peut par exemple être effectuée en utilisant le protocole ICMP. Les agents relais signalent leur présence avec des ICMP Agent Advertisements. La réception par un mobile de l'un de ces messages lui permet de savoir s'il est ou non sur son réseau mère et de démarrer la procédure d'enregistrement. Cette procédure se décompose en quatre messages **[Per01]** :

- demande d'enregistrement au Foreign Agent (Registration Request),
- traitement de la demande par le Foreign Agent, puis transmission de la demande au Home Agent,
- notification d'acceptation ou de refus de prise en charge du mobile du Home Agent au Foreign Agent (Registration Reply).
- traitement de la notification par le Foreign Agent et transmission de l'information au mobile.

La procédure d'enregistrement est accompagnée d'une authentification afin d'assurer que les paquets destinés au mobile ne sont pas détournés par une autre machine. Lorsque le mobile est situé sur son réseau mère, il n'utilise pas les services de mobilité. L'une des extensions des messages ICMP Router Discovery, appelée Mobility Agent Advertisement Extension, permet d'utiliser un champ durée de vie (TTL), présent dans les messages ICMP Router Advertisement envoyés régulièrement par un Foreign Agent. Quand un mobile se connecte sur un réseau, il enregistre la durée de vie associée à son Foreign Agent. Lorsque cette valeur expire, le mobile suppose qu'il n'est plus géré par ce Foreign Agent. Par contre, si le mobile reçoit un Router Advertisement d'un autre Foreign Agent, il s'enregistre auprès de celui-ci.

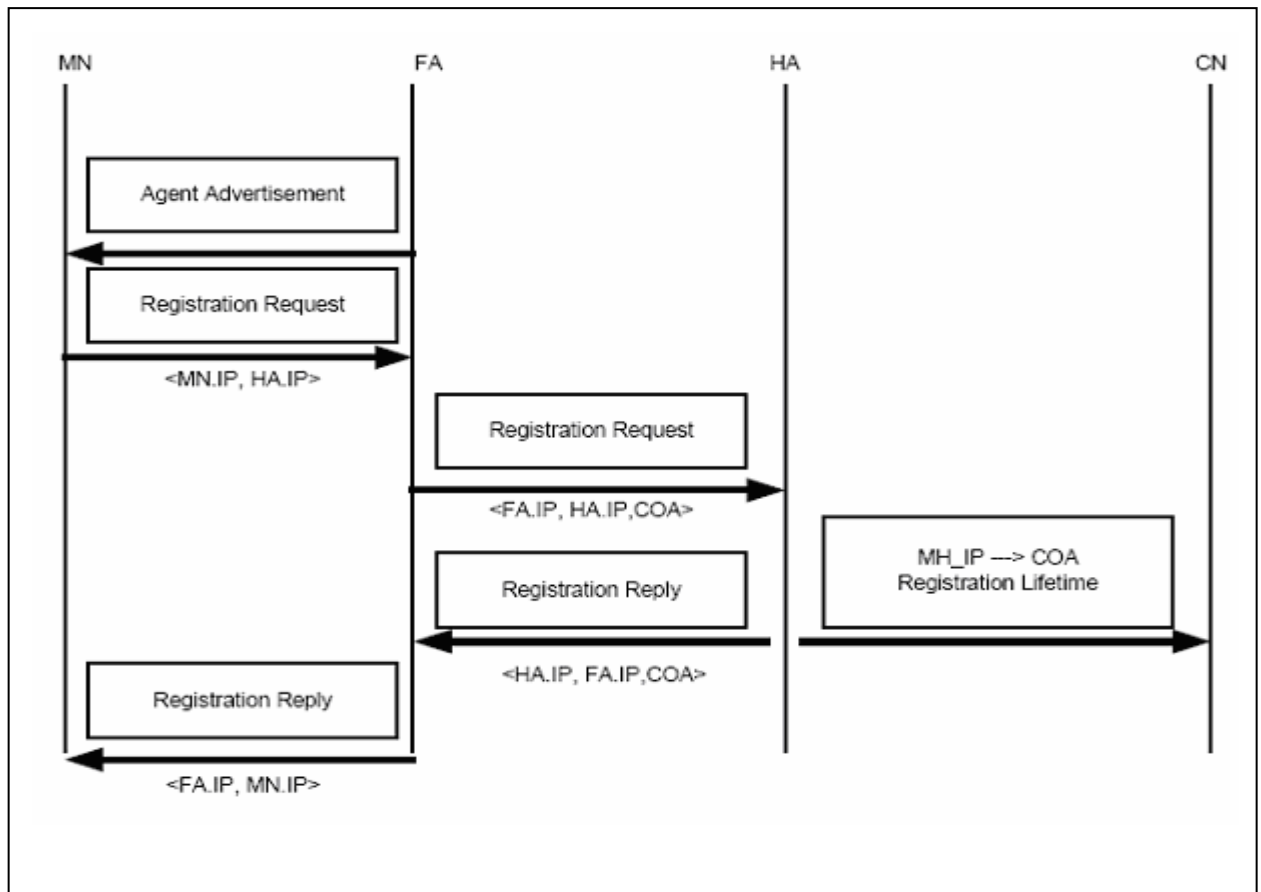


Fig 3.2 Enregistrement.

3.3.2.5. Extensions

Mobile IP définit un mécanisme d'extension général pour permettre aux messages de contrôle Mobile IP et aux messages de découverte de routeurs ICMP de transporter des informations supplémentaires. Ces extensions sont codées avec le format Type-Longueur-Valeur, à une exception près.

Certaines extensions n'apparaissent que dans les messages des contrôles Mobile IP :

- Authentification mobile / Home Agent.
- Authentification mobile / Foreign Agent.
- Authentification Foreign Agent / Home Agent.

Les extensions présentes uniquement dans les messages de recherche de routeur ICMP sont :

- Bourrage d'un octet
- Annonce d'agent de mobilité (Mobility agent advertisement)
- Longueurs des préfixes

3.3.2.6. Routage

Lorsqu'un correspondant veut dialoguer avec un mobile (sens descendant), il ne connaît pas la position courante de ce dernier et envoie les datagrammes à l'adresse mère du mobile. Ceux-ci sont

interceptés par le Home Agent et envoyés dans un tunnel du Home Agent vers l'adresse temporaire du mobile. A la sortie du tunnel (au niveau soit du Foreign Agent, soit du mobile lui-même), les datagrammes sont transmis au mobile. Si le correspondant est mobile, il doit également passer par son propre Foreign Agent.

Dans l'autre sens (sens montant), les datagrammes envoyés par le noeud mobile sont généralement remis aux destinataires en utilisant des mécanismes de routage IP standard, en passant par le Home Agent (mode bidirectionnel) ou non (mode standard).

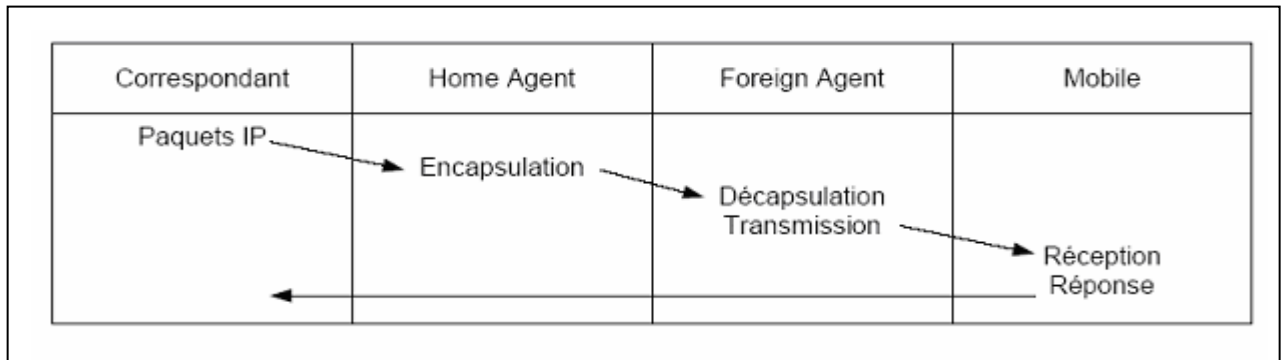


Fig 3.3 Dialogue avec un Mobile

Ce mode est très peu utilisé. Il peut servir lorsqu'un mobile reçoit des paquets multicast, si son Home Agent est un routeur multicast. Le mobile encapsule les messages IGMP vers son Home Agent et ce dernier qui envoie les messages multicast pour le mobile dans le tunnel. Pour des paquets tunnelés au HA, l'adresse source dans l'entête IP doit être l'adresse mère du mobile.

La station émettrice (E) envoie un datagramme destiné au mobile. E utilise l'adresse mère du mobile sans se préoccuper de la mobilité.

Cas 1 :

Le mobile est actuellement dans son réseau mère, son HA ne procède à aucune action particulière, le paquet est acheminé normalement jusqu'à lui.

Cas 2 :

Le mobile s'est déplacé et ne dépend plus de son réseau mère. Le HA intercepte le paquet et l'encapsule dans un nouveau datagramme IP en mettant la dernière adresse temporaire connue du mobile. Il s'agit ici de l'adresse du FA et non d'une adresse attribuée dynamiquement. Le paquet arrive au FA qui le décapsule et envoie le datagramme initial sur son réseau.

Exemple :

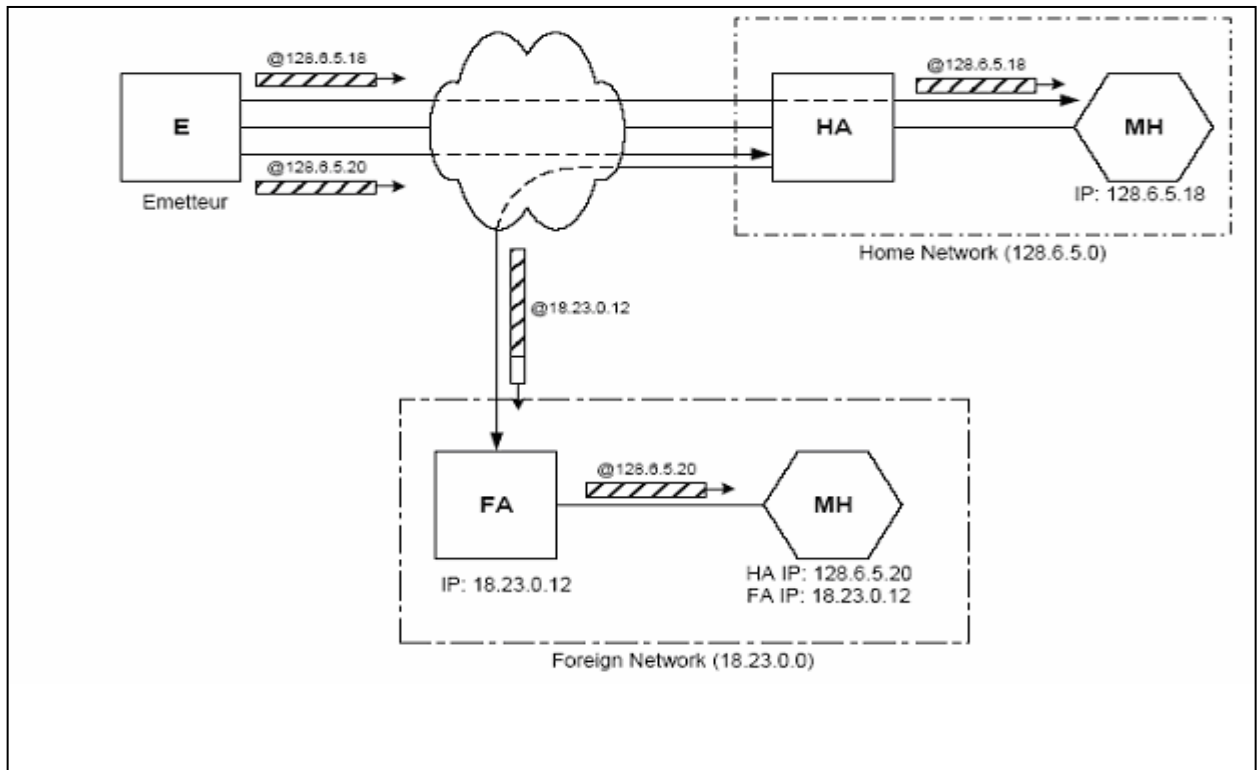


Fig 3.4 Exemple de routage

Routage triangulaire

Du point de vue du routage, la principale faiblesse de Mobile IP est le routage triangulaire, c'est à dire un routage via le Home Agent dans le sens correspondant vers mobile.

En effet, lorsqu'un mobile hors de son sous-réseau mère cherche à joindre une machine sur le même réseau visité, les paquets doivent néanmoins transiter par le sous-réseau mère du mobile. Ce problème n'apparaît pas dans le sens mobile vers correspondant car Mobile IP définit un routage direct.

Routage optimisé

Afin d'obtenir un routage performant dans le sens correspondant vers mobile, [Per01] propose d'ajouter des mécanismes supplémentaires sur les correspondants d'un mobile et de modifier la pile de protocoles TCP/IP. Mais cette solution supprime la transparence du protocole par rapport aux utilisateurs, c'est à dire qu'un correspondant doit à présent savoir s'il dialogue avec une machine fixe ou mobile.

Dans cette proposition, seul le premier paquet envoyé à un mobile passe par son Home Agent; après réception de ce paquet, le Home Agent indique au correspondant que la machine qu'il cherche à joindre est un mobile et fournit l'adresse par laquelle il peut être joint (son adresse temporaire).

Un mobile peut avoir plusieurs correspondants susceptibles d'utiliser des routes non optimales (c'est à dire qui passent par le Home Agent). Le Home Agent mémorise donc la liste des correspondants de chacun des mobiles qu'il gère pour informer les correspondants des mouvements des mobiles.

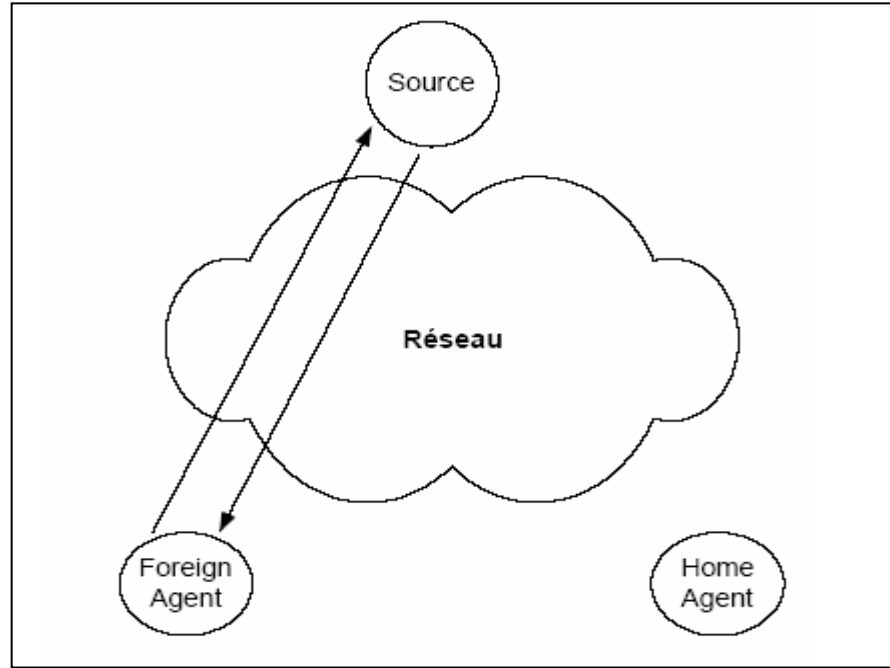


Fig 3.5 *Routing Optimisé*

3.4. MIPv6

Une nouvelle version du protocole IP est en train d'émerger depuis quelques années : il s'agit de la version 6 du protocole IP. Ce protocole inclut entre autres la mobilité en standard. L'objectif de MIPv6 est d'offrir une communication directe entre un noeud mobile et ses correspondants (élimination du routage triangulaire) et éviter les ruptures des communications pendant les déplacements. Bien que MIPv6 reprenne des mécanismes de MIPv4, de nombreuses fonctionnalités supplémentaires ont été mises en place.

3.4.1. Fonctionnalités requises

Dans MIPv6, l'agent visité décrit dans MIPv4 n'existe plus. Par contre, l'agent mère est encore un routeur d'accès du sous-réseau principal du noeud mobile. Son rôle est le même que dans le cas de MIPv4, à savoir capturer les paquets à destination du mobile et les lui tunneler à sa localisation courante.

Par contre, les correspondants doivent mettre en oeuvre certains mécanismes supplémentaires : tout d'abord, ils doivent disposer d'un cache d'association tout comme l'agent mère ; dans ce cache sera

stockée la correspondance entre l'adresse principale d'un noeud mobile avec lequel il a une communication et son adresse temporaire courante. Il devra donc être capable de traiter des messages de registration envoyés par un noeud mobile. De plus, il devra être capable d'effectuer le routage directement vers le mobile (**routing header**). Ceci constitue un apport important dans le fonctionnement de la mobilité puisque les paquets des correspondants n'auront pas à passer par le réseau mère systématiquement. Mais toutes ces fonctionnalités supplémentaires ne sont faites qu'au niveau de la couche IP ; l'adresse identifiant la communication au niveau applicatif sera toujours l'adresse principale du noeud mobile, la couche IP cahant l'adresse temporaire source (ou destination selon qu'on se situe sur le noeud mobile ou le correspondant).

D'un autre côté, un noeud mobile doit toujours conserver la liste des correspondants auxquels il envoie un message de registration (pour les mises à jour éventuelles) et doit être capable de décapsuler lui-même les paquets qui lui sont transmis ; au niveau application, un noeud mobile utilise toujours son adresse principale, c'est pourquoi la couche IP doit pouvoir décapsuler l'en-tête indiquant l'adresse temporaire. Cette opération était exécutée par le agent visité dans MIPv4.

3.4.2. Découverte des routeurs d'accès

Le protocole de **découverte des voisins** [Std99] offert par IPv6 joue un rôle important dans MIPv6. Il permet entre autres à des équipements situés sur le même lien physique de se découvrir mutuellement, de découvrir leurs adresses niveau 2 et de localiser les équipements de routage. Le processus de découverte des routeurs d'accès se déroule de manière similaire au protocole de découverte des agents ; tout routeur d'accès émet périodiquement des *Router Advertisement* contenant la liste des préfixes sur le lien. Un noeud mobile peut éventuellement en demander un explicitement, à travers un *Router Solicitation*.

3.4.3. Enregistrement

De la même manière que dans MIPv4, lorsqu'un noeud mobile se déplace hors de son sous-réseau mère, il doit en informer son agent mère. Le noeud mobile signale la correspondance entre son adresse principale et son adresse temporaire courante dans un message *Binding Update*. Ce message peut éventuellement être envoyé en « piggybacking8 ». En réponse à une telle requête, l'agent mère envoie un *Binding Acknowledgement* pour indiquer s'il peut répondre à la requête du mobile. Pour le moment, tout se passe comme dans MIPv4. Cependant, le mobile a par la suite la possibilité d'informer ses correspondants de sa position courante ; Lorsqu'il reçoit un paquet d'un correspondant, il détermine si le paquet a transigé par le réseau mère en regardant si le paquet contient un routing header ou s'il a été tunnelé par l'agent mère (encapsulation). S'il ne contient pas de routing header, le noeud

mobile en déduit que le correspondant émetteur n'a pas d'entrée dans son cache d'association pour lui. Il peut alors lui envoyer un *Binding Update* pour qu'il lui envoie les paquets directement, sans plus passer par son sous-réseau mère.

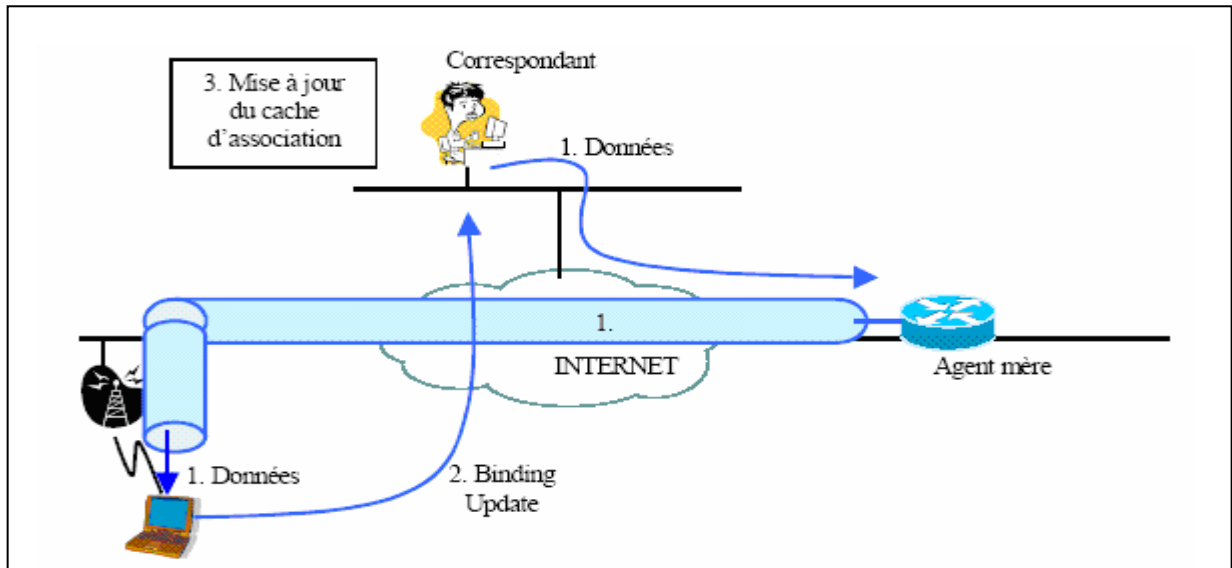


Fig 3.6. Communication IPv6

Le fait que les correspondants aient la possibilité d'envoyer les paquets directement au mobile offre une meilleure résistance au facteur d'échelle et fiabilité. La communication entre noeuds mobiles et correspondants engendre moins de charge sur le réseau et est plus rapide. Comme l'agent mère est peu sollicité pour la retransmission des paquets, il y a beaucoup moins de risque de congestion au niveau de l'agent mère et une panne de l'agent mère aura un effet moindre.

Un noeud mobile peut détenir plus d'une adresse temporaire à un instant donné. Celle enregistrée auprès de l'agent mère (une seule) est dite principale. L'utilisation de plusieurs adresses temporaires peut être utile pour améliorer les performances lors d'un déplacement lorsque par exemple deux cellules de points d'accès se recouvrent fortement ; le noeud mobile peut alors acquérir une nouvelle adresse temporaire tout en utilisant son ancienne le temps de l'opération.

3.5. Limites de Mobile IP

Mobile IP est conçu pour permettre aux noeuds de se déplacer d'un sous-réseau IP à un autre ; il résout le problème de la "macro" mobilité. Lorsque les mobiles effectuent des handoffs fréquents dans une même zone géographique. Tant que le déplacement du noeud ne se situe pas entre des points d'accès sur des sous-réseaux IP différents, des mécanismes de mobilité de niveau liaison offrent probablement une convergence plus rapide et nécessitent bien moins d'overhead que Mobile IP. Dans la version standard de Mobile IP, la nouvelle localisation d'un mobile est toujours signalée à son Home Agent. Ce dernier est ainsi averti de tous les déplacements des mobiles qu'il gère. Ces opérations

gènèrent une grande quantité de signalisation. De plus, les pertes de paquets pendant les handoffs peuvent être importantes puisque la procédure d'enregistrement est longue, en particulier si le Home Agent se trouve à l'autre bout du monde. La durée d'un handoff peut atteindre plusieurs secondes dans l'Internet actuel. La micro mobilité représente donc la gestion de la mobilité locale. La gestion de la micro mobilité permet à des mobiles de s'enregistrer localement à l'intérieur du réseau qu'ils visitent. Cet enregistrement local réduit le délai de signalisation, ce qui peut améliorer les performances des handoffs.

D'autre part, pour des mobiles bénéficiant d'une certaine QoS, l'acquisition d'une nouvelle care-of address pour chaque handoff implique de remettre en place des réservations entre le Home Agent et le nouveau Foreign Agent. Cependant, une grande partie du chemin entre le correspondant et le mobile risque d'être identique avant et après le handoff (si le mobile ne change pas de domaine).

Plusieurs solutions pour gérer la micro mobilité ont été proposées. Toutes reposent sur le principe suivant : le réseau visité par un mobile se charge des déplacements locaux ; le Home Agent n'est donc pas prévenu de tous les changements de localisation du mobile qu'il gère.

Ces propositions peuvent être évaluées selon différents critères :

- o la quantité de signalisation nécessaire par rapport au protocole Mobile IP standard,
- o le nombre de paquets perdus pendant un handoff,
- o le délai de signalisation : durée d'un handoff, temps nécessaire pour qu'un mobile reçoive des paquets à sa nouvelle localisation,
- o les changements d'implémentation des noeuds mobiles,
- o les modifications à apporter au Home Agent, aux correspondants et aux Foreign Agents,
- o les modifications à apporter aux autres noeuds dans le réseau visité.

Certains des paramètres ci-dessus peuvent dépendre d'autres facteurs. Par exemple, la technologie de transmission a un effet sur la durée d'un handoff.

Les routeurs d'accès offrant des fonctionnalités pour la mobilité émettent des *Router Advertisement* quelque peu modifié (pour avertir les mobiles de leur capacité). En outre, l'information contenue dans ces *Router Advertisement* permet aux noeuds mobiles de créer une adresse temporaire (auto configuration offerte par IPv6). Ensuite il leur faudra vérifier l'unicité de celle-ci grâce au protocole de détection de duplication d'adresse [Ind98]. La découverte des voisins ainsi que la découverte de l'adresse de niveau 2 d'un équipement voisin s'avère aussi très utile dans la mobilité, notamment pour effectuer des registrations plus rapides. L'utilisation des ces données sera détaillée plus tard dans le rapport car le protocole MIPv6 ne prend pas encore en compte ces données.

3.6. Conclusion

Nous avons présenté dans ce chapitre le protocole Mobile IP qui supporte la mobilité des utilisateurs du Wlan. Mobile IP est conçu pour permettre aux nœuds de se déplacer d'un sous réseau IP à un autre, il résout le problème de la macro mobilité. Il est important de permettre la coexistence de mobile IP et tous ses extensions possibles (chapitre 6) pour surmonter les limites de MIPv6. Si nous supposons que cette coexistence est possible, chaque administrateur de réseau peut alors choisir quelle solution il souhaite utiliser à l'intérieure de son réseau. Ce choix ne doit pas affecter les mobiles, les Home Agents. Pour cela nous souhaitons offrir des garanties de qualité de service des transmissions vidéo dans les environnements mobiles. Les chapitres 6, 7,8 sont consacrés à l'étude et la proposition des mécanismes et extensions de Mipv6

Chapitre 4

Transmission du trafic vidéo MPEG sur un réseau WLAN 802.11b

4.1. Introduction

Au cours de ces dernières années, les avancées technologiques dans les réseaux sans fil 802.11b et le développement du domaine multimédia et en particulier, les applications multimédia telles que vidéo streaming, vidéo conferencing entraînent une augmentation de la demande en capacité de transmission. Une des principales difficultés réside dans le choix du support physique et de la technique de transmission. Les performances de telles applications rencontrent de nouveaux problèmes dus à la nature des liens sans fil et certains schémas d'accès.

L'étude de la transmission vidéo sur WLAN 802.11b devient différente concernant la manière à résoudre les problèmes de conception et le déploiement de l'infrastructure sans fil et à accomplir les besoins pour la transmission vidéo.

Le standard IEEE 802.11b contient des mécanismes pour la transmission des données en temps réel, ces mécanismes sont appelés « Point coordination : PCF, DCF » (voir chapitre1). Beaucoup de travaux sont réalisés dans l'étude des protocoles d'accès au médium 802.11b et les améliorations possibles pour une transmission temps réel [Bou03] [Tay01][Cou00] .

Dans cette partie nous allons présenter les différents mécanismes pour une transmission sans fil en temps réel, Ainsi que les problèmes de cette transmission dus aux liens sans fil. Après nous définissons un modèle de simulation de la transmission vidéo *MPEG* sur un réseau WLAN 802.11b sur la base des travaux déjà réalisés [Mas03] [Apo02] [Lia01] [Wan02] [Bou03]. Nous nous limitons à des scénarios bien définis, toutes les simulations seront basées sur le schéma d'accès *PCF*, la nouvelle norme de compression *MPEG4 AVC10*, et les débits 11Mb/s, 5,5Mb/s. Ensuite nous prendrons à la fin les mesures en terme de retard et de dégradation *PNSR*

4.2. La diffusion de la vidéo sur un Wlan

4.2.1. La préparation des flux MPEG

La transmission de données sur un réseau se fait par paquets et non par flux continu. Il faut donc diviser le flux de vidéo en paquets pour permettre leur envoi par le biais d'un réseau IP [Try96].

Pour ce faire, on utilise les Transport Stream (TS) qui sont utilisés pour la transmission multiple dans un environnement sujet aux erreurs. Les flux vidéo sont divisés en paquets, les PES

(packetized elementary streams - flux élémentaires découpés). Ces PES sont dans le cadre du TS, de taille constante, et ne sont donc pas corrélés avec les différentes images du flux car une image P ne fait pas la même taille qu'une image I par exemple.

Un paquet de transport est composé de 188 octets, et contient 4 octets d'en-têtes et les données du flux. Comme un paquet au sens du réseau ne peut contenir les données que d'un seul PES à la fois[Hur03].

Les en-têtes du paquet de transport sont composées de :

- Un code de début de paquet, pour faciliter la synchronisation et l'identification d'un nouveau paquet.
- Un code indiquant le début d'un nouveau PES, ou la transmission de données hors flux.
- Un numéro d'identification du flux (utilisé dans le cas de plusieurs flux multiplexés).
- Un indicateur de priorité de transmission.
- Un compteur de continuité, incrémenté à chaque fois qu'un nouveau paquet de transport est généré pour un flux donné, utilisé pour détecter la perte de paquets.

4.2.2. RTP et RTCP

La diffusion de vidéo doit se faire autant que possible en temps réel; c'est-à-dire que la diffusion doit se faire à des vitesses comparables aux vitesses d'exécution de la tâche. Cela ne pose pas forcément de problème dans le cadre d'un véhicule d'exploration lent, mais peut en poser plus dans le cadre de la télé chirurgie[Hur03]. Or, il est difficile de garantir un trafic donné pour la diffusion vidéo. Bien évidemment, on dispose en général pour les applications critiques de réseaux dédiés et sécurisés. Mais la surveillance de la transmission est d'autant plus nécessaire que le canal est peu fiable.

Ainsi, les transmissions peuvent être sujettes à des retards dûs au routage et à des pertes qui sont le plus souvent des retards trop importants. La bande passante allouable à la diffusion vidéo peut aussi être modifiée au cours du temps. Ces paramètres modifient la qualité de la vidéo retransmise. C'est pourquoi il est nécessaire d'encapsuler la vidéo dans une entité plus large contenant aussi ces informations.

RTP (real-time transport protocol) et RTCP (RTP control protocol) sont définis dans [Hur03]. RTP permet d'encapsuler les données vidéo ou audio dans un format valide pour la transmission réseau, mais n'offre pas de garanties quant à la qualité de la transmission réseau. RTCP permet de contrôler la bonne marche de la transmission par RTP. RTP divise les données à envoyer (dans le cas qui nous intéresse, de la vidéo MPEG) en paquets. Ces paquets contiennent les en-têtes RTP et la vidéo MPEG comme traité en (partie II). Les en-têtes RTP contiennent principalement des informations de synchronisation et de datage. Des en-têtes supplémentaires peuvent également être définis suivant les besoins de l'application considérée. Mais RTP ne garantit pas la délivrance des paquets, ni l'ordre d'arrivée des paquets, ni la qualité du réseau sous-jacent.

C'est précisément à cela que sert RTCP, qui ne transporte pas de données au sens qui nous intéresse (vidéo) mais qui propose un certain nombre de mécanismes :

- Retour sur la qualité des données reçues – ce qui peut être intéressant dans le cadre d'une compression adaptative.
- Gestion des identifiants des clients : chaque client reçoit un identifiant unique, qui peut également servir à associer des flux multiples (par exemple assurer la synchronisation vidéo/audio).
- Contrôle du nombre de paquets RTCP que les clients doivent envoyer. Ce nombre peut être aussi élevé si le nombre de clients est faible pour assurer un meilleur suivi de l'état des connexions, mais doit être plus faible si le nombre de clients est élevé pour ne pas surcharger le réseau.
- Stockage d'informations de session minimales, qui sans pouvoir se substituer à un réel contrôle des sessions au niveau applicatif, permet un contrôle de base relativement simple.

Comme les en-têtes RTP, les paquets RTCP sont constitués d'un certain nombre de champs fixes et de champs supplémentaires utilisables selon l'application.

La combinaison des protocoles RTP et RTCP permet donc la diffusion vidéo en boucle fermée au sens de l'automatique, avec un retour possible par RTCP de la qualimétrie réseau finale, et, par l'intermédiaire des en-têtes supplémentaires, de la qualimétrie vidéo finale. Il est tout de même à noter que ce retour induit forcément un retard non négligeable (à cause des temps de transmissions réseau) par rapport aux contraintes temps réel de la compression vidéo.

4.3. Caractéristique du trafic vidéo encodé

Après avoir étudié le codage (compression) de la vidéo dans le chapitre 2. Nous nous intéressons maintenant à la vidéo encodée. En particulier nous étudions les caractéristiques du trafic vidéo aussi bien que la qualité de la vidéo. Le trafic vidéo est gouverné par les tailles des fragments vidéo. Le trafic vidéo codé est souvent enregistré dans des *traces vidéo* qui décrivent les caractéristiques du flux vidéo, tel que le nombre du fragment (index), le type (par exemple, I, P, ou B), les *playout time* (quand le fragment est décodé et affiché sur l'écran), et la taille des fragments, comme illustré par l'extrait suivant d'une trace de la vidéo après l'encodage [Fit03].

<i>Num-fragment</i>	<i>Type-fragment</i>	<i>Temps [ms]</i>	<i>Taille [octet]</i>
0	I	0.0	700
1	P	12.0	1000
...	...	...	...

Les traces vidéo sont la base de l'étude d'un trafic vidéo [Fit03] [Fen97]. Les traces sont généralement utilisées pour les simulations du trafic vidéo et ses mécanismes. (Par exemple NS [NS],

Omnet++ [Var02]). Pour cela les traces vidéo seront par la suite utilisées dans la définition du modèle de simulation d'une transmission vidéo sur un réseau Wlan(section 4.7).

4.4. Mesure de la qualité de la vidéo

4.4.1. Introduction

La vidéo doit être maintenue à une qualité minimale que l'opérateur soit en mesure de l'exploiter. S'il n'est par exemple pas possible de faire la différence entre le sol et un obstacle, il sera difficile de piloter un robot mobile dans l'environnement considéré. Comme la vidéo subit une dégradation au cours de la compression nécessaire à sa transmission, la mesure de la qualité finale est essentielle au bon déroulement des opérations.

Mais ce problème est complexe : comment quantifier la qualité d'une vidéo ? Un être humain pourra parler d'une « bonne qualité », d'une « mauvaise qualité », d'une « bonne qualité dans l'ensemble »... L'association d'une échelle numérique à ces critères peut être une idée valable, mais reste fortement dépendante de la sensibilité du spectateur. Cela reste tout de même le meilleur moyen de savoir si la vidéo est visuellement acceptable ou non.

Inversement, un critère « automatisé » pose le problème de la cohérence avec la vision humaine, mais il serait a priori plus facilement quantifiable.

4.4.2. Indicateurs objectifs

Ce sont des paramètres de mesure de qualité vidéo **basés sur la qualité du signal image** , Une image peut être considérée comme un signal (image en niveaux de gris) ou un ensemble de signaux bidirectionnels (image couleur) $F(x, y)$ (voir [Adj01]). Ces signaux peuvent être donnés dans différents systèmes de couleurs : RGB, Yuv (voir la section 2.2.1)

Les indicateurs de la première catégorie quantifient la différence du signal bidirectionnel de l'image compressée puis décompressée par rapport au signal de référence.

En particulier, un indicateur très utilisé est le PNSR (Peak Signal to Noise Ratio - rapport) défini dans [Adj01]. Si on considère une image $F(x, y)$ de dimensions $N \times M$ et sa reconstituée après compression et décompression $\hat{F}(x, y)$, on peut définir l'erreur moyenne par pixel par :

$$RMSE = \sqrt{\frac{1}{NM} \sum_{x=1}^N \sum_{y=1}^M |F(i, j) - \hat{F}(i, j)|^2}$$

Le MSE (mean signal error - erreur moyenne du signal) = RMSE peut être considéré en soi comme un indicateur.

Si F est mono signal (image en niveaux de gris), F est un scalaire. Sinon, on considère en général

$F = \sum \sum |F_i(x; y) - F_j(x; y)|$, Où les i et j représentent les différentes composantes de l'espace de couleurs choisi. On peut également travailler indépendamment sur ces composantes. On définit alors le PNSR par :

$$PNSR(F \rightarrow \hat{F}) = 20 \log_{10} \frac{\text{valeur maximale d'un pixel}}{RMSE}$$

Où la valeur maximale d'un pixel est la valeur maximale de la plage de valeurs de la composante d'un pixel. Ainsi, dans le système RGB, on utilisera 255 comme valeur maximale des trois composantes R, G et B.

Ces indicateurs sont assez largement utilisés car leur complexité algorithmique est très faible. Néanmoins, ils reflètent en général assez mal la qualité réelle de la vidéo ou de l'image. En effet, l'oeil humain n'est pas forcément sensible aux mêmes distorsions (c'est-à-dire à un PNSR égal) selon le type de la distorsion ou selon la zone de l'image considérée. Ainsi, à PNSR égal, une distorsion est en général moins gênante à l'arrière- fond qu'à l'avant-plan.

4.5. Les caractéristiques des canaux sans fil

Avant que nous entamions l'étude des protocoles et mécanismes utilisés pour la transmission vidéo sur un Wlan, nous citons les plus importantes caractéristiques d'un canal Wlan.

- ✓ Les affaiblissements du signal sur le canal sans fil
- ✓ Les erreurs sur les canaux sans fil
- ✓ Les *burstiness* du trafic vidéo

• Impact des erreurs des canaux sans fil sur la qualité vidéo

La variabilité du trafic vidéo et la nature du canal sans fil affectent de façon directe la qualité vidéo lors d'une communication sans fil. Nous illustrons dans la section 4.7 leurs effets combinés sur la qualité de la vidéo. Nous simulons ceci en infligeant les erreurs sur la vidéo codée

Cependant, comme les décodeurs sont différents dans leur *error-resilience*, quelquefois peu d'erreurs ont le même effet. Nous introduisons dans la section 6 un ensemble de mécanismes de correction d'erreur et qui peuvent améliorer la transmission sans fil de la vidéo.

4.6. Mécanismes pour le flux vidéo

Pour une meilleure amélioration de performance de transmission, nous citons dans ce qui suit un ensemble de techniques et mécanismes hors standards (MPEG, IEEE802.11) pour une transmission vidéo encodée robuste. Nous ne rentrons pas dans les détails car, bien qu'étant complémentaires aux techniques et standard étudiés dans cette thèse (voir chapitres 1, 2), elles ne font pas l'objet du modèle de simulation développé dans la section 4.7.

A. Entrelacement des données

La définition de l'entrelacement donnée ici est tirée de [Sal99]. Sur les canaux à mémoire, l'entrelacement des données est utilisé pour se rapprocher du cas où le canal est sans mémoire.

Cette méthode vise à disperser les pertes de façon à ce qu'elles apparaissent comme indépendantes les unes des autres au niveau du décodage. Pratiquement, le processus d'entrelacement fonctionne de la manière suivante. Une matrice ($L \times n$) est remplie horizontalement par des symboles encodés. La transmission se fait ensuite colonne par colonne. Au récepteur, le schéma réciproque est appliqué. L'opération d'entrelacement est illustrée sur la figure 4.1. Avec l'entrelacement, deux symboles consécutifs d'un mot de code sont séparés de L transmissions. Ceci permet de diminuer l'effet de mémoire sur les performances de codage, au prix d'une augmentation notable de la latence de transmission. Pratiquement, concernant la mise en paquets de données vidéo pour la transmission sur un WLAN, l'entrelacement permet de disperser les effets des pertes de paquets à plusieurs images, plutôt que de corrompre une zone étendue d'une image unique. Ceci limite très sensiblement l'impact visuel des pertes, dans l'espace et dans le temps.

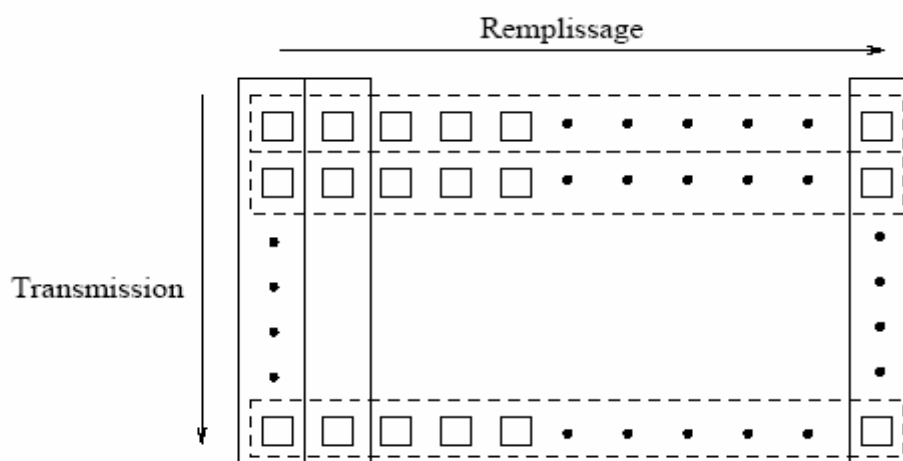


Fig.4.1. Illustration de processus d'entrelacement

B. Contrôle de pertes

L'objectif du contrôle d'erreur est d'augmenter la résistance aux pertes des flux transmis sur un canal. Nous présentons brièvement les techniques de contrôle d'erreur existantes.

B.1. Mécanismes d'ARQ

Les mécanismes d'ARQ *Automatic Repeat Request* [Hai97] consistent à répéter la transmission de paquets originaux perdus. Pour cela, la source doit disposer des numéros de séquence des paquets perdus, via la réception d'acquittements positifs ou négatifs.

L'ARQ semble inadaptée à des contraintes de transmission temps -réel. Elle peut cependant être utilisée sous forme hybride, associée à des codes correcteurs d'erreur (voir B.3)

B.2. Correction d'erreur par anticipation (FEC)

Les techniques de correction d'erreur par anticipation ou FEC consistent à rajouter de la redondance au signal à transmettre. Elles peuvent prendre la forme de codes correcteurs d'erreur ou bien consister à inclure au flux transmis des versions de données vidéo compressées à plus fort taux que les données du flux principal, et donc de moindre qualité.

Codes correcteurs d'erreurs Les FEC à codes correcteurs d'erreurs utilisées sur l'Internet sont souvent basées sur l'utilisation de **codes de parité** ou de **codes en blocs**.

B.3. Hybridation ARQ/FEC

Les FEC ne garantissant pas des taux de pertes nuls, elles sont souvent associées l'ARQ pour pallier les pertes résiduelles. On appelle ce schéma ARQ-hybride. Le but est de limiter au maximum le nombre de renvois de paquets par la source. Le processus peut prendre plusieurs formes.

Une première forme d'ARQ hybride consiste à demander retransmission de paquets non corrigés par les FEC. Une autre méthode, plus adaptée multicast, consiste à indiquer à la source le nombre de pertes détectées par le récepteur. Ainsi, le codeur peut renvoyer des données de redondances utiles à plusieurs récepteurs.

C. Contrôle de congestion et de débit

Les techniques de contrôle de congestion et de débit [Fab01] ont pour but de limiter le nombre de pertes de paquets en ajustant la contrainte de débit du codeur vidéo à la bande passante disponible sur le réseau. Le principe de base du contrôle de congestion point à point (ou unicast) est un processus adaptatif mettant en oeuvre une source et un récepteur pour contrôler le débit de la source. En observant l'état du réseau, la paire émetteur-récepteur peut détecter des congestions et réagir en diminuant le débit de la source. Inversement, un état de sous-utilisation de la bande passante du réseau conduit à une augmentation du débit pour mieux exploiter les ressources disponibles. Dans ce mécanisme, le récepteur est responsable d'observer l'état du réseau et de le rapporter périodiquement via des messages de retour à l'émetteur. Ce dernier utilise ce retour pour déterminer le bon débit cible à émettre, et régler les paramètres du codeur vidéo afin de l'atteindre. L'observation de l'état du réseau doit être aussi stable que possible.

4.7. Le modèle de simulation

Dans le but d'évaluer la transmission vidéo sur 802.11b, nous définissons le système illustré par la figure 4.2. L'idée de base est d'utiliser des traces vidéo définies précédemment (on spécifie les caractéristiques d'une séquence Mpeg dans le fichier trace). Une fois la trace vidéo est préparée on peut la faire entrer dans le simulateur après la précision des différents scénarios de simulation et nous réalisons les simulations possibles. Ensuite nous aurons comme résultats des simulations des fichiers trace. Ces fichiers seront utilisés par la suite pour extraire les profils de retard et de perte qui permettent de calculer les valeurs de PNSRs qui représentent la qualité vidéo à la réception.

En changeant les caractéristiques d'un flux vidéo dans la trace (la taille, le format, le type...) et le nombre de source vidéo ainsi que les caractéristiques des liens sans fil (les erreurs, débit, ...) nous aurons un ensemble de scénarios de simulation qui nous permet de mesurer la qualité vidéo (PNSRS) pour chaque scénario.

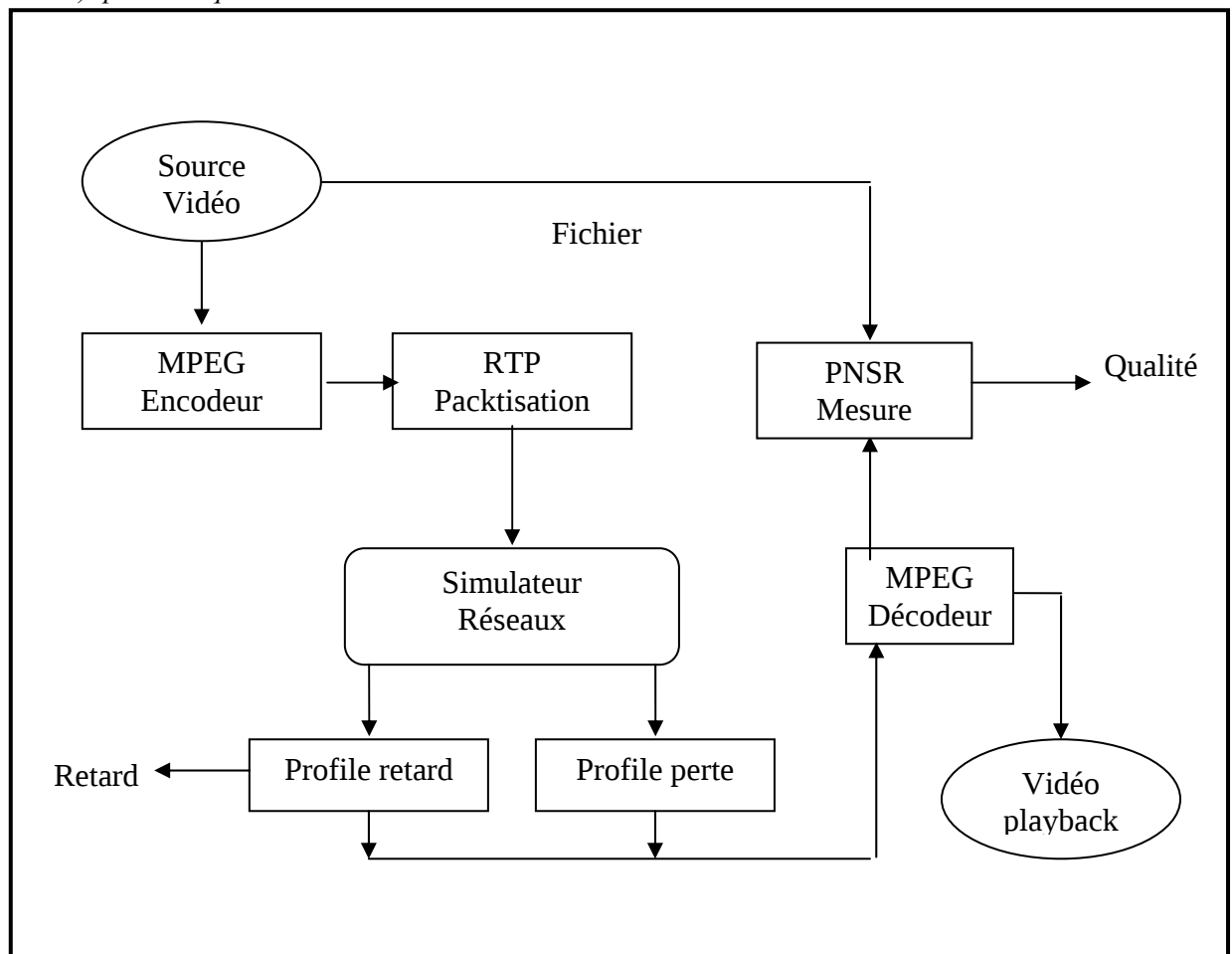


Fig 4.2. Modèle de simulation

Nous nous sommes basés sur le modèle présenté ci dessus (**Fig 4.2**), pour la définition du diagramme de flux de contrôle et de données présenté dans la **Fig 4.3**, qui donne plus de précision sur les étapes et procédures du modèle de simulation d'une transmission vidéo sur des liaisons sans fil.

Les étapes à suivre dans notre système sont :

- 1- Préparer la séquence vidéo sous forme de trace (préciser le format QCIF dans la trace - le rang de séquences : QCIF [4 :2 :0]-).
- 2- Générer le fichier trace fiche- **YUV** (préciser les caractéristiques d'une séquence vidéo **YUV** - voir chapitre MPEG- dans le fichier trace préparer dans l'étape précédente).
- 3- Générer le fichier « **RTP-fiche** ».
- 4- Générer le fichier trace
- 5- Entrer le fichier trace dans le simulateur et exécuter les simulations
- 6- Récupérer le fichier « **output.tr** » (trace) ensuite le traiter pour extraire le profil des paquets en retard.

- 7- Pour la mesure de qualité nous utilisons les fichiers « **RTP-Fich** » généré durant la phase 3 et les fichiers profils **perte/retard**.
- 8- Finalement, générer les résultats de la séquence vidéo et les valeurs **PNSRs**.

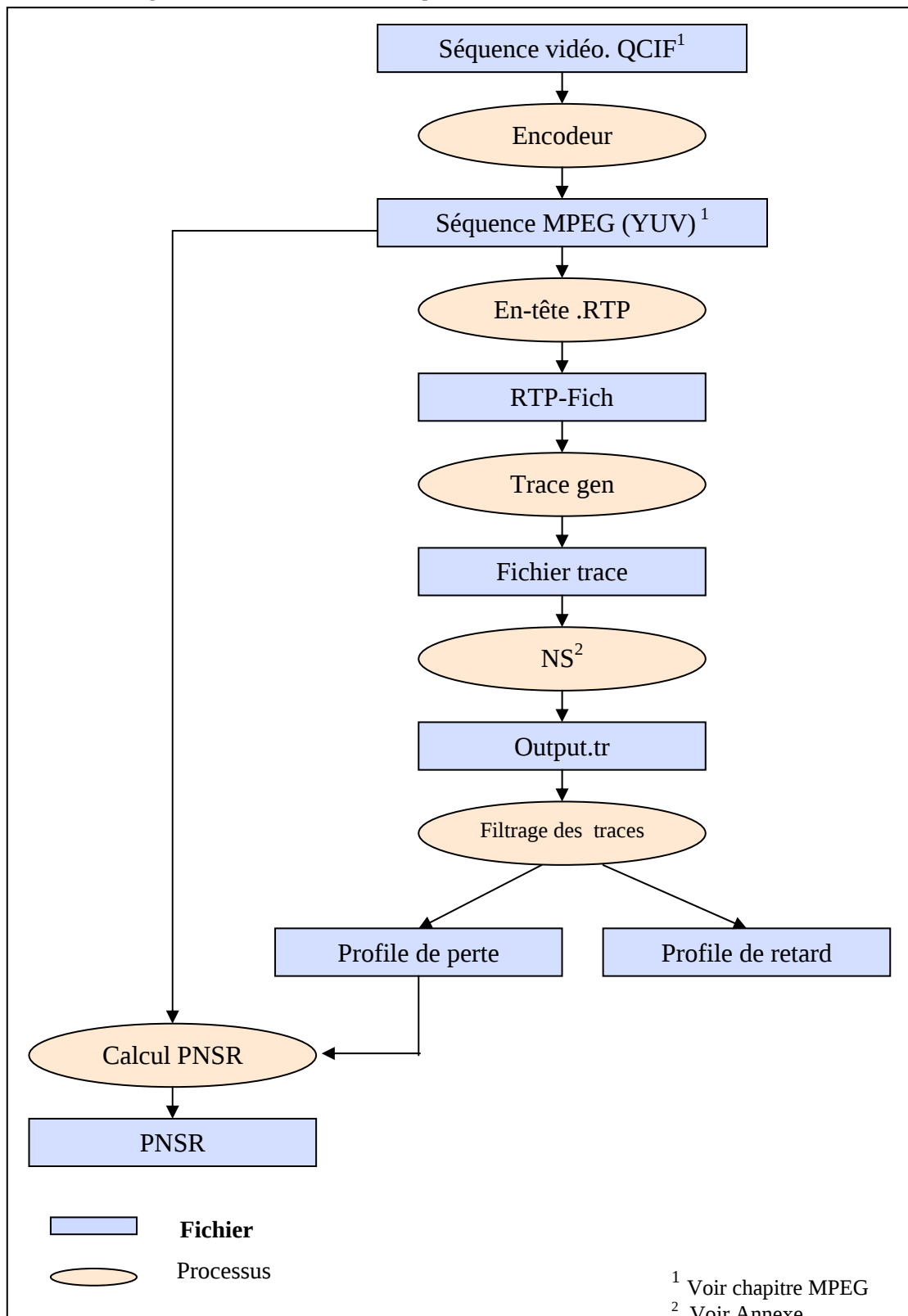


Fig 4.3. Diagramme de flux de contrôle et de données.

4.7.1. Génération de la trace d'une source vidéo

Pour faire le traitement du comportement d'une séquence vidéo, on doit générer le fichier trace du trafic pour le faire entrer dans le NS (Network Simulation). Ce fichier trace est généré à la base du fichier « *trace-gen* » et le fichier « *RTP-Fich* ». Le rôle du fichier « *trace-gen* » est de lire dans le fichier *RTP-Fich* afin de récupérer la taille des paquets RTP et les informations « *time stamp* » du fichier *RTP-Fich*. Ces informations seront utilisées dans la création et le remplissage du nouveau fichier trace, celui ci aura un format générique et donc peut être utilisé dans différents types de simulateurs.

Tous les simulateurs peuvent produire un profil de perte de paquets comme suit :

< paquet sequence number > < paquet size > <delay /loss >

Le décodeur dans MPEG a la dissimulation de l'erreur (error concealment) pour prendre soin des pertes des données. Cependant il n'y a pas une interface commune prête entre le simulateur et l'outil **MPEG**. Par conséquent pour refléter la perte de paquets, le schéma est proposé pour prendre en compte la perte de paquets :

Traiter le fichier « *RTP-Fich* » avant le faire entrer dans le décodeur durant le traitement, nous déplaçons ces paquets dans le fichier « *RTP-Fich* » si le paquet correspondant est marqué comme perdu dans la simulation, nous écrivons ensuite le fichier résultat comme un nouveau paquet « *RTP-Fich* ».

Le PSNR est le paramètre de mesure et d'évaluation de la qualité vidéo, cependant PSNR ne peut pas prendre en compte les paquets arrivés en retard. Cela paraît raisonnable dans les résultats parce que d'une part l'encodeur/décodeur ne peut pas attendre pour une longue durée et de l'autre la taille du buffer n'est pas assez importante pour stocker les paquets arrivés. Ce ci n'est pas convenable pour les applications temps réels, cependant il faut définir un seuil dans le programme « *RTP-Fich* » [Woo03] pour marquer comme perdu tous les paquets si leur retard dépassent ce seuil.

4.7.2. Topologie utilisée

Pour réaliser nos scénarios de test sur les applications vidéo nous considérons la topologie suivante :

Nous supposons aussi que :

- Toute transmission sur le médium sans fil est de 11Mb/s ou 5.5Mb/s.
- Le débit entre le PA et l'autre partie du réseau est de 100 Mb/s.
- Le délai total *end to end*: **mobile-PA** ≤ 25ms.

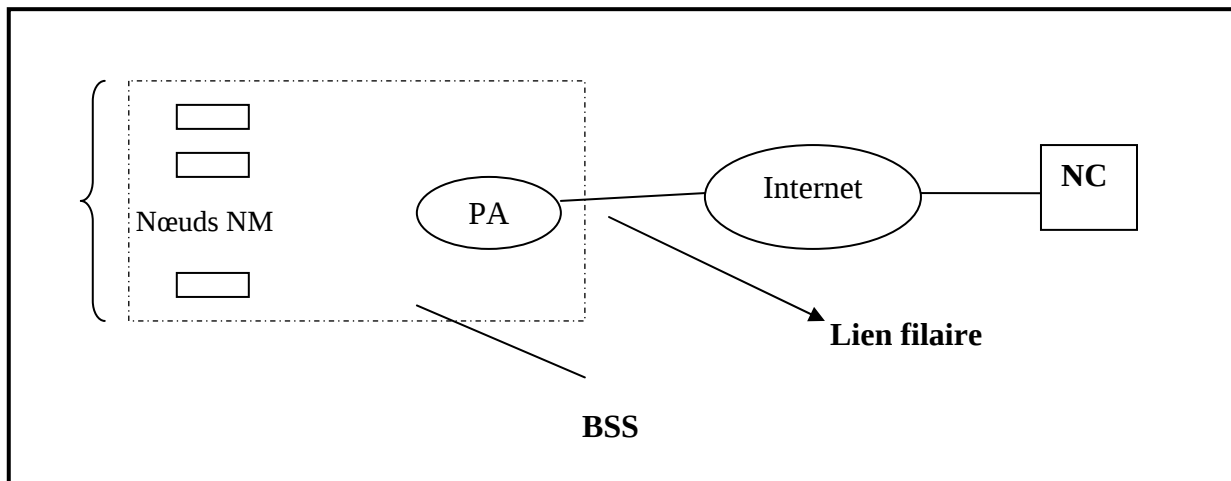


Fig 4.4. *La topologie de simulation*

PA : Point d'accès

BSS : base station set

NM : Nœud Mobile

NC : nœud correspondant

La mise on œuvre des différents scénarios sur la topologie présenté précédemment implique la réalisation les étapes suivantes :

- 1) La mise on œuvre des spécifications sans fil couche MAC / PHYSIQUE
- 2) La création d'un nouveau simulateur et le fichier trace
- 3) La création des PAs et la mise des spécifications (le mode PCF).
- 4) La création du nœud vidéo (Notons que le mode PCF est utilisé et que le nœud vidéo doit être ajouté au « polling list » du PA.)
- 5) Trafic vidéo : Génération de la trace vidéo (séquencece.tr)¹ et son attachement au nœud vidéo.

4.7.3. Analyse des résultats de simulation

Les résultats de simulation sont filtrées pour récupérer toutes les données telle que la trace UDP, nous exécutons des traitements semi-automatiques pour générer le profil du retard et de perte, à partir de ces profils nous allons créer le profil PSNR et nous tracerons les graphes pertinents dont nous aurons besoin pour l'analyse.

Nous avons simulé sous un débit de 11Mbps et 5.5Mbps dans le mode PCF en modifiant à chaque fois le nombre de sources vidéo, taille des paquets et la probabilité d'erreur sur le canal.

Les graphes suivants représentent les résultats des différents scénarios de la simulation.

¹ Pour le choix des séquence on peut consulter le site www.acticom.info.de On trouve de fichier de caractéristiques de séquences voir l'annexe pour plus de détail.

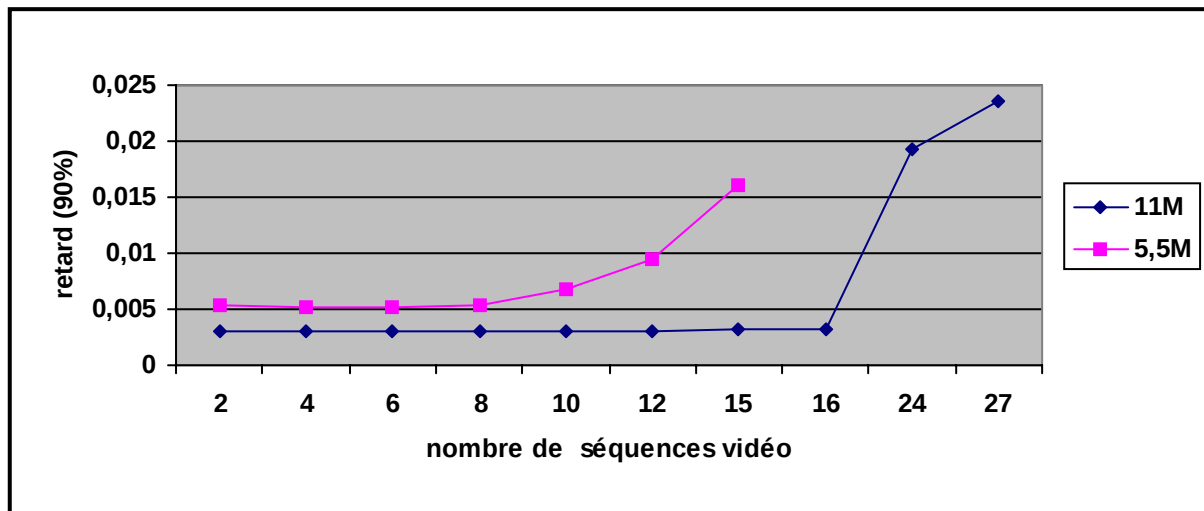


Fig 4.5 Retard des paquets en fonction des séquences vidéo (retard 90%)

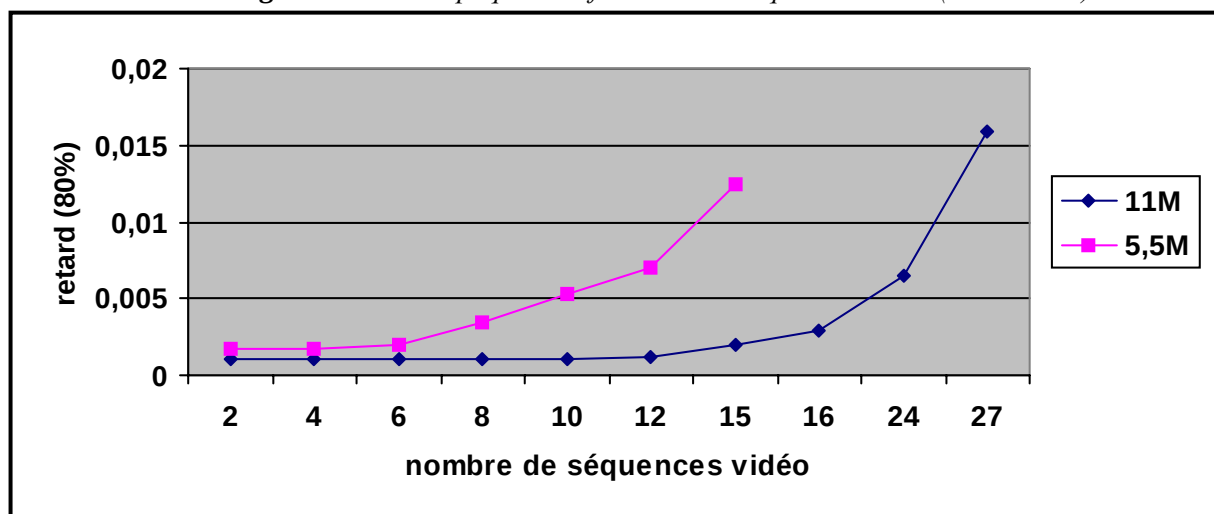


Fig 4.6. Retard des paquets en fonction des séquences vidéo

Remarque: 90%,80% : représente le pourcentages des paquets arrivés avec un certain retard.

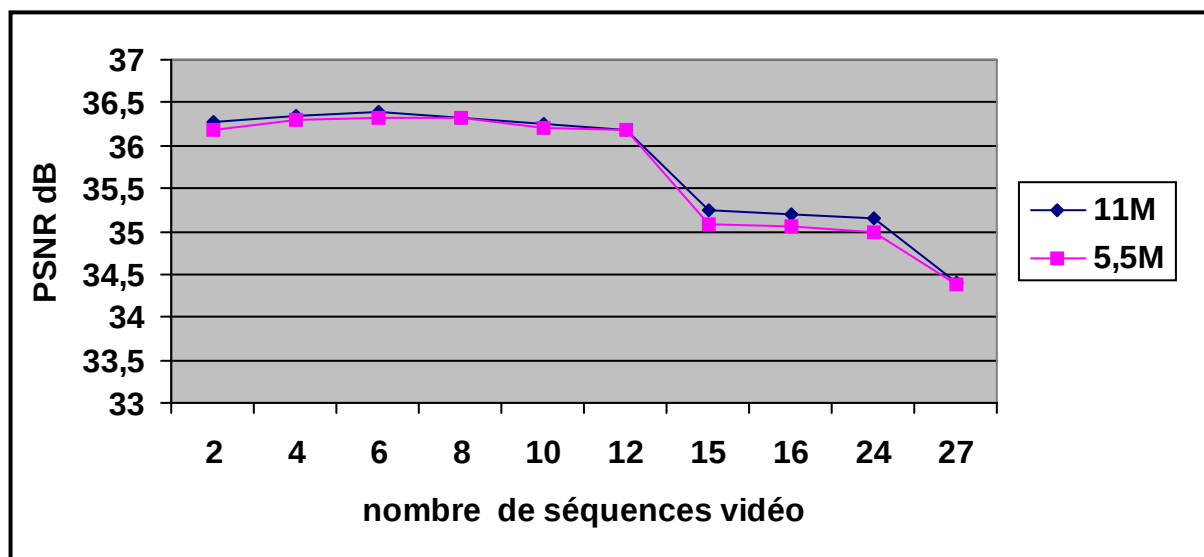


Fig 4.7. Qualité vidéo PSNR en fonction des séquences vidéo

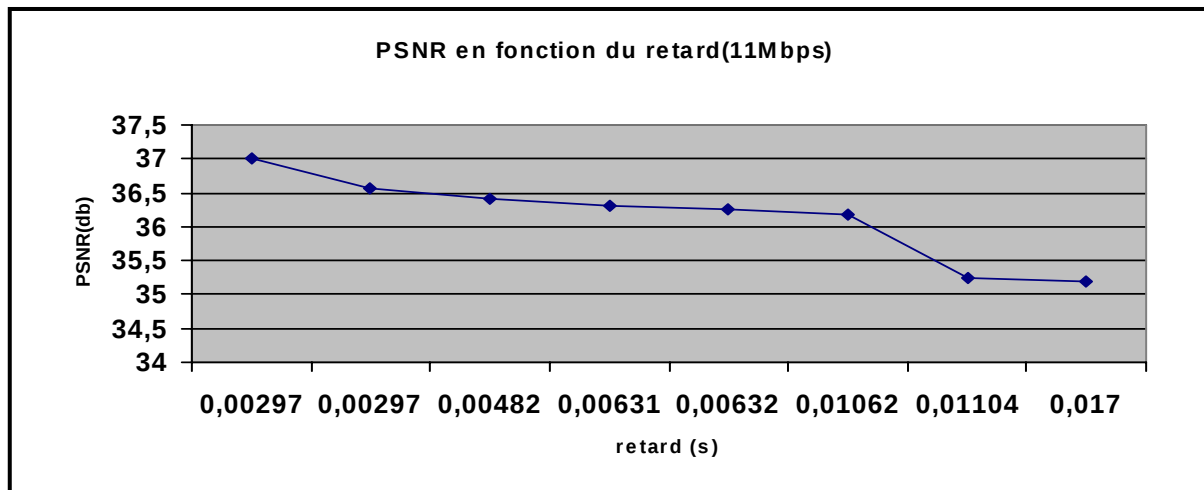


Fig 4.8. PSNR en fonction du retard des paquets

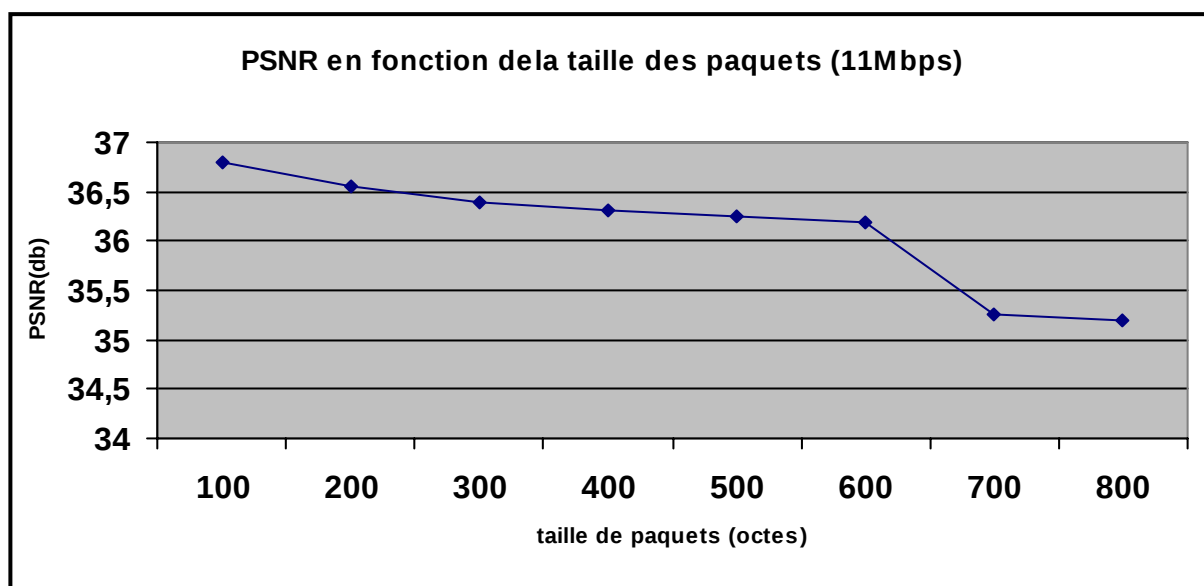


Fig 4.9. PSNR en fonction des tailles des paquets

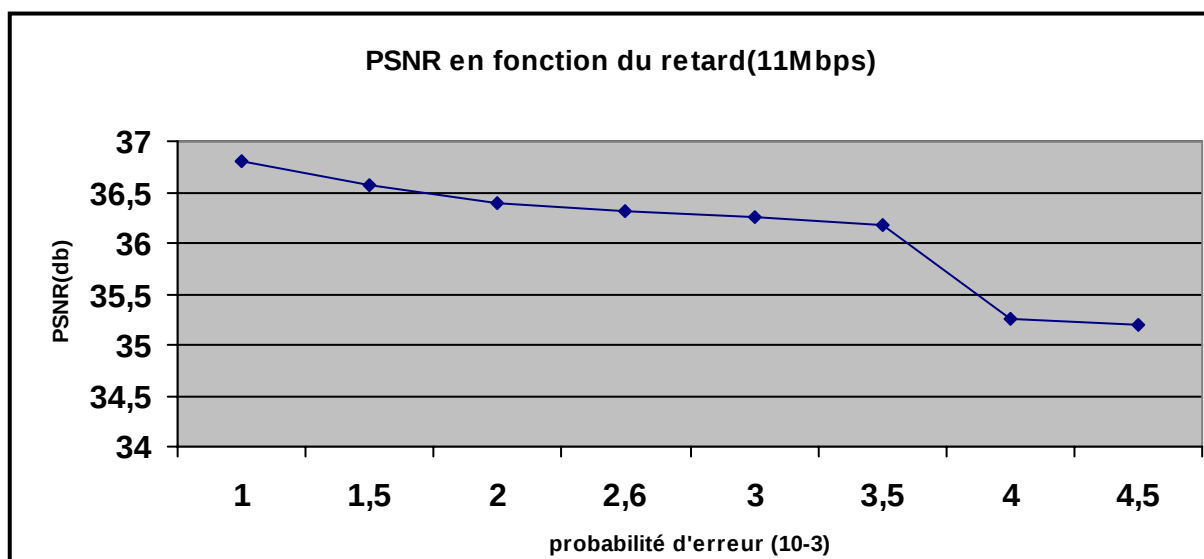


Fig 4.10. PSNR en fonction des probabilités d'erreurs

Interprétation et analyse des résultats de simulation

Avant d'interpréter les figures présentées précédemment nous poserons les points suivants :

- Le délai 80%, 90% équivalent à dire : 80% ou 90% des paquets arrivent avec un certain retard.
- Les canaux radio 802.11 sont modélisés en modèle de Gilbert (Gilbert suppose deux états du lien de transmission sans fil (**GOOD, BAD**)).
- La probabilité d'erreur: c'est la probabilité d'erreur sur le canal en se basant sur le modèle de Gilbert (canal=good → le paquet est reçu correctement, sinon il y a une erreur sur le canal).
- Le standard de compression utilisé est le **MPEG AVC part 10(H 264)** (sa plus importante caractéristique : la perte d'un macrobloc n'influx pas sur le décodage des autres)
- Les séquences de test sont des petites séquences en format (Cif, Qcif)
- la probabilité d'erreur est affectée par les conditions des canaux.

Les résultats de la simulation sont récupérés pour les données de la trace UDP pour les différents scénarios, nous avons regroupé les différents résultats dans les 6 graphes précédents après un traitement semi-automatique pour récupérer les profils de pertes des paquets et du retard, à partir de ces profils nous avons tracé les graphes dont nous avons besoin

A partir des graphes (1, 2, 3, 4), nous pouvons observer clairement que dans un Wlan, le PCF achève d'une manière acceptable (et plus rapidement la communication (nœud- PA) avec des petits retards, mais il faut dire que le schéma PCF est limité et que beaucoup d'améliorations peuvent lui être apportées (voir les travaux réalisés pour le 802.11 e). Une autre remarque concerne l'ajout d'un autre trafic (audio, données) qui influence d'une façon directe sur la qualité de la vidéo, et donc induit à une dégradation importante des valeurs du PSNRs.

Pour étudier l'impact des erreurs de transmission, nous choisissons une séquence vidéo en petit format (QCIF), et nous envoyons cette séquence plusieurs fois (taille des paquets UDP est de 800Ko, la pile protocolaire IP/UDP/RTP, utilisation du buffer dans le récepteur)

Les graphes (5, 6) illustrent que dans le système de transmission le PSNR se dégrade en fonction de la valeur de la probabilité d'erreur sur le canal et la taille des paquets. Si la probabilité d'erreur est importante ou la taille des paquets émis sur le canal est importante une dégradation des valeurs du PSNRs est aperçue (moins 36DB) mais restent toujours acceptable. Ce ci qui confirme l'impact direct de ces deux paramètres (tailles des paquets, probabilité d'erreur) sur la qualité de la vidéo à la réception.

Pour pouvoir augmenter la tailles des paquets sans causer des augmentations dans la probabilité d'erreur il existe des techniques de compression d'entête (MAC, IP) qui seront nécessaires pour une

telle transmission. Nous devons toujours noter l'utilisation des formats vidéo CIF et QCIF (séquence de petite taille).

En plus nous pouvons aussi constater que la probabilité d'erreur sur le canal influe de façon directe sur les retards des paquets, si cette probabilité est grande, une augmentation du taux de retard est aperçue, celle-ci implique une dégradation des valeurs PSNR.

Nous remarquons un comportement similaire dans le cas d'envoi de plusieurs séquences en parallèle (Fig. 4.3), à chaque fois que nous augmentons le nombre de sources vidéo la qualité de la vidéo reçue se dégrade.

4.8. Conclusion

Dans cette partie nous avons établi une relation entre le **WLAN 802.11** et l'outil **MPEG AVC** pour définir le modèle de simulation en se basant sur les modifications apportées par [woo03] sur le schéma d'encodage et de décodage pour réaliser la simulation d'un trafic vidéo sur un WLAN et nous avons mesuré les performances en termes de retard, de perte et de qualité vidéo (**PSNR**).

Nous avons présenté dans cette partie le comportement de la vidéo MPEG sur un WLAN selon différents scénarios et analyses, nous remarquons l'influence de quelques paramètres impliqués dans le modèle de simulation et du système de transmission sur la qualité de la vidéo à la réception. Les résultats fournissent une claire indication pour le choix des paramètres et des seuils pour une telle transmission (taille des paquets, probabilité d'erreur, valeur PSNR seuil). Nous voyons qu'une taille entre **700** et **800 Ko** est meilleure pour une telle transmission et nous pouvons conclure aussi:

- Le nombre de vidéos sources ne doit pas dépasser 20 sources pour éviter une augmentation significative du retard qui peut arriver aux paquets transmis sur le WLAN.
- Les contraintes de QOS consistent à envoyer des paquets avec un retard inférieur à 60 ms
- Nous pouvons définir un seuil acceptable de PSNR pour les séquences vidéo de petit format (32-34 DB).

Chapitre 5

Etude de performance et l'impact du handover 802.11 sur le trafic vidéo

5.1. Introduction

L'accès sans fil risque de devenir de plus en plus fréquent dans les prochaines années. Les utilisateurs mobiles demanderont évidemment des niveaux de qualité de service identiques à ceux offerts aux utilisateurs de postes fixes. Le handoff est le processus enclenché quand un nœud mobile (en cours de communication) change son point d'attache à l'Internet. Nous pouvons définir le handoff de la manière suivante : handoff de niveau 3 (couche IP) et handoff de niveau 2 (couche liaison). Dans cette partie nous étudions le handoff de la couche 2, ce processus est l'opération effectuée par un nœud mobile qui change de point d'accès sans fil, c'est-à-dire le passage d'un point d'accès à un autre sans changer de routeur. Ce handoff n'engendre pas un handoff de la couche supérieure. Les travaux de recherche concernant l'amélioration de ce type de handoff, sont présentés dans la partie 5.3

Quand un nœud mobile NM (nœud mobile) exécute le handover entre deux PAs (on considère dans cette partie la 2^{ème} catégorie de la mobilité où le NM change de point d'accès sans changer le retour d'accès). Le but dans ce chapitre est d'exécuter le handover couche liaison (L2) avec peu ou pas de dégradation de performances de la qualité de la vidéo à la réception sur un WLAN 802.11

Dans cette étude nous présentons une analyse de performances du handover L2 (couche liaison) et son impact en terme de perte et de retard sur le trafic vidéo. Le NM n'arrête pas la communication durant le handover selon [Olo03] il est capable d'envoyer et de recevoir des paquets via l'ancien (PA). La cause de la perte des paquets est en général due à la nature du WLAN ou au handover. Si un paquet sera perdu pendant la communication, l'influence sur la perte des paquets suivants diminue linéairement et 3 à 4 paquets seront affectées [Olo02], si plusieurs NMs sont concurrents pour l'accès à un nouveau PA le retard impliquer augmente.

Il existe un grand changement dans le comportement du handover parce que le NM commence souvent à chercher le nouveau PA, après il exécute le handover et durant le handover, plusieurs messages sont échangés entre les 2 phases du handover (phase de recherche, phase d'exécution)

5.2. Principe de base de cette partie

- 1- Dans cette étude nous avons limité l'étendue du WLAN où tous les PAs utilisés appartiennent au même (ESS).
- 2- Le NM prend l'avantage de lancer la handover avant de perdre la connectivité avec son ancien PA.
- 3- Le NM effectue le balayage actif pour la recherche du candidat PA.
- 4- Utiliser l'authentification ouverte « open authentication » pour le contrôle d'accès.
- 5- Minimiser la perte des paquets d'une transmission vidéo durant le handover L 2.

5.3. Related work

Il existe un nombre de recherche [Olo02][Olo03] [Val03][Mis03] qui concerne les performances du WLAN et le handover couche liaisons (handover L2) les deux buts principaux de cette étude est d'établir et de présenter une étude convenable des mesures de latence du handover dans WLAN IEEE 802.11b.

[Mis03] [Olo03] présentent des mesures pratiques du handover (utilisation d'un *Sniffer* des paquets *Netperf*). [Mis03] mesure le handover du nœud mobile qui se déplace autour d'un réseau WLAN. Le handover est mesuré par l'observation de la gestion 802.11b du trafic via un ensemble de 'mesures pratiques', ils ont observé que la latence du handover est significative, en plus la phase de recherche est le premier contributeur dans la latence du handover, ce handover dépend du type de la carte WLAN utilisée (PCMCA, Lucent STA, Cisco...).

[Vel03] propose de diviser le handover en 3 phases : phase de détection, phase de recherche, phase d'exécution.

[Vel03] mesure aussi la latence du handover par l'observation des messages échangés quand le NM se réassocie avec le nouveau PA, il réclame que la détection d'un mouvement est un grand contributeur à la latence de handover et de recherche pour réduire la perte durant le handoff.

D'après [Vel03] le handover obtenu en temps réel peut changer de manière significative cela est dû au *probe-wait latency* (c'est le temps d'attente sur un canal après le probe request)

Quelques paramètres sont nécessaires pour améliorer les tests :

- 1- **PRT** (probe response timeout) : le temps d'attente après l'active probe.
- 2- **PET (Probe Energy Timeout) temps d'attente pour un prob reponse**

Selon [Vel03] montre que la carte sans fil Cisco produit un handover avec **PRT= 35ms** et **PET=10ms**.

Le temps du handover peut être mesuré dans la théorie par l'envoi du flux vidéo entre le NM et le nœud correspondant NC ; quand le NM se déplace entre deux PA, On définit le temps handover théorique suivant

$$H_t = \text{nbr}(c)(\text{nbr de canaux libre}) * PET + \text{NBR}(C \text{ occupé}) * PRT$$

Le handover réel fournit un temps plus grand que dans la simulation ou la théorie. Avant d'implémenter et tester cet aspect, il faut bien préciser le fonctionnement de la couche Mac et la couche physique

[Woo03] fournit un overview pour l'implémentation d'un modèle d'ieee802.11b (couche MAC, Physique) en utilisant un simulateur. Le modèle est approprié pour l'évaluation de performance de wlan.

Le but du travail de [Vel03] et [Mis03] est l'analyse du temps du handover couche liaison 802.11b d'une (procédure handover). Les mesures indiquent que la phase de détection et de recherche sont les principaux contributeurs dans le latence handover. En plus, la réduction du temps de la phase de détection est possible avec les réactions rapides aux paquets perdus et l'utilisation des petits intervalles de *beacons*. Aussi pour la phase de recherche, on peut la réduire en utilisant la **recherche active**. Dans ce cas, on calcule les valeurs des deux *timers* qui contrôlent la durée de la **recherche active** dans le but de réduire le temps de recherche.

5.4. Procédure handover couche liaison

Quand un NM opère dans le mode de l'infrastructure il essaiera de s'associer avec un PA dans son voisinage. Chaque PA constitue une station de base et tout le trafic vers et du NM ira par le PA. En plus chaque cellule est contrôlée par un PA et forme un BSS, toutes les cellules (BBS) sont reliées à une épine dorsale système de distribution (DS) qui permettra le transfert de données au sein d'un même ESS ainsi que la communication avec des réseaux filaires. Pour couvrir une plus grande zone, les PAs peuvent être connectés par un DS et forme un service étendu (ESS).

Un NM qui change de zone de couverture (cellule) d'un PA pourrait se réassocier avec un autre PA (dans le même ESS), donc nécessite l'exécution d'un handover couche 2 (couche liaison). Les PAs sont configurés pour opérer sur des fréquences de différents canaux pour éviter les interférences entre les cellules. Le standard 802.11 ne spécifie pas la conception du DS [STD99], mais la solution est de connecter l'ensemble de PAs par un ou plusieurs ponts Ethernet comme montré dans Figure 5.1.

Quand un NM se déplace loin de son PA courant, la qualité du signal des messages reçus de PA diminue. À un seuil de la qualité signal (configurable), le NM commencera à chercher un meilleur PA pour se réassocier avec, donc déclencher une procédure du handover. Le standard spécifie qu'un NM

peut être associé avec seulement un PA au même temps, donc il y a un risque que la communication soit interrompue en exécutant le handover. La durée de la période quand le NM est incapable d'échanger le trafic des données par son ancien et nouveau PA est souvent connue sous le nom de *latence du handover* ou *délai du handover*. Cependant, la définition complète de la latence du handover a besoin d'être faite avec plus de soin.

Pour faciliter les explications, la figure 5.3 montre que les messages échangés représentent les parties classiques d'un handover.

Si la qualité du signal se dégrade durant la communication (NM avec son PA) sous un certain seuil, la procédure du handover se déclenche. La valeur du seuil du handover est configurée afin qu'un handover est déclenché avant que la connectivité avec le PA courant soit perdu, c'est ce qu'on appelle un *soft handover*. Par conséquent le temps de **détecter le mouvement n'affectera pas la latence total du handover couche liaison**. Pour trouver un PA candidat et le réassocier avec le NM nous commençons une *phase de recherche* (scan) cette phase consiste à parcourir les différents canaux de la radio; ce parcours peut être fait passivement (en écoutant les messages beacons des PAs) ou activement en envoyant un message (*prob request*) sur chaque canal et écouter sur ce canal pour un message *probe réponse* des PAs, (Figure 5.3.). Une autre amélioration est possible pour la recherche passive est de réduire l'intervalle du beacon à 50 ou 60 ms plutôt que 100ms [vel03].

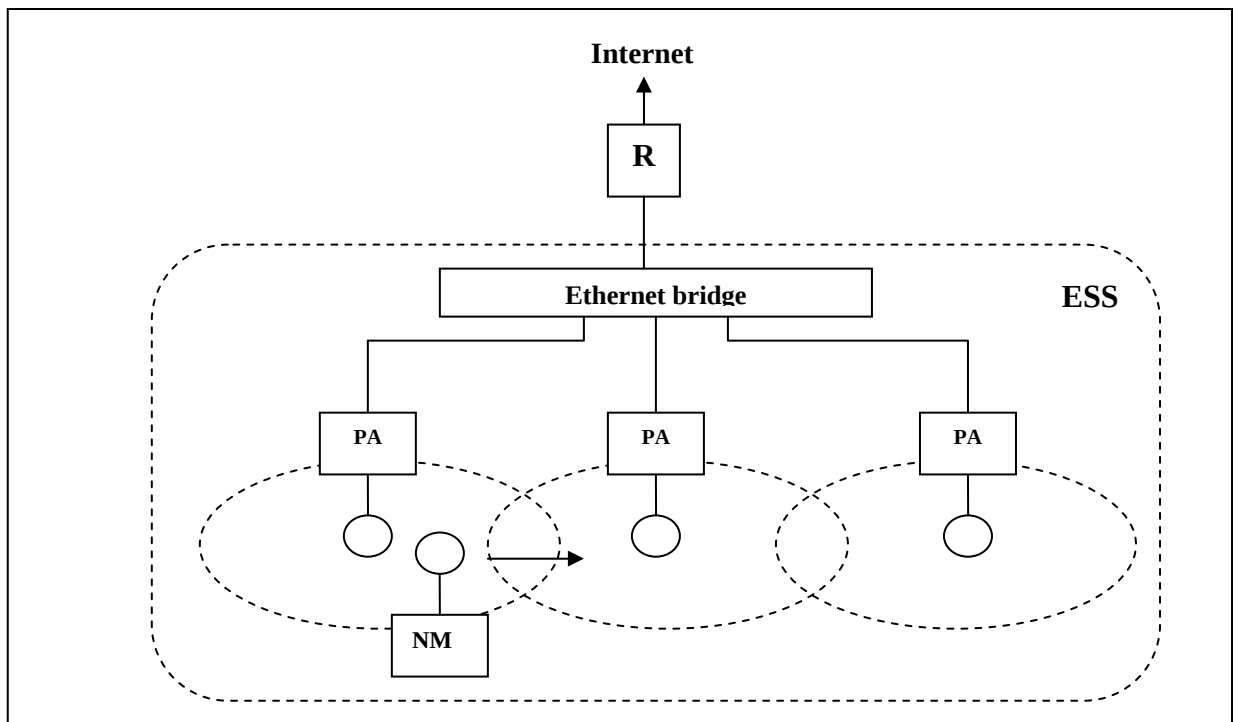


Fig 5.1. La topologie proposé pour l'exécution du handover L2.

PA : point d'accès

NM : Nœud Mobile

R : Routeur

ESS : Extended set service

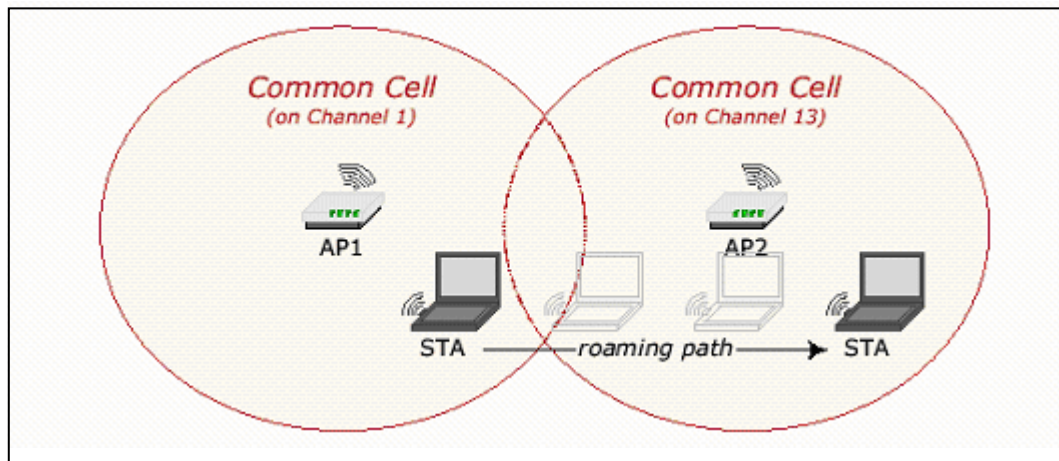


Fig5.2. La procédure de roaming.

Quand le NM a fini la recherche d'un nouveau candidat PA, il commencera la procédure de réassociation avec le meilleur candidat PA. Cette réassociation peut être divisée en deux sous procédures, l'échange du message de l'authentification, (messages 5 et 6 dans la figure 5.2) et l'échange du message du réassociation (message 7 et 8 dans figure 5.3). Nous faisons référence à tout de ceci comme la *phase d'exécution* du handover L2. Cependant, nous présentons après une description plus précise de la phase de l'exécution.

Il y a deux méthodes d'authentification différentes:

Système d'authentification ouverte : Il s'agit d'une authentification nulle, pour le principe. Celle-ci consiste seulement à fournir le bon SSID (Service Set ID) (figure 5.3).

Authentification par clé partagée (n'est pas considéré dans cette étude) lequel implique 4 messages entre le NM et le nouveau PA. Le NM est permis de s'authentifier avec multiple PAs, et le NM peut faire donc en gardant son association avec l'ancien PA (connu comme pre - authentification). Par la suite la réassociation avec le nouveau PA, le NM, pourrait envoyer aussi un message Dé association à son ancien PA.

Si nous considérons la figure 5.3 et supposons que le NM n'est pas capable d'envoyer ou recevoir les paquets pendant la phase de recherche ou d'exécution, alors la latence du handover sera égale à la somme de la latence de recherche et d'exécution.

Mesurer la latence du handover, revient à mesurer le temps entre la décision de déclencher un handover jusqu'à la réception du message *Réponse Réassociation* (après la réception de ce message tous les paquets sauvegardés dans des buffers seront envoyés au nouveau PA via un tunnel), nous pouvons bien estimer la latence du handover en dirigeant la gestion des paquets de données qui sont envoyés entre le NM et les PAs sur les liens de Wlan.

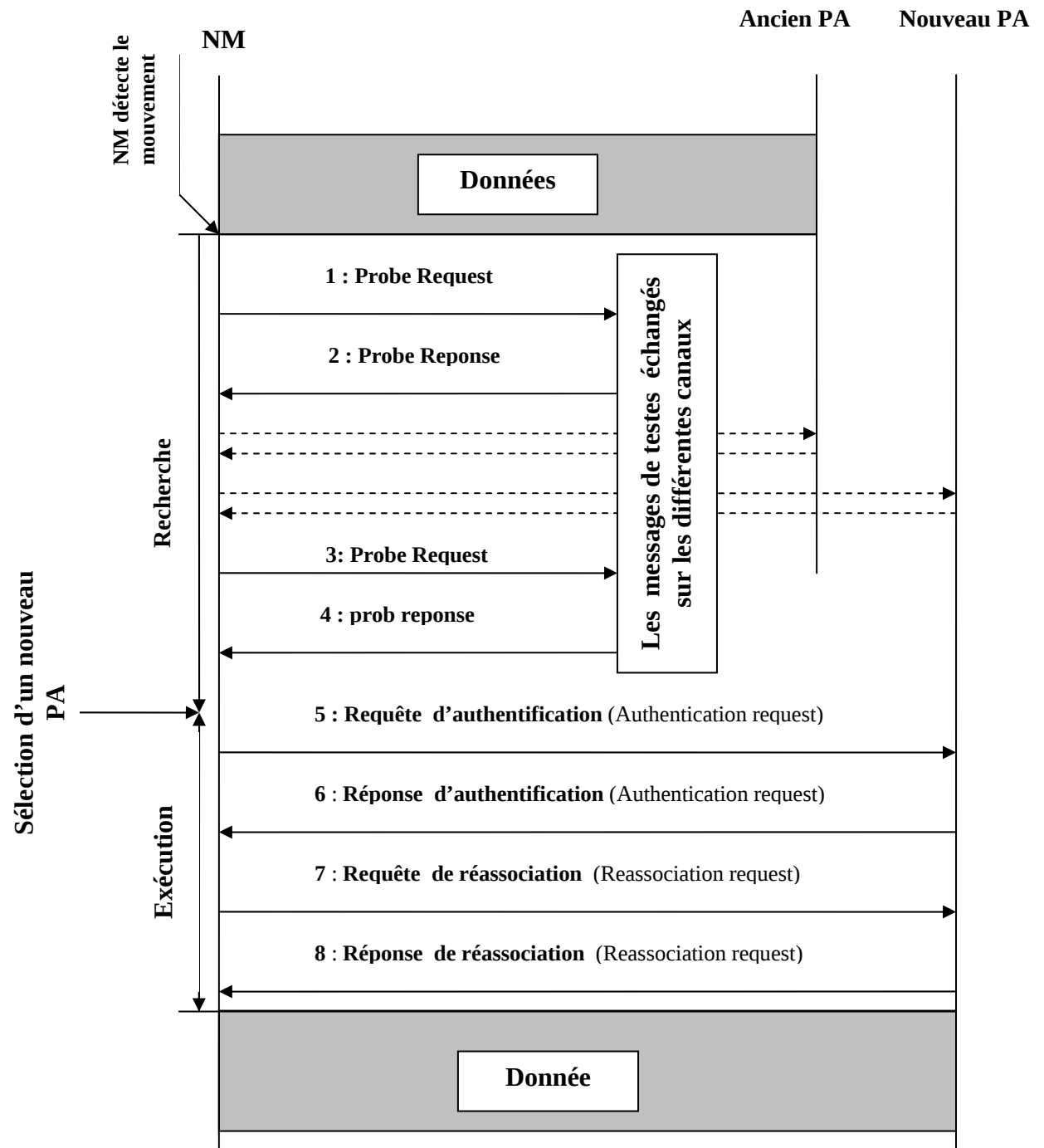


Fig 5.3. Handover L2 classique.

Cependant, quand nous analysons l'échange de messages dans la figure 5.3, il n'est pas vraiment similaire à l'échange de message réel vu sur les liens du Wlan, et que le NM était capable d'échanger quelques paquets pendant la période où commence la phase de recherche jusqu'à la fin de la phase de l'exécution.

Nous proposons dans la figure ci-dessous (figure 5.4) une amélioration du schéma du handover proposé précédemment (Figure 5.3).

Si le NM transmet les données de façon continue, les paquets du flux ascendant (ce sont les paquets envoyés par le NM à un correspondant via le point d'accès) seront bufférisés dans le NM pendant les phases de recherche et d'exécution. Le scheduling des paquets ascendant, pour être envoyés après la phase de recherche (**1ere phase de buffering**), peuvent être faits à partir de l'ancien PA entre les phases de recherche et d'exécution ou envoyé au nouveau PA après la phase de l'exécution (**2 phase de buffering**).

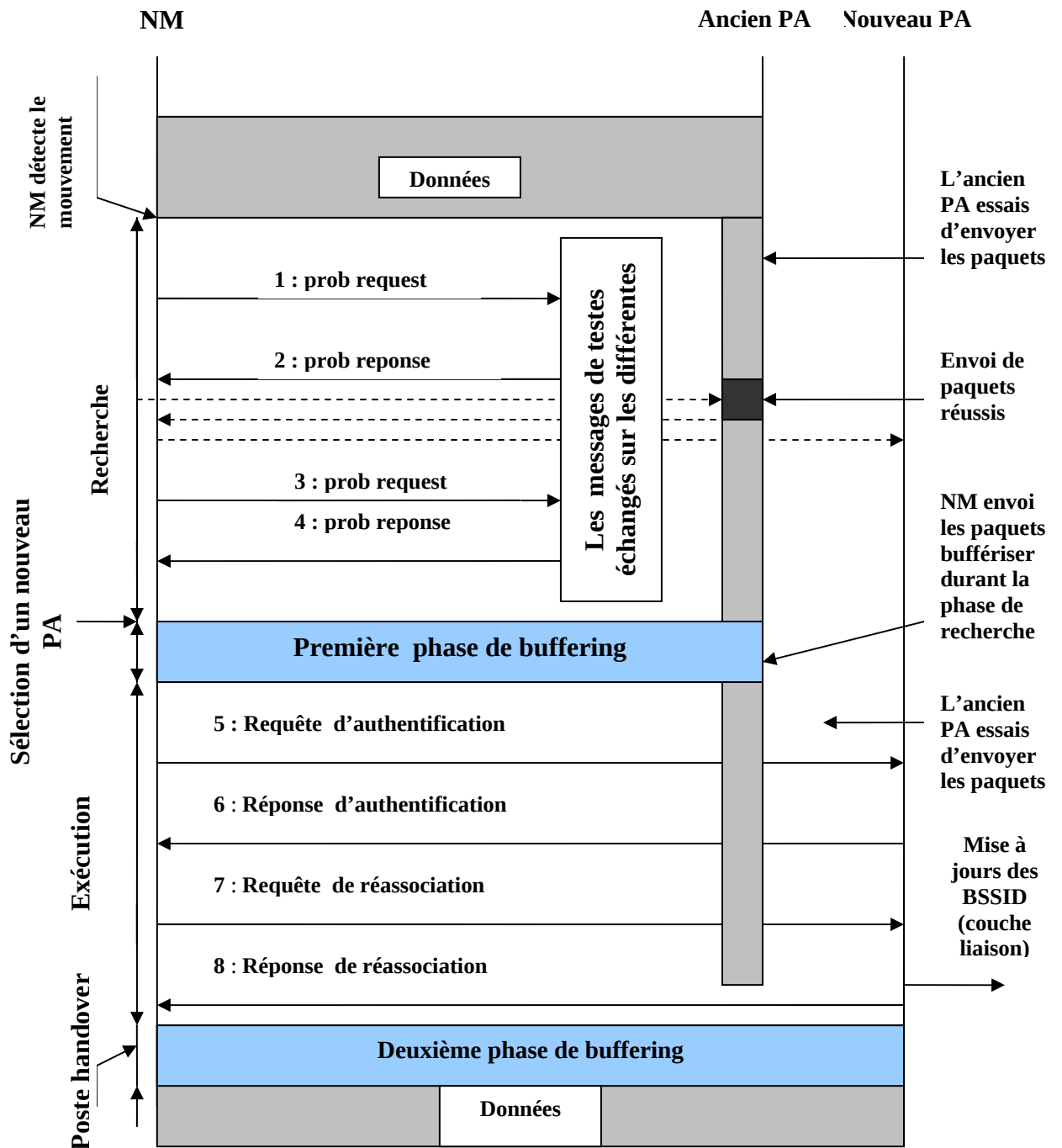


Fig 5.4. Handover L2 proposé.

Nous définissons ces deux phases de buffering comme la phase du **1ere phase de buffering** et la **2 phase de buffering** (figure 5.4). Autre amélioration possible concerne l'introduction **du message de mise à jour couche liaison c'est un message de mise à jour des adresses BSSID**.

Si nous considérons le cas opposé où l'ancien PA reçoit les données pour le NM (cas du trafic descendant), le PA courant peut sauvegarder ces paquets dans un buffer et les envoyer via un tunnel au nouveau PA.

A la fin de la phase de recherche le NM pourrait informer le PA pour qu'il envoie les paquets sauvegardés au buffer **« 1ere phase de buffering »** au NM avant que le NM Commence la phase d'exécution. Le PA essaiera de transmettre les paquets de données au NM comme ils arrivent. Pendant la phase de recherche il y a une chance que quelques paquets du flux descendant peuvent arriver au NM réellement. Cela peut se passer si le PA envoie ces paquets pendant que le NM parcourt sur le canal utilisé par l'ancien PA.

Le comportement du handover exact varie non seulement selon le matériel et logiciel utilisé, mais ce comportement peut varier s'il y a un flux ascendant ou descendant.

Alors Le handover n'est pas seulement l'intervalle de temps des différentes phases mais il dépend aussi de la direction du trafic. Pour cela nous allons définir quelques paramètres que nous supposons utiles :

La définition 1 (Intervalle Handover) L'intervalle du handover (T_{handover}) est l'intervalle de temps quand la transmission des paquets sur le lien WLAN peut être affecté (retard ou perdu) dû à l'exécution du handover.

La définition 2 (Données Handover) Ces paquets de données transmis sur le lien WLAN ou bufférisés au NM ou (ancien) PA pendant l'intervalle du handover.

La définition 3 (Retard du flux ascendant) Le délai imposé sur le flux des paquets ascendant durant le handover.

Definition 4 (Paquet Chanceux) Un trafic descendant de paquet des données transmis par l'ancien PA sur le lien WLAN pendant l'intervalle du handover et qui réussit à atteindre le NM.

5.5. Conception des PAs dans un Wlan 802.11

Nous sommes intéressés par le comportement du handover Ieee802.11b. Selon [Mis03] le problème c'est la phase de recherche (les messages de test toutes pour trouver le PA avec le meilleur signal). Durant le déplacement le NM détermine la puissance du signal du PA courant via les *paquets beacons* reçus, quand la portée du signal est en dessous de la valeur seuil du handover (L_2 calculé), le handover est initié.

La puissance de signal reçu par un NM (seuil de puissance de réception dans le cas optimal) S' , tout fragment reçu avec une puissance de signal inférieur au seuil S' sera considéré comme perdu.

$$S' = T/R^2$$

T : Puissance de transmission et

R : La distance entre PA et le NM

Dans cette étude nous considérons le handover des PAs attaché au même LAN et servir le même ESS (Fig5.1). Dans cette section nous discutons l'usage des ponts dans les PAs et le reste du système de distribution DS. Il devrait être noté que bien que nous ayons limité l'étendue de cette partie à la situation quand un NM se déplace dans un ESS, le cas quand le NM exécute le handover entre les PAs connectés aux LAN différents est aussi d'intérêt (voir les chapitres 6,7,8).

Comme mentionné auparavant, les PAs sont connectés aux ponts Ethernet communément et dans le même ESS. La réalisation des tests réels sur un réseau wifi (wlan 802.11) est possible grâce aux outils NetPerf/ IPerf, depuis qu'ils sont capables de récupérer les informations sur les paquets transmis dans le réseau (possibilités de récupérer dans un journal le comportement du trafic UDP, TCP, les débits).

Les ponts Ethernet ont la capacité d'apprendre où résident les adresses MAC différentes, afin que les paquets aient besoin d'être seulement envoyés sur le lien qui conduit vers la destination. Les Ponts stockent les informations sur les adresses MAC dans un cache d'information (cache station), ces adresses peuvent informer les dispositifs seulement si les noeuds transmettent le trafic activement ou si les entrées sont placées dans la table manuellement (par exemple par SNMP).

Les trafics à des destinations inconnues (*c.-à-d., quand l'adresse de la destination de la couche MAC est inconnue*). Le broadcast sera un dispositif important, depuis qu'il permet au trafic d'atteindre un NM qui a déplacé dans le LAN, même s'il ne transmet pas activement toutes les données. Les PAs sont fixes et ne se déplacent pas très souvent, et quand leurs déplacements sont possibles l'utilisateur sera informé. Cependant, un utilisateur du WLAN et au cours d'une communication ne tolère pas une interruption de l'ordre de minutes ou même de secondes. La façon d'informer le DS des changements dans le BSS n'est pas standardisée toujours, mais cela devrait être fait en ayant le nouveau PA et envoyant une mise à jour à la couche MAC, *c.-à-d., avec l'adresse MAC du NM comme adresse source*. Le flux descendant des données sera réacheminé au nouveau PA immédiatement. Bien que, nous ayons supposé que le nouveau PA envoie un message de mise jour de la couche MAC au DS.

5.6. Envoi de Flux ascendant et descendant durant le handover L2

Cette section présente une étude du comportement du handover et son impact durant un flux vidéo descendant ou ascendant. Nous nous intéressons au cas où un NM envoie ou reçoit un trafic vidéo au même temps que le NM exécute un handover

5.6. 1. Envoi du flux ascendant

Le comportement du handover durant le flux ascendant est représenté dans la figure 5.5, à un moment précis dans le temps le NM décide d'initier la procédure du handover et commence la recherche actif d'un nouveau candidat (PA) sur tous les canaux possibles par des messages « prob request ». Le scheduling des paquets d'un flux ascendant pour la transmission pendant cette phase de recherche est bufferisé (stockés) pour une transmission plus tard (exécuté un smooth handover en cas de perte). Après avoir fini la phase de recherche d'un candidat PA; le NM transmet un paquet seul (le premier paquet bufferisé) à l'ancien PA avant de commencer la phase de l'exécution.

La première phase de buffering est mesurée de l'instant d'observation du paquets « donnée k+1 » jusqu'à l'observation du premier message dans la phase de l'exécution.

Cela devrait donner une estimation raisonnable du temps ajoutée entre la phase de recherche et d'exécution pour envoyer le paquet intermédiaire (donnée K+1) à l'ancien ancien PA. Dans la phase d'exécution le NM peut transmettre (incorrectement) un message de réassociation initial qui pourra être rejeté par le nouveau PA. Après ces *faux messages de réassociation*, le message d'authentification est envoyé et le nouveau PA l'accuse positivement, ensuite la réassociation aura lieu. Après que la réponse du réassociation soit reçue par le NM qu'il essaie de commencer la deuxième phase de buffering par l'envoi du paquet « donnée k+2 » figure 5.5.

Les paquets « données (i) » ce sont des paquets du flux vidéo ; c.-à-d., les paquets bufferisé pendant l'intervalle du handover.

Le paquet donnée k+2 peut être ignoré par le nouveau PA (figure 5.5), et la raison pour ceci est que bien que ce paquet soit transmis sur le canal de la fréquence du nouveau PA, le paquet est adressé à l'ancien PA, c.-à-d., le paquet emporte le BSSID de l'ancien PA. Le NM communément retransmet ce paquet un nombre de fois avant de se déplacer au prochain paquet dans la file de paquets du bufferisé (donnée k+3). Le temps de la recherche a été estimé comme l'instant de l'observation du dernier paquet (donnée k) jusqu'à l'observation de paquet de données (donnée k+1) 1ere paquet bufferisé, (voir la fig 5.5).

Le dernier paquet envoyer avec succès avant le handover (donnée k) désigne le début de la 1ere phase du buffering (**en faveur de l'instant d'envoi du premier message prob request**).

le paquet de données (data k+1) dénote la fin de la phase de recherche; ce paquet est bufferisé dans le NM et sera transmis dès que possible après que la phase de la recherche soit terminée.

Améliorations possibles

Améliorer la performance du handover, il y a quelques améliorations possibles qu'on pourrait essayer. En premier, si les messages fausses *réassociation request* n'auraient pas été envoyés, la phase d'exécution pourrait être diminué par quelques ms en moyenne. De plus, si on utilise le

BSSID du nouveau PA le paquet donnée k+2 aurait été envoyé à nouveau PA. Alors, on pourrait prédire le retard causé par l'envoi du paquet donnée k+2 s'il avait été reçu correctement.

Améliorer les performance du handover même plus loin il aurait été préférable si tous les paquets bufferisés pendant la phase de la recherche peuvent être envoyés par l'ancien PA. Pour le NM il aurait été bénéfique si donnée k+1 et donnée k+2 (en plus de donnée k) ont été envoyés par l'ancien PA avant d'entamer la phase d'exécution. Alors le délai moyen prédit pour ces paquets (2 et 3) réduit le temps de l'exécution.

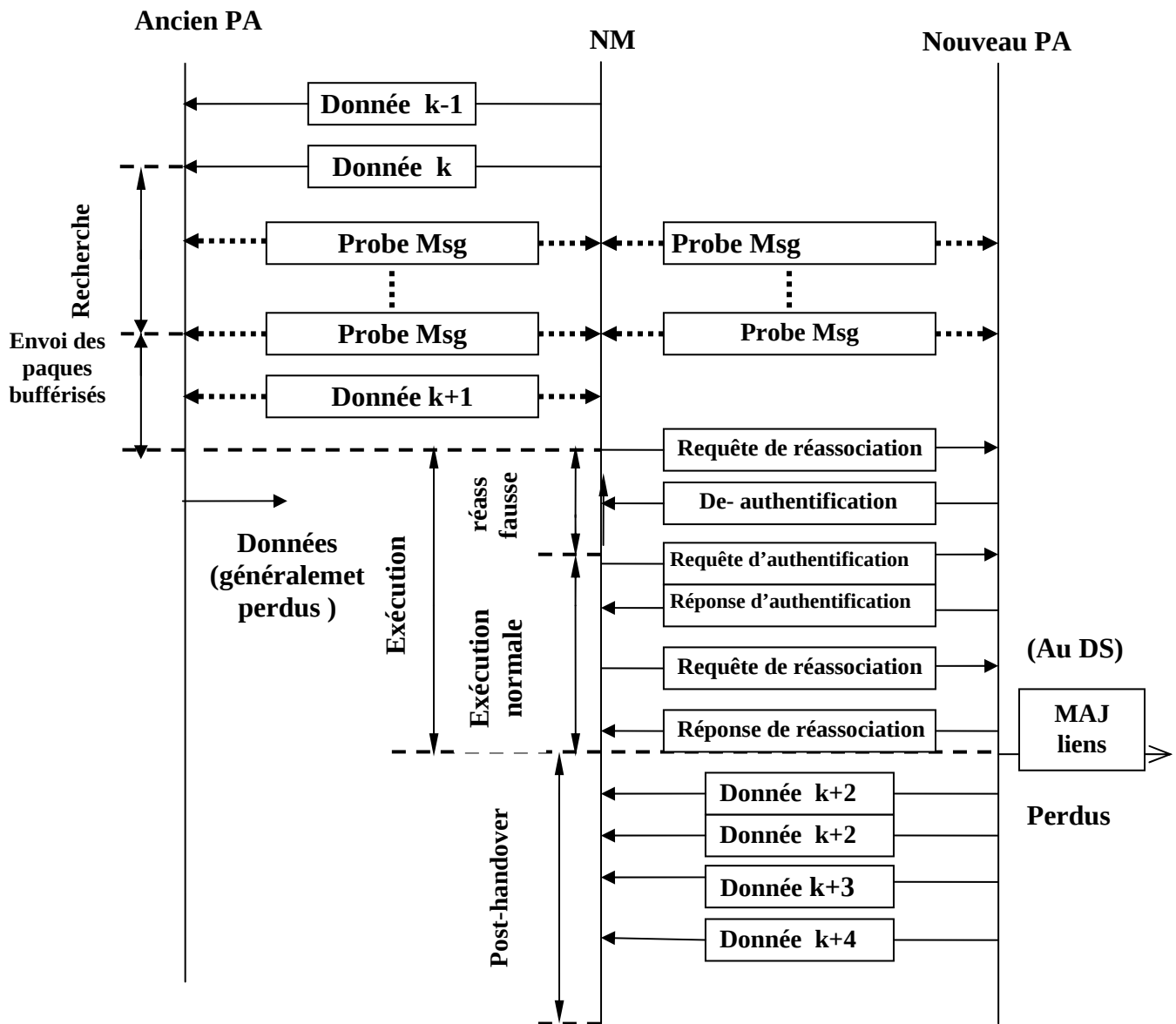


Fig 5.5. La procédure du handover pour un trafic ascendant.

5.6.2 Envoi du flux descendant

Quand le NM exécute un handover pendant la réception du trafic, l'ancien PA continue à envoyer les paquets arriver au NM pendant les phases de recherche et d'exécution, (voir la figure 5.6) .nous considérons qu'on a pas la première phase de buffering comme dans le flux ascendant, Une autre différence du cas ascendant est que la fin de la phase d'exécution est défini par l'observation du message de mise à jour de la couche MAC, depuis que c'est ce message (plutôt que le message réponse *réassociation*) cela réachemine des paquets du trafic descendant au nouveau PA. Le temps pour que le premier paquet arrive par le nouveau PA dépend de l'intervalle de la transmission (ce devrait être environ un demi intervalle de la transmission en moyenne). Les paquets de données devraient venir à un égal intervalle de transmission.

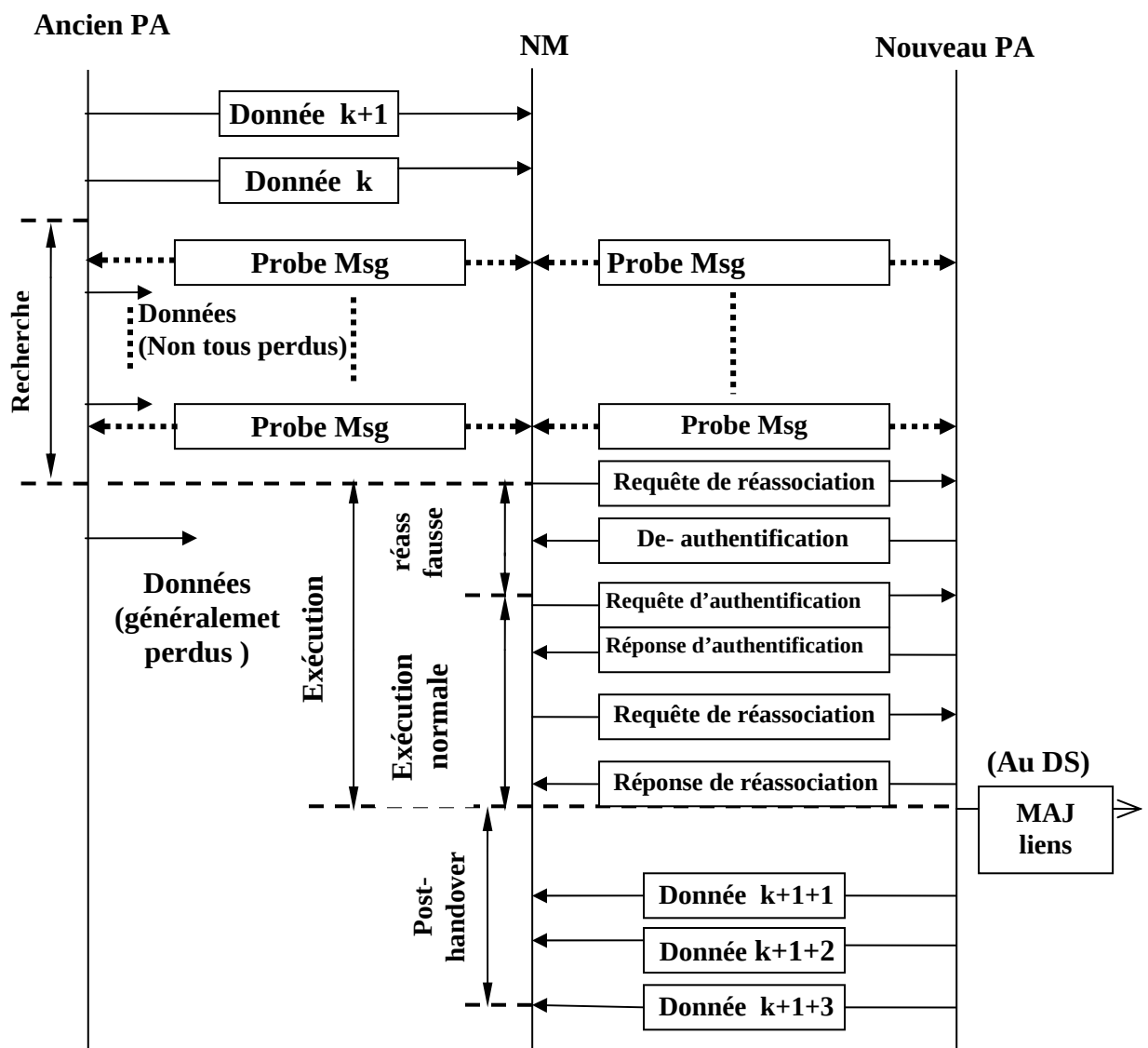


Fig 5.6. La procédure du handover pour un trafic descendant

La plupart des paquets du flux descendant durant le handover peuvent être perdus plutôt qu'en retard comme pour le cas ascendant.

La raison pour la réception des paquets chanceux dans le cas descendant est probablement que le NM en cherchant les PAs sur différent canal cherche dans le canal sur lequel que son ancien PA réside, et recevoir par conséquent ses paquets envoyé via l'ancien PA.

5.6.3. Évaluation de la portée des PA

Afin d'évaluer la portée du point d'accès et pour voir l'effet du handoff sur les performances du transfert de données, nous avons réalisé des tests pratiques, on utilise la topologie définie dans la figure 5.1.

Description des tests

Pour ce genre de mesures qui nécessite le déplacement du NM entre des PAs, le choix de l'outil le plus adéquat s'est porté à IPerf (permet de lancer en mode démon sur la station fixe de telle sorte qu'on peut lancer les tests à distance sans avoir besoin de relancer le serveur et qu'on peut récupérer les résultats sur la NM (qui sera un micro portable dell P IV). Les paramètres qu'on fait varier à chaque session de test sont la distance séparant le NM à la station de base et la taille des paquets transférés dans le cas d'un flux UDP.

- 1500 octets pour des distances ne dépassant pas les 100 mètres, elle correspond au MSS (Maximum Segment Size) déduite à partir de la MTU (Maximum Transmission Unit de ethernet à laquelle on enlève l'overhead ($IP+UDP=20+8$)).
- Pour les courtes et moyennes distances le débit croît avec l'augmentation de la taille du paquet car le surdébit devient de plus en plus négligeable par rapport à la charge utile lors des transmissions des gros paquets.
- Pour les distances limites du handoff le débit décroît avec l'augmentation de la taille du paquet car les paquets de tailles importantes sont plus susceptibles d'être corrompu si la durée d'occupation du support est assez longue (délais important aux handoff) à cause des variations brusques de l'état du support tel que l'apparition du bruit de mouvement des objets. En plus, pour les distances importantes le signal est atténué (la régression rapide en t/r^2 de la puissance du signal) et le niveau du bruit devient considérablement, pour mieux tester le handover il faut utiliser des distances plus grandes pour déterminer les vrais rapports du débit et déterminer la relation entre la taille des paquets et le handover
- On favorise les paquets de petites tailles (retransmission que de petits segments) du fait qu'il y a un grand taux d'erreur de réception du signal au niveau physique pour les distances relativement importantes.

- Généralement, le débit diminue lorsque le système procède au mécanisme de fragmentation appliqué sur les paquets de tailles supérieures à la MTU.

Interprétations :

- A l'approche du handover dans notre cas les BSS sont fortement chevauchés pour déterminer les position critiques du handover il faut minimiser le chevauchement dans notre cas si le signal atteint une certaine puissance (calculer) le handover est déclenché sans casser les lien avec l'ancien PA (un soft handover) théoriquement on peut dit que le débit chute considérablement si la distance dépasse une certaine valeur , cela est dû à l'importance des taux d'erreur pour ces distances qui provoque des retransmissions excessives. Pour remédier à ces retransmissions qui chargent le réseau sans fil surtout s'il est partagé par plusieurs utilisateurs, les stations en cause, peuvent baisser automatiquement leur bande passante d'émission. En effet, une interface configurée en mode bande passante automatique, peut basculer d'une capacité de transfert de 11 Mbits à 5 Mbits quand elle s'aperçoit qu'il y a beaucoup d'erreurs. Cela constitue une explication de l'effet de la chute brusque du débit pour les grandes distances,. En plus, cet effet à été confirmé par l'outil de configuration de la carte Wireless qui affiche instantanément la capacité effective d'émission sur l'interface mise en mode auto.

Discussions

- Pour UDP, les différences entre les débits des deux types de flux commencent à être apparentes à partir d'une certaine taille de paquet (800 octets d'après le chapitre précédent) pour laquelle l'interface sans fil fragmente les paquets en plusieurs fragments 802.11, cela augmente les délais de transmissions sur l'interface wireless du point d'accès et par conséquent diminue le débit du trafic descendant vers cette interface. En effet cette fragmentation de trames ne peut s'effectuer que dans le sens filaire → sans fil. Ce mécanisme se déclenche automatiquement quand l'interface sans fil reçoit beaucoup de trames erronées (ou manuellement par la configuration du paramètre *fragmentation_threshold*), cela permet de diminuer le nombre de retransmissions des trames erronées qui augmente considérablement avec la taille de la trame. La chute du débit ainsi que l'augmentation du taux de perte de paquets à partir de la taille 1500 octets qui est égale au MTU, est du à la fragmentation de paquets IP.
- Durant les tests nous avons montré que tous les flux de paquets UDP de taille inférieure à 500 octets subissent des pertes allant jusqu'à 80-90% (le chapitre précédent) pour des tailles de l'ordre de 100 octet. Ces paquets perdus n'étaient jamais transmis par UDP, et il n'ont même pas atteint l'interface selon le pilote de la carte qui reporte des statistiques sur

les paquets en over run (ceux qui sont générés plus rapidement que le temps de service de la dernière interruption de la carte wireless par le noyau).

5.7. Conclusion

Dans ce chapitre nous avons étudié la procédure du handover 802.11b, ainsi que les améliorations possibles sur les différentes phases et messages qui constituent cette procédure. La principale contribution dans cette partie est la compréhension du comportement du handover pour le trafic vidéo, ainsi que l'étude des performances de la couche mac à l'envoi et la réception du trafic temps réel durant le handover, de plus nous avons détaillé tout les messages et les métriques pertinents dans un handover L2, et nous avons proposé quelque modifications des messages et phases du handover, enfin nous avons réussi a installer et utiliser un réseau Wlan 802.11b(Figure5.1) pour des testes pratiques des débits d'un trafic vidéo durant le handover.

Comme perspectives nous proposons d'introduire les mécanismes de la transmission vidéo présenter dans le chapitre précédant pour minimiser la perte pendant le handover L2.

Chapitre 6

La gestion du handover couche IP (macro mobilité)

6.1. Introduction

De nos jours, différentes technologies d'accès réseaux sans-fil coexistent. À quelques exceptions près, elles permettent à un NM de se déplacer tout en gardant une connectivité avec un réseau cœur (réseau d'interconnexion).

Plusieurs points d'attachement sont répartis dans le réseau, et assurent la connectivité des terminaux mobiles. Une caractéristique primordiale d'une technologie d'accès est la distance maximale (la portée) entre un RA et un NM, à partir duquel la réception des informations devient précaire. La zone délimitée par cette portée est appelée cellule de bas niveau. La Figure 6.1 présente un NM connecté à un RA, relié à un réseau filaire. La portée du RA est symbolisée par la cellule A. Tant que le NM est situé au centre de la cellule A, la qualité de réception de la communication sans-fil est optimale. Plus le NM se rapprochera du bord de la cellule, plus la qualité de réception se dégradera. En dehors de la cellule A, le NM ne peut plus communiquer avec le RA A, et le NM perd sa connectivité.

Le déplacement du NM entraîne la déconnexion du NM avec le RA A, et la connexion avec le point B. Ce phénomène est le *handover*. Les technologies d'accès sans-fil étant incompatibles, ce type de *handover* de bas niveau ne permet pas au NM d'effectuer un mouvement entre cellules hétérogènes (technologies d'accès distinctes). Le délai D de déconnexion au réseau est susceptible de provoquer des pertes de paquets, et aucun paquet ne peut être transmis vers ou depuis le NM.

Le protocole IP (Internet Protocol) est très largement utilisé comme protocole d'interconnexion de réseaux filaires. Il est indépendant des technologies d'accès sous jacentes. Cette propriété est fondamentale pour pouvoir effectuer un mouvement entre cellules hétérogènes. IP apparaît comme le protocole permettant de fédérer les diverses technologies d'accès filaires et sans-fil.

Un RA au réseau IP est un routeur d'accès. Il permet d'accéder à une cellule IP constituée d'une ou plusieurs cellules homogènes de bas niveau (BSS). Le *handover* IP est le processus qui permet à un NM de changer de routeur d'accès. Le délai D de déconnexion IP du NM désigne le temps d'indisponibilité du mobile au niveau IP provoqué par le *handover* (Figure 6.2).

Durant ce délai, aucun paquet IP ne peut être transmis depuis ou vers le NM. Cette indisponibilité débute au moment précis où le mobile perd sa connectivité avec son ancien routeur d'accès, jusqu'au moment où il est capable d'envoyer et recevoir les paquets IP à travers son nouveau routeur d'accès.

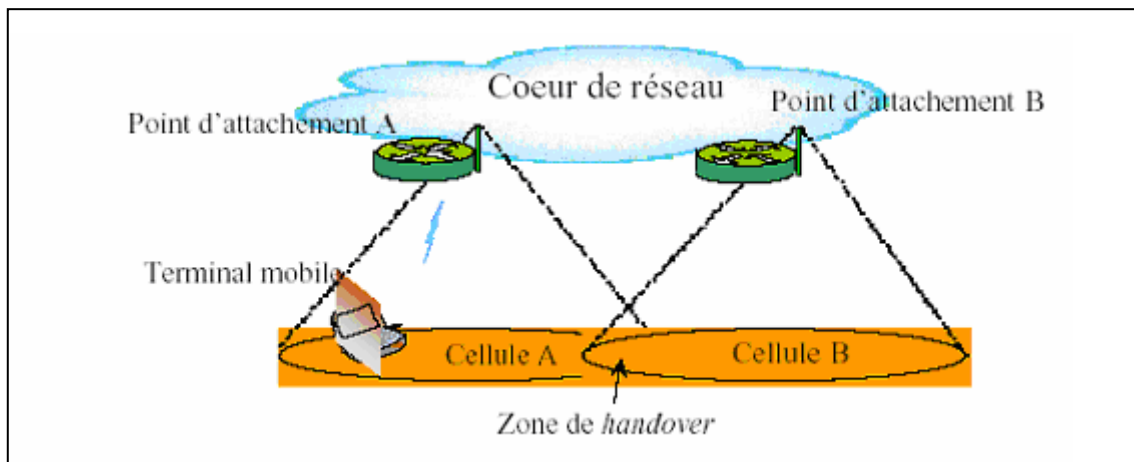


Fig 6.1 : Interaction entre un NM et le cœur de réseau.

La quantité de paquets IP perdus, provoqués par le *handover* IP, est proportionnelle au délai D . Ce délai est composé de deux grandes périodes τ et Δ . La période τ peut précéder ou succéder à la période Δ . Cette période correspond au délai nécessaire de configuration des paramètres IP du mobile. L'enregistrement d'une adresse IP est un exemple typique de cette configuration.

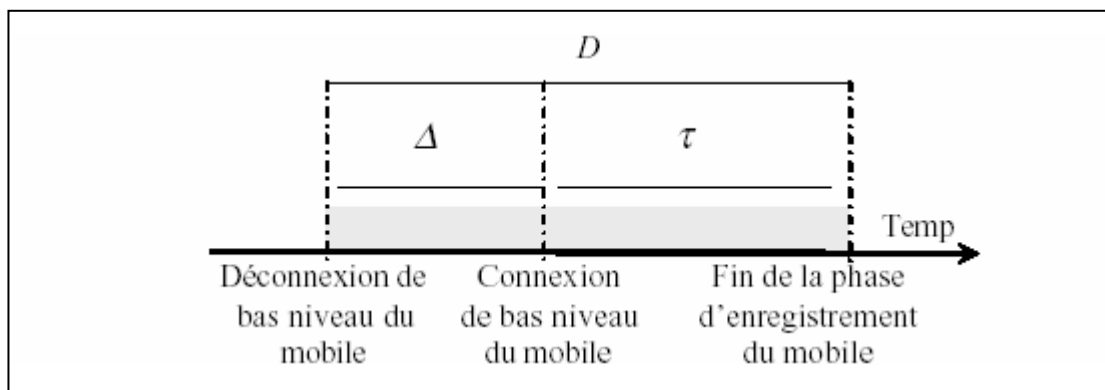


Fig 6.2 : La latence du handover.

Les mécanismes de *handover* IP induits par *Mobile IPv4* [Per02] ou *mobile IPv6* sont similaires. Le délai D , et la quantité de paquets IP perdus, sont relativement importants. De nouvelles techniques comme le *fast-handover* [Koo02][Mon02], et *hierarchical mobile IP* [Cas02], réduisent le délai. De fait, le délai D et la quantité de paquets perdus sont diminués.

Contrairement à ces dernières techniques, le *smooth-handover* [Koo00] [Per99] a pour objectif de réduire le taux de pertes en sauvegardant, pendant la déconnexion, tous les paquets à destination du mobile. Après le *handover* IP, ces paquets sont renvoyés au NM par l'intermédiaire du nouveau routeur d'accès.

Si le NM a la possibilité technique de se connecter simultanément à deux points d'attachement, il lui est possible d'éviter une déconnexion totale du mobile lors d'un *handover*. [Zha01] énonce les bases de cette technique (*soft handover*).

Dans ce chapitre, nous allons analyser et comparer les diverses techniques de *handover* IP, en considérant Le délai de transfert, le nombre de perte, et la gigue (variation du délai de transfert) des paquets seront considérés [Mor02] [Bel04].

6. 2. Les diverses techniques de *handover* IP

Après avoir détaillé les fonctionnalités de Mobile IP (chapitre3) et le handover L2 dans le (chapitre 5), nous allons décrire les diverses techniques de *handover* basées sur le protocole IP. Une analyse de ces techniques, nous permettra de définir le comportement théorique d'un flux UDP entre un NM (nommé mobile) et un terminal fixe (nommé CN). Le mobile dispose d'au moins une interface radio (802.11) lui permettant de se connecter à un ou plusieurs routeurs d'accès disposant du même type d'interface radio. Le correspondant dispose d'une interface filaire (Figure8.1) lui permettant de se connecter au réseau. La communication (flux UDP) entre le correspondant et le mobile est unidirectionnelle et soutenue (il n'y a pas d'interruption de la transmission de la part du correspondant). Le comportement du flux vidéo (UDP) sera caractérisé par le délai de transfert, le nombre de perte et la gigue induits par le *handover* sur un flux de type UDP. Le délai de transfert représente le délai entre l'émission d'un datagramme UDP par le correspondant, et la réception de ce datagramme par le mobile. Le nombre de pertes représente le nombre de datagrammes UDP reçus par le mobile moins le nombre de datagrammes émis par le correspondant durant le *handover*. La gigue représente la variation du délai de transfert (valeur absolue du délai de transfert d'un paquet en réception moins le délai de transfert du paquet précédemment reçu). Afin de simplifier notre analyse des *handover*, nous supposons qu'une cellule IP n'est constituée que d'une cellule de bas niveau, que le RA de bas niveau et le routeur d'accès sont les mêmes entités, que le taux d'erreur bit est négligeable, et que les routeurs ne sont pas congestionnés.

Les Figures 6.1 et 8.1 seront des figures de référence pour l'étude du handover couche IP. Le NM initie le processus de *handover* IP dans le routeur d'accès A, et le routeur B stoppe ce processus. Les techniques de *handover* que nous allons présentées dans ce document fonctionnent avec IPv6.

T	Délai de transfert
N	Nombre de paquets perdus lors du <i>Handover</i>
R	Débit moyen de transmission entre le mobile et le correspondant
D	Délai de déconnexion
Δ	Délai du <i>Handover</i> de bas niveau
Ω	Délai d'enregistrement du mobile
Φ	Délai de configuration de l'adresse temporaire
T_{CA}	Délai de transfert entre le correspondant et le routeur d'accès A
T_{CB}	Délai de transfert entre le correspondant et le routeur d'accès B
T_{CH}	Délai de transfert entre le correspondant et l'agent mère
T_{AB}	Délai de transfert entre les routeurs d'accès A et B
T_{HR}	Délai de transfert entre l'agent mère et le routeur d'accès B
T_{radio}	Délai de transmission sur une interface radio
J_{radio}	Gigue de transmission sur une interface radio
B	Taille du tampon de sauvegarde

Tab 6.1. : Notations générales.

6.2.1. Hierarchical mobile IP

Hierarchical mobile IP [Sol00] est une extension de *mobile IPv4* et *mobile IPv6*. Il peut être associé à toutes les autres techniques de *handover*. Ce protocole est particulièrement conçu pour les mobiles changeant fréquemment leurs points d'attache au réseau dans un périmètre réduit.

Cette approche est basée sur l'introduction d'une architecture hiérarchique d'agents mobiles dans le réseau. Ces agents sont considérés comme des agents mère locaux, capables d'effectuer des opérations d'enregistrement du mobile.

L'architecture hiérarchique du réseau, consiste à le structurer en domaines, chaque domaine ayant son propre agent de mobilité. Ce domaine à son tour est découpé en plusieurs régions ayant chacune son propre agent de mobilité. Ces agents sont hiérarchiquement inférieurs à l'agent mobile du domaine. Une région peut à son tour être découpée en sous-régions. Dès lors, les agents des sous-régions sont hiérarchiquement inférieurs aux agents de la région. Lors de la première connexion du mobile dans un domaine, il enregistre son adresse temporaire une première fois avec son agent mère et à l'ensemble des agents de mobilité hiérarchiquement supérieurs au point d'accès courant. L'agent de mobilité d'une région a une table de liaison permettant d'associer l'adresse temporaire du nœud mobile dans cette région, et l'adresse temporaire de la région de niveau hiérarchique supérieur. Les paquets IP sont véhiculés de région en région par l'intermédiaire de tunnels.

Chaque mouvement du mobile au sein d'une région, entraîne un enregistrement local de l'adresse temporaire. Cet enregistrement est transparent pour les agents hiérarchiquement supérieurs, pour l'agent mère et pour les correspondants. Ceci permet de réduire la charge d'enregistrement des adresses temporaires auprès des diverses entités. Ω_{HMIP} est égal au délai aller et retour entre le mobile et l'agent de mobilité. L'agent de mobilité étant vraisemblablement plus proche que le correspondant, nous obtenons :

$$\Omega_{HMIP} < \Omega_{MIP6}$$

L'intérêt principal d'une telle approche est d'accélérer l'opération d'enregistrement des mobiles en traitant localement les déplacements. Le délai nécessaire pour l'obtention d'une nouvelle adresse temporaire étant identique à *mobile IPv6*, le délai de déconnexion est :

$$D_{HMIP} = \Delta + \Phi_{MIP6} + \Omega_{HMIP}$$

Et le nombre de paquets perdus pendant le *handover* devient :

$$N_{HMIP} = R(\Delta + \Phi_{MIP6} + \Omega_{HMIP})$$

Nous déduisons que $N_{HMIP} < N_{MIP6}$. Cependant, l'introduction des agents de mobilité MA_i et des tunnels successifs augmente le délai de transfert des paquets IP :

$$T_{HMIP} = T_{MAi} + \sum T_{MAi,MAi+1} + T_{MA_n} + T_{radi}$$

avec $T_{MAi,MAi+1}$ le délai de transfert entre les agents de mobilité MA_i et MA_{i+1} , T_{CMAi} le délai entre le correspondant et l'agent MA_i , et $T_{MA_n,B}$ le délai entre l'agent de mobilité MA_n et le routeur B. or,

$$T_{HMIP} = T_{MAi} + \sum T_{MAi,MAi+1} + T_{MA_n}$$

de ce fait $T_{HMIP} > T_{MIP6}$. La gigue induite par le *handover* est :

$$J_{HMIP} = |T_{MA_n,A} - T_{MA_n,B}| + J_{radio}$$

avec $T_{MA_n,A}$ le délai de transfert entre l'agent MA_n et le routeur d'accès A. Les agents de mobilité étant proche des routeurs d'accès et si A et B appartiennent à un même domaine de gestion, il est fortement probable que la gigue de *hierarchical Mobile IP* soit inférieure à celle de *Mobile IPv6*. Si A et B n'appartiennent pas au même domaine de gestion, il n'y a pas de comparaison possible.

Toutefois, si l'on suppose que la gigue J_{radio} est prépondérante, la gigue globale de *hierarchical Mobile IP*, de *mobile IPv4* et de *mobile IPv6* serait similaire.

6.2.2. Smooth-handover

Le *smooth-handover* [Koo00] est une technique censée réduire considérablement le nombre de pertes de paquets IP lors d'un *handover*. Pour réduire ces pertes, elle fait appel à un mécanisme de sauvegarde dans le routeur d'accès A. Le routeur A effectue une copie de tous les paquets qu'il transmet au mobile avant, pendant et après la déconnexion. Bon nombre de paquets perdus en transmission ont été sauvegardés. Dès que le mobile obtient sa nouvelle connexion avec le routeur B et sa nouvelle adresse temporaire, le mobile dévoile au routeur B l'adresse du routeur A.

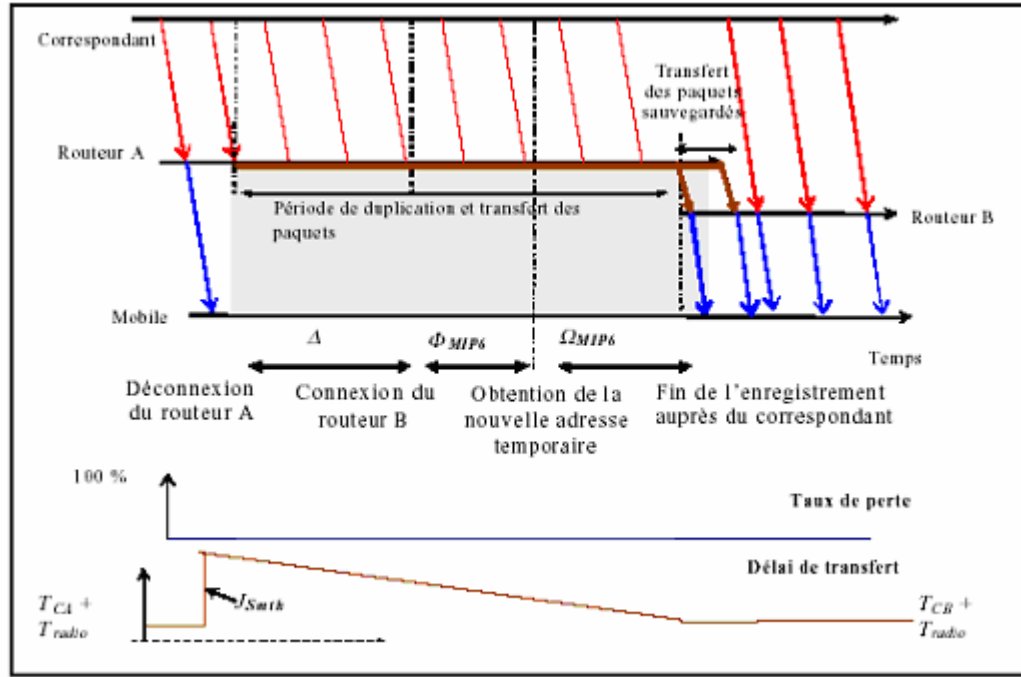


Fig 6.3 : Diagramme temporel pour le smooth handover.

Le routeur B établit un tunnel IP avec le routeur A et lui indique la présence du mobile. A cet instant, le routeur A cesse de sauvegarder et transmettre les paquets sur l'interface sans-fil à destination du mobile. Il transmet au mobile, et à travers le tunnel, tous les paquets précédemment sauvegardés. Le routeur B se chargera de relayer ces paquets vers le mobile. Les paquets dupliqués et sauvegardés lors de la déconnexion sont récupérés localement.

Selon la technique de gestion du tampon, le mobile peut recevoir deux fois le même paquet. Il existe plusieurs techniques de sauvegarde afin d'optimiser les ressources des routeurs d'accès [Per99]. La Figure 6.3 présente le diagramme temporel concernant le *smooth-handover*.

Le délai de déconnexion du mobile lors d'un *handover* est exactement le même que celui d'un *handover* pour *mobile IPv6*. Toutefois, le nombre de paquets perdus est :

$$N_{Smth} = \max(0, R.D_{MIP6} - B)$$

avec B la taille du tampon de sauvegarde. Si B est supérieur à $R.D_{MIP6}$, il n'y a aucune perte due au *handover*. En contre-partie, le mécanisme de sauvegarde et la tunnelisation des paquets du routeur A vers le routeur B augmentent notablement le délai de transfert et la gigue. Si nous négligeons le délai de transmission de A vers B, le délai de transfert des paquets sauvegardés dans le routeur A est :

$$T_{Smth} = T_{CA} + x.D_{MIP6} + T_{AB} + T_{radio}$$

avec T_{AB} le délai de transfert entre les routeurs d'accès A et B, et $0 \leq x \leq 1$. Le délai de transfert des autres paquets est identique à *mobile IPv6*. La gigue induite par le *handover* est :

$$J_{Smth} = D_{MIP6} + T_{AB} + J_{radio}$$

Or, $T_{AB} \geq 0$ et $D_{MIP6} > 0$. De ce fait, nous obtenons :

$$J_{Smith} > J_{MIP6}$$

Nous pouvons conclure que le *smooth-handover*, peut éviter ou diminuer la perte de paquets, mais introduit un délai de transfert additionnel et une gigue importante.

Son utilisation pour des transmissions de données ayant des contraintes temporelles parait impossible.

6.2.3. Fast-handover :

Le *fast-handover* [Koo 02] réduit le délai de déconnexion du *handover*. Une des lacunes du *handover* dans *mobile IPv6* est la détection tardive du mouvement du mobile, et l'enregistrement tardive de l'adresse temporaire. Dans le cadre de *mobile IPv6*, le mobile détecte son mouvement après l'obtention de la nouvelle connexion physique avec le routeur B. A partir de ce moment, le mobile obtient sa nouvelle adresse temporaire, et l'enregistre. Afin d'anticiper le mouvement du mobile, le *fast-handover* suppose une interaction entre la couche IP et la couche de bas niveau. Cette anticipation du *handover* de bas niveau permet au mobile, et aux routeurs d'accès de préparer la nouvelle liaison avec le routeur B avant même de perdre la liaison avec le routeur A.

Avant l'exécution du *handover* de bas niveau, le mobile identifie le routeur d'accès B. Il communique cette information au routeur A. Le routeur A établit un tunnel avec le routeur B. Le mobile peut obtenir sa nouvelle adresse temporaire, et entamer le processus d'enregistrement de cette adresse à travers le routeur A. Ainsi, dès que le mobile est lié au routeur B, il est immédiatement en mesure d'envoyer et de recevoir des paquets IP. Juste après la déconnexion, tous les paquets reçus par le routeur A, et destinés au mobile seront transmis au routeur B par l'intermédiaire du tunnel. Le routeur B les transmettra au mobile. Il est à noter que la charge due à la signalisation IP du *fast-handover* est importante. La Figure 6.4 présente le diagramme temporel pour le *fast-handover*.

Ce mécanisme réduit considérablement la durée de déconnexion du mobile lors du *handover* IP. Si le temps de déconnexion de bas niveau est inférieur au délai de traversée du tunnel, tous les paquets redirigés par le routeur A arriveront au mobile après sa liaison avec le routeur B. A cette condition, aucun paquet n'est perdu. Dans le cas contraire, un certain nombre de paquets arriveront au routeur B avant l'établissement de la nouvelle connexion et seront perdus. Dans tous les cas, le nombre de paquet perdu N_{fast} a la propriété suivante :

$$N_{Fast} > R.\Delta$$

Contrairement au *smooth-handover* les paquets ne sont pas sauvegardés lors du *handover*. Le délai de transfert est donc largement diminué. Pendant le *handover*, ce délai est :

$$T_{Fast} = T_{CA} + T_{AB} + T_{radio}$$

Ensuite, ce délai devient équivalent à celui de *mobile IPv6* :

$$T_{Fast} = T_{CB} + T_{radio}$$

La gigue est :

$$J_{Fast} = T_{AB} + J_{radio}$$

Or, nous avons $T_{AB} \geq 0$ ce qui implique :

$$J_{Fast} \geq J_{MIPv6} \text{ et } J_{Smth} > J_{Fast}$$

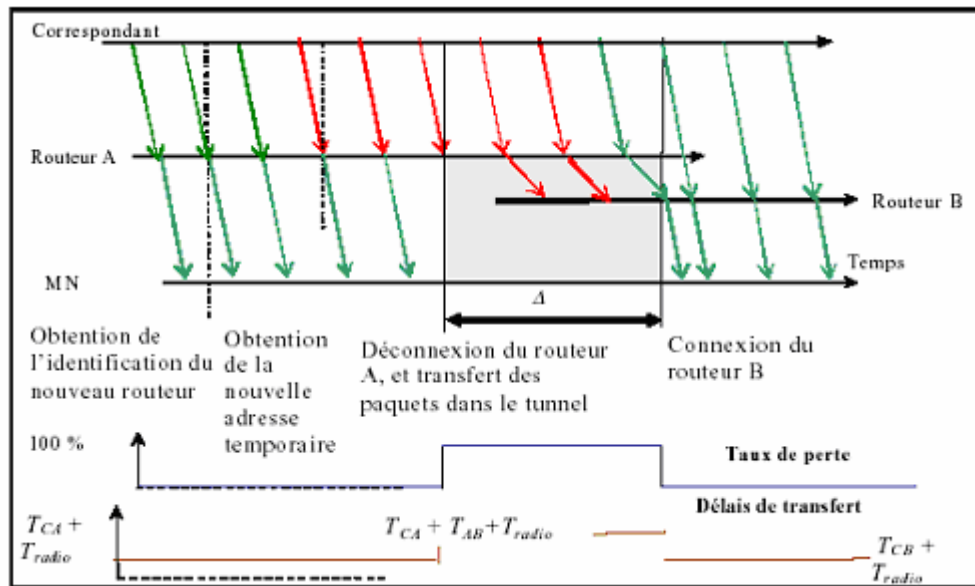


Fig 6.4 : Diagramme temporel pour le fast handover.

6.3. Analyse générale des diverses techniques de handover

Dans ce qui suit, Nous proposons de résumer notre analyse des diverses techniques de handover. *Mobile IPv6* est une amélioration de *mobile IPv4* pour IPv6. Il profite des avantages d'IPv6 comme la configuration automatique des adresses. Dès lors, le temps de déconnexion du mobile est plus faible, et le nombre de paquets perdus est réduit proportionnellement. Il évite le routage triangulaire, ce qui diminue le délai de transfert des paquets. La gigue est identique à *mobile IPv4*.

La restructuration du réseau, proposé par *hierarchical mobile IP*, a le mérite d'améliorer les performances du handover lors d'un déplacement local du mobile.

La charge des messages de signalisation nécessaire aux enregistrements des mobiles est réduite. Dans le cas général, le délai de la déconnexion est plus faible que pour *mobile IPv6*, et donc le nombre de perte est lui aussi réduit. De plus, la gigue est plus faible que pour *mobile IPv6*. Toutefois, l'utilisation successive des tunnels entre chaque agent de mobilité augmente sensiblement les délais de transfert des paquets. Les mécanismes de stockage introduits dans le *smooth-handover* réduisent considérablement le risque de perte de paquets au détriment du délai de transfert et de la gigue. Il n'est donc pas adapté à la transmission de données ayant de fortes contraintes temps réel. Le *fast-handover* réduit le délai de déconnexion au minimum. Les pertes sont donc plus faibles, voir inexistantes si le délai de reconnexion physique est prompt.

Les délais de transfert sont plus longs que pour *mobile IPv6*, mais plus courts que pour le *smooth-handover*.

Le *fasthandover*, comme toutes les autres techniques décrites précédemment, sont incapables de contrôler des mobiles se déplaçant trop rapidement. Toutes les techniques de *handover* IP ont une influence sur le flux UDP. Il n'y a aucune isolation parfaite du flux.

En se basant sur les trois techniques présentées dans la section 6.2 beaucoup de solutions sont proposées [Far04] [Cle04] [Fen04] [Cam02] [Bel03] [Viv04] et qui décrivent un ***Fast handover avec bicasting***, un schéma de smoothing basé sur les déclencheurs de niveau 2 dans un réseau sans fil 802.11 une solution HMIPv6, un protocole de gestion de mobilité hiérarchique basée sur les informations de la couche liaison et un algorithme de fast handover adapté pour une solution hiérarchique de MIPv6

Nous pouvons conclure que toutes ces extensions de mobile IPv6 suivent les trois axes illustrés par la figure suivante :

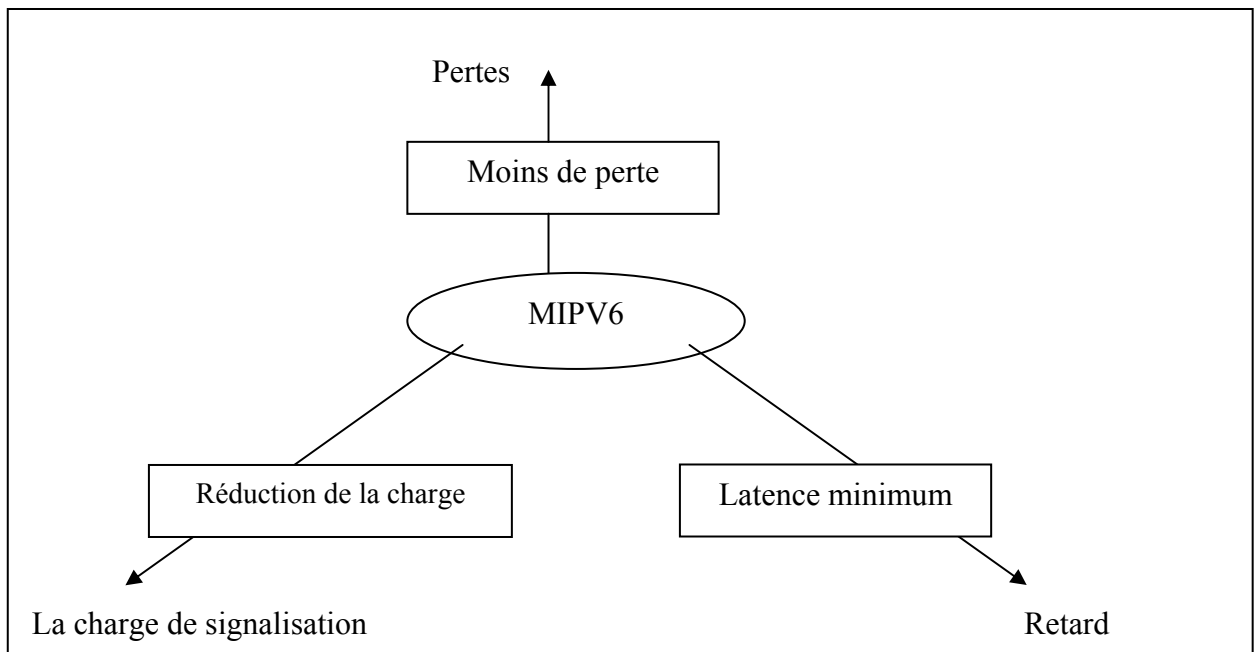


Figure 6.5- *Stratégie du Handover*

Pour cette raison nous allons concevoir dans le chapitre 7 un modèle générique de gestion du handover qui couvre les besoins des trois axes de la figure ci-dessus.

6.4. Conclusion

Dans ce chapitre nous avons présenté les différentes définitions et hypothèses, ainsi que les extensions et techniques pour la gestion de la mobilité dans le but de motiver leurs utilisations dans notre modèle de gestion du handover (chapitre 7). Ce chapitre ne couvre pas tous les travaux réalisés dans le domaine de la gestion de mobilité, mais nous permet d'illustrer les principaux axes qui définissent la base de notre solution. En plus nous avons cité les travaux de critère bien précis et qui tend vers l'amélioration du handover dans les trois axes: moins de perte, moins de retard, moins de charge de signalisation. Ces critères et techniques sont par la suite utilisés pour définir les hypothèses et la conception du modèle générique (présenté dans le chapitre suivant).

Chapitre 7

Modèle générique pour le Contrôle du handover IP

7.1. Introduction

Nous présentons dans cette partie un modèle générique flexible pour le contrôle de la mobilité dans les réseaux sans fil Wlan 802.11. Ce modèle permet l'amélioration des performances du handover mobile IPV6, de plus il offre une satisfaction aux exigences de la **QoS** des différentes applications temps réel telles que : la vidéo streaming, vidéo conférence. En particuliers, nous visons à maximiser le niveau de satisfaction de l'utilisateur pendant l'usage de ressources des réseaux mobiles.

7.2. Les besoins du modèle générique pour le contrôle de mobilité

Mobile IP peut surmonter les différents challenges et problèmes. Les plus importants de ces défis sont: les performances du handover, les performances de transport¹, la sécurité [Tea03].

Plusieurs mécanismes pour la gestion du handover couche IP sont proposés tels que le Smooth et le Fast handover, pour le support d'un trafic multimédia durant le handover. Cependant, une solution ne peut pas être équitable pour tout trafic et pour toutes les exigences de la QoS.

Il y a un besoin de penser à des mécanismes génériques et plus flexibles qui peuvent contrôler les variations et les caractéristiques du comportement des applications temps réel durant le handover de telle sorte qu'une améliorations ou changement d'un composant n'implique pas forcément les changements des autres composants du modèle.

7.3. Latence du handover pour des applications temps réel

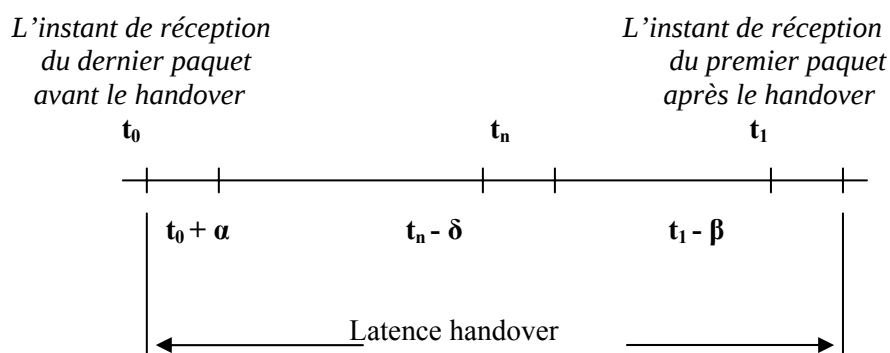


Fig 7.1 : Définition de la latence du handover mobile IP.

¹ - les performances de transports sont étudiées dans la contribution 1

Le **Fast handover** est une condition nécessaire pour le trafic temps réel [Tea03][Mal03] plus spécialement pour le handover très fréquent dans le WLAN [Fen04].

[Mon02] définit « low latency handover » : le temps minimum dans lequel le mobile est incapable de recevoir ni envoyer les paquets.

En général la latence du handover est définie comme la différence en temps ($t_1 - t_0$) (Fig7.1), les autres ($t_0 + \alpha$, $t_1 - \beta$, t_n , δ) décrivent les différentes phases du handover global, ces phases sont (handover couche liaison, détection de mouvement, découverte des routeurs d'accès, l'enregistrement...)

$t_0 + \alpha$: l'instant quand la connectivité via l'ancien lien sera perdue

$$0 < \alpha < \max_interval_time$$

$t_1 - \beta$: nouvel enregistrement terminé, et le NM est prêt à recevoir les données via le nouveau lien

$$0 < \beta < \max_interval_time$$

t_n : l'instant quand le handover couche lien est réalisé après que le contact avec le nouveau lien est établi à $(t_n - \delta)$ (δ la latence du handover L2)

δ : est en général plus petit que la latence du handover L3.

t_n pourrait venir avant $t_0 + \alpha$ si les zones de couverture radios de deux points d'accès (ancien/nouveau) sont recouvertes.

$t_1 - \beta - t_n$: le temps nécessaire pour compléter la détection de mouvement et la procédure de l'enregistrement Mobile IP.

Les réseaux Wlan impliquent des changements fréquents de PA et de RA et des fois à de trop petits intervalles de temps, ce qui cause l'intolérance du handover à la longue découverte des RA, les délais d'enregistrement et les charges de signalisation significatifs.

Plusieurs approches et mécanismes sont développées et qui aident à réduire la latence handover ainsi que la réduction de la charge globale du réseau et l'optimisation de découverte des RAs.

Nous pouvons dériver les directives suivantes (ces directives (entités) définiront la base de notre modèle) pour le développement du modèle générique pour la gestion du handover.

- ❖ Solution basée sur mobile IPv6
- ❖ La gestion hiérarchique de la mobilité (HMIPv6) est désirable pour réduire la charge de signalisation.
- ❖ le Fast handover (réduction de la latence du handover).
- ❖ Le Smooth handover (Minimiser la perte).
- ❖ Le couplage entre Layer3, Layer2 est important pour avoir de meilleures performances

7.4. Modèle générique du handover IP

7.4.1. Introduction

Le modèle générique du handover IP contient plusieurs composants tels que : le contrôle du handover (CH), l'anticipation du handover (AH), l'optimisation du handover (OH) , et la détection du mouvement (DD). Nous allons implémenter ces composants dans le RA, le NM et le HA (voir fig7.2).

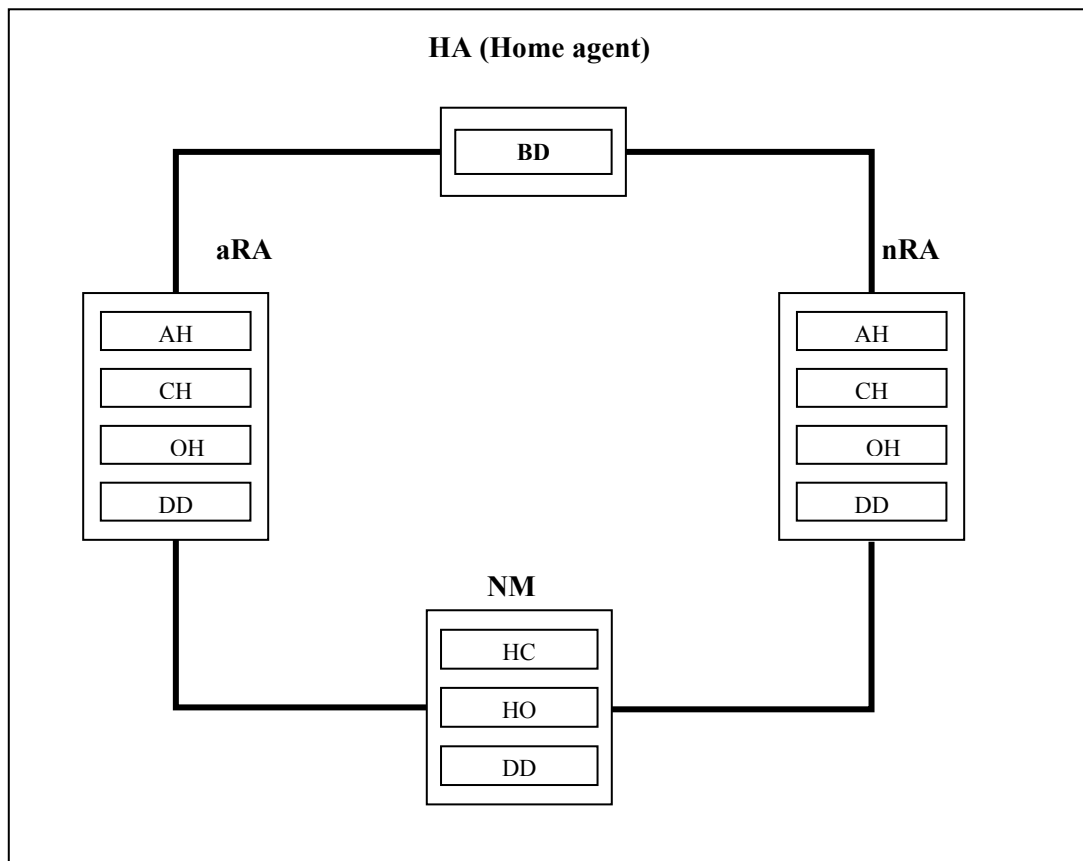


Fig 7.2 : Modèle générique du handover IP.

aRA : Ancien RA.

nRA : Nouveau RA

CH : Contrôle du handover

DD : Détection du mouvement et sélection et découverte des RA

OH : Optimisation du handover

AH : Anticipation du handover

7.4.2 Les étapes de contrôle du handover dans le modèle générique

À la base des différents composants du modèle générique, nous définissons les étapes de base pour la gestion du handover. Quand la décision du handover est faite, nous devons exécuter les étapes suivantes :

a- Sélectionner le meilleur RA (routeur d'accès) : Nous définissons le composant DD (Découverte et sélection des RAs) pour réaliser la sélection du meilleur RA parmi le groupe des RAs candidats. Nous citons dans la section 7.4.4 les différents protocoles de découverte et de sélection des RAs.

b - Décider du meilleur temps pour l'exécution du handover : Cette étape nécessite l'intervention des composants AH (Anticipation du Handover) et CH (Contrôle du Handover), puisque la décision du meilleur temps dépend des informations fournies par la couche MAC et le mode d'enregistrement (post – pré enregistrement).

c - Choix des meilleurs mécanismes disponibles pour le handover : Dans cette étape nous introduisons les techniques d'optimisation du handover (OH) : le buffering, le Bicating ...

7.4.3. Les buts de la conception du modèle générique

Le modèle est inspiré des recherches et des solutions et les approches de base du handover proposées précédemment du mobile IP [Abd01][Fen04][Ens04][Tea03] et il a pour but les points suivants

- Fournir un service flexible
- *Seamless handover* (moins de perte, moins de retard)
- Dégradation acceptable de la qualité de vidéo lors d'une transmission sur un Wlan
- Compatibilité avec MIPv6 et les autres extensions
- Extensibilité pour couvrir les futures propositions du handover
- Modèle générique qui satisfait les besoins des applications multimédia et qui peut commander tous les scénarios IP qui existent
- Le modèle doit être simple pour le déploiement (Simple nœud, simple composant) et inclut des mécanismes simples, suffisants pour faire un *seamless handover* entre deux RA

7.4.4. Définition des différents composants du modèle générique

Dans cette partie nous proposons une description des différents composants du modèle générique qui aident à améliorer le traitement du handover couche IP. Ces composants sont basés sur les recherches présentées dans le chapitre 6. L'idée de base est la conception du modèle générique pour le Handover L3 qui pourra améliorer les performances globales du handover dans les trois axes définis dans le chapitre précédent Figure (6.5).

A. Le contrôle du handover (CH)

Le CH est un composant du modèle générique du handover IP, c'est le plus important composant parmi les autres. Toutes les solutions de la gestion du handover doivent disposer d'un ensemble de fonctionnalités qui correspondent au composant de contrôle du handover. Le rôle majeur du CH est d'initier le handover, l'enregistrement (mise à jours des adresses) et installer les nouveaux chemins de routage.

Initiation du handover

C'est le déclenchement de la procédure du handover (lancer les déclencheurs du handover), le meilleur instant pour commencer le handover c'est quand les sessions de communication actives

subissent la moindre altération ; ie la puissance de réception diminue en dessous du seuil S' (chapitre 5). Dans ces cas, l'instant d'initier le handover peut être en retard à cause des stratégies mises en œuvre, l'entité d'initiation du handover peut être le NM ou le RA, dépend de la nature du handover ; ie : une initiation réseau ou une initiation du mobile.

La mise à jours de localisations « Location update »

C'est pour la mise à jour des emplacements courants du NM dans une base de données cela définit la procédure d'enregistrement. L'enregistrement comporte la création, modification et la suppression des informations de localisation.

Chemin de routage

C'est pour installer les nouveaux chemins de routage, selon les informations de mises à jour à propos des emplacements courants des NMs et les points d'attachement au réseau qui peuvent être obtenus de la base de données. Ce processus est dynamique dans le sens que le nouveau chemin est établi sur la réception des paquets destinés au NM.

La figure ci-dessous montre les scénarios possibles pour les processus des handover L3 et L2. Comme illustré par les cas a, b, c, d. Nous définissons le pre-enregistrement dans les cas (b) (d) et le post-enregistrement (a)(c)

Pre- Enregistrement

Le handover L3 (établissement de nouvelle connexion) est fait via l'ancien lien avant l'établissement et l'installation des nouvelles connexions L2 (handover L2) Cas (b) (d). Il est possible que la connexion L3 soit établie sans un lien L2 direct, en plus si un autre chemin L2 indirect il peut servir la connexion L3 (b). La *Pre-Enregistrement* est basé sur un tel mécanisme.

Post-Enregistrement

La nouvelle connexion L3 (handover L3) est faite après que le nouveau lien L2 est fait. Cas (a)(c). Les deux méthodes d'enregistrement sont définies indépendamment du nombre des liens courants.

- (a) **Connexion de lien Multiple** la nouvelle connexion L3 (handover L3) est établie après l'établissement du lien L2 (handover L2), c'est le cas le plus général. Quand le NM peut entendre de plusieurs PAs en même temps il peut faire un ***soft handover*** (*make before break*).
- (b) **Connexion de lien Multiple** la nouvelle connexion L3 (handover L3) est faite via l'ancien lien avant l'installation du nouveau lien L2, l'établissement du lien L2 se fait via l'ancien lien L2.

- (c) **Connexion de lien simple** : casser l'ancien lien avant ou juste après la connexion du nouveau lien, la connexion L3 est faite après l'exécution du handover L2 (connexion L2) c'est le **hard handover** (break-before-make). Il existe le **semi soft handover** (entre le hard et le soft handover). La condition d'un tel handover est que le chevauchement des zones de couvertures des PAs est imperceptible.
- (d) **Connexion lien simple** la nouvelle connexion L3 est faite via l'ancien lien avant l'installation du nouveau lien L2. Si l'ancien L2 est cassé avant de terminer la connexion L3, cette dernière continue via le nouveau lien L2.

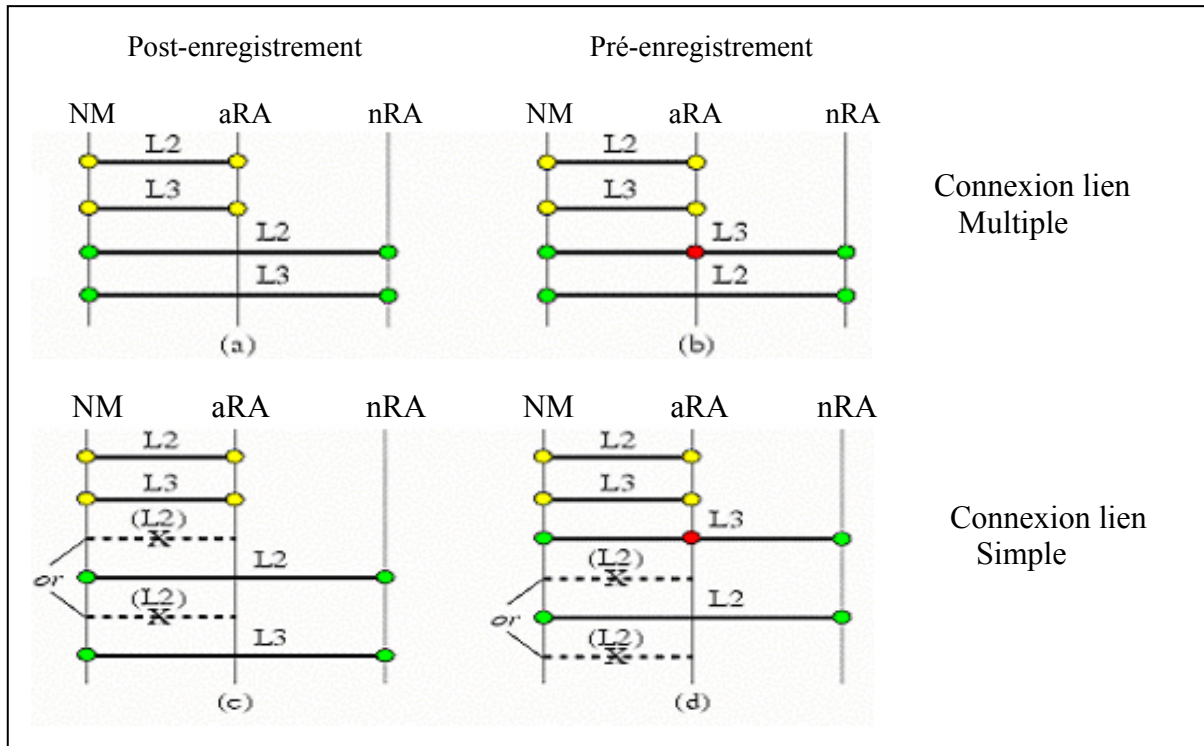


Fig 7.3 : Différents scénarios de l'enregistrement.

B. L'anticipation du handover (AH)

Le composant de l'anticipation du handover est un composant qui a pour but de réduire la latence du handover, en cours de communication, entre des points d'accès sans fil. Lorsqu'un mobile entre dans un nouveau sous-réseau IPv6, les mécanismes définis par Mobile IPv6 [Per01] lui permettent de poursuivre ses communications, mais les délais observés peuvent être importants pour des applications de type temps réel.

Les protocoles de Fast Handover [Mont01] ou de Transfert de Contexte [Per01], qui sont des extensions à Mobile IPv6 sont des protocoles d'anticipations du handover. Ils proposent d'anticiper des opérations de la couche réseau par des déclencheurs de niveau 2. Ces déclencheurs de niveau 2 sont des messages génériques à tout protocole de la couche liaison qui avertissent d'un changement dans l'état de liaison du mobile avec les points d'accès.

1- Transfert de contexte

Le *transfert de contexte* [Per01] a pour objectif de réduire les effets d'un «handover» sur les applications temps réel des NMs. Le principe est de mettre en place le plus rapidement possible le même niveau de service dont bénéficiait une NM dans son sous-réseau précédent. Ce protocole propose donc un échange de messages entre RA pour les différentes sessions du NM. A chaque flux d'un NM est donc associé un contexte différent et identifiable. Le but majeur du Transfert de contexte est de fournir des services au mobile durant le handover [Ens04]. Celui-ci peut-être :

- La compression d'en-tête, souvent utilisée dans la propagation radio puisque le coût d'accès au médium est plus coûteux et les débits offerts sont en général inférieurs à ce que l'on observe en filaire.
- La mise en tampon des paquets, par exemple pour la gestion des flux TCP.
- Le niveau de qualité de service (IntServ , DiffServ [Gra01]).

De la même manière que pour le *Fast Handover* (défini dans le chapitre précédant), les deux RA établissent un tunnel entre eux leur permettant de s'échanger des informations sécurisées sur les NM. La mise en œuvre de ce tunnel peut être faite avant même la connexion du NM au nouveau point d'attache en utilisant les déclencheurs de niveau 2.

C. Détection de mouvement et Découvertes des RAs (DD) :

C.1. Contributions de détection de mouvement et la découverte des RAs à la latence du handover total :

Deux composants majeurs contribuent dans la latence totale : la découverte d'agent et la détection de mouvement. On définit ces deux composants comme suit :

Td : le temps nécessaire pour la détection de mouvement et la découverte de RAs dépend toujours des mécanismes adoptés dans les deux composants.

Treg : le temps nécessaire pour terminer la procédure d'enregistrement impliquant le NM, RA et HA. Selon [Tea03], on peut définir la latence du handover comme suit

$$TH = Td + Treg \quad (1)$$

S' il existe une longue distance ou retard important entre le RA et HA, le **Treg** devient plus grand que le **Td**. La détection du mouvement est un mécanisme indépendant des autres composants du modèle générique, c'est un composant en général indépendant du cycle de gestion de la mobilité, la justification de l'inclure dans le modèle vient du fait que la détection du mouvement a un impact significatif sur les performances générales du handover.

C.2. Les mécanisme de découverte des Agents et la détection du mouvement

Il existe un nombre de mécanisme suggérés pour mobile IP tel que LCS (Lazy Cell Switching), ESC (Eager Cell Switching) [Tea03]. Ces deux mécanismes dépendent de la couche IP et le Fast Active Cell Switching, CARD (Candidate Access Router Discovery) [ens04], qui sont des mécanismes hybrides (couche liaison et couche IP).

ESC (Eager Cell Switching)

Le NM initie le handover immédiatement après la découverte d'un nouveau RA, ie quand le NM reçoit l' *advertisement agent* du nouveau RA. Alors la latence maximal de découverte d'agent et détection de mouvement dans cette méthode est égale au maximal de l'intervalle d'envoi du *advertisement agent*.

$$T_{ECS_Max} = T_{ADV_Max_Send_Interval}$$

$$T_{ECS_Mean} = T_{ADV_Max_Send_Interval}/2$$

LCS (Lazy Cell Switching)

Autorise le handover au nouveau RA seulement après la perte de trois paquets consécutifs de l'ancien RA.

$$T_{LCS_Max} = 4 \times T_{ADV_Max_Send_Interval}$$

$$T_{LCS_Mean} = (7 \times T_{ADV_Max_Send_Interval})/2$$

Fast Active Cell Switching (ACS)

C'est une solution basée sur ECS et assistée par les déclencheurs niveau2 (*L2 trigger information*). ACS sollicite un message *Advertisment Request* du nouveau RA à la réception des déclencheurs L2 (basé sur les valeurs des seuils de puissance du signal). Cette technique aide à réduire au maximum la latence de découverte des RAs et la détection du mouvement au délai RTT (Round Trip Time) qu'est le temps d'aller retour entre le NM et le RA.

Candidate Access Router Discovery Protocole (CARD)

Le NM peut choisir le RA dans un futur proche, après la formation du groupe de candidats.

Il existe deux approches :

- 1- Le NM peut sélectionner le **RA en se** basant sur le besoin précédemment détecté pour le handover.
- 2- Les [RAs] peuvent informer pro activement le NM via un chemin direct.

D. Optimisation du Handover (OH)

C'est un composant optionnel du modèle générique, son but est d'améliorer les performances globales du handover.

D.1. Buffering

Le buffering est généralement considéré comme la technique du « *smoothing handover* », ce terme veut dire réduire la perte des paquets due au handover, il est généralement indépendant des mécanismes du Fast handover. Le concept de base de l'utilisation du buffering est d'achever le *smoothing handover* et de bufferiser les paquets destinés au NM ou nouveau RA, quand le handover est en progression (la livraison des paquets est impossible) et les envoyer à la fin du processus du handover.

Après que le NM se déplace au nouveau « point d'attachement » et avant qu'il exécute les nouveaux enregistrements (Mobility Binding) avec le HA ou (MAP) via le nouveau RA, les paquets bufferisés seront envoyés au NM en utilisant deux agents : l'agent de buffering et l'agent d'envoi.

Le nœud demandeur du buffering peut être le NM, nouveau RA, ancien RA pendant que le buffering réduit énormément la perte des paquets.

Durant le handover, le buffering a aussi quelques effets négatifs tels que la *duplication de paquets*, augmentation des retards des paquets dûs au délai du buffering.

La duplication peut arriver quand les paquets bufferisés sont déjà délivrés au NM. La méthode la plus utilisée pour éliminer les duplications est l'utilisation de l'identificateur des paquets (dans l'entête IP : un champ pour l'identification du paquet).

Le bénéfice majeur de l'utilisation du buffering est l'élimination des pertes des paquets durant le handover, il peut y avoir encore quelques paquets qui peuvent être perdus durant le début et la fin du buffering si les deux(handover, buffering) ne sont pas synchronisés; la solution possible pour ce problème est de commencer le buffering avant de commencer le handover.[Wan99] suggèrent de commencer le buffering juste après la réussite de l'enregistrement.

[Tea03] propose trois cas pour commencer le buffering :

1. Commencer le buffering au nouveau RA quand la RRQ *Registration Request* est faite
2. Avec l'aide de prédiction du prochain handover, commencer le buffering juste avant d'exécuter le prochain handover
3. Commencer le buffering à l'ancien RA, quand le RRQ est exécuté au nRA, (bénéfique dans le cas où le temps de l'enregistrement est très grand dû à la longueur du chemin entre le RA, HA

D.2. Bicasting

Le Bicasting c'est une autre technique du smoothing handover, le Bicasting peut être catégorisé selon les entités du Bicasting. Premièrement et le cas le plus connu, c'est quand le HA ou le MAP envoie des paquets au NM par les deux routeurs l'ancien RA et le nouveau RA. Ce Type de Bicasting peut être achevé facilement sans n'importe quels mécanismes additionnels dans le mobile IP de base qui supporte une mise à jour simultanée via le message RRQ (registration request). Un deuxième type est quand le RA courant agit comme entité de Bicasting, le RA courant envoie les paquets au NM et aux nouveaux RA s. Un exemple de ce type de Bicasting est expliqué dans le *Bidirectionnel Edge Tunnel* [Mon02]. L'avantage de l'utilisation du Bicasting est de réduire la perte des paquets durant le handover aussi bien que minimiser les délais. Cependant le Bicasting peut causer des effets indésirables quand le hard handover est exigé. Il y a une possibilité d'augmenter la charge du réseau qui va augmenter le délai du trafic puisque il envoie des paquets identiques à des destinations multiples.

7.5. Conclusion

Dans ce chapitre nous avons présenté un modèle générique pour la gestion de mobilité. Ce modèle propose d'inclure les différents composants qui permettent d'améliorer les performances en terme de perte, de retard et charge de signalisation. En plus le modèle défini est un modèle flexible et extensible. Par conséquent nous pouvons définir sur la base de ce modèle des solutions optimales pour la gestion du handover (chapitre8).

CHAPITRE 8

GHMM (Generic Hierarchical Mobility Management)

8.1. Introduction

Grâce à la popularité d'Internet, le protocole IP s'impose comme incontournable pour les futurs des télécommunications. Par ailleurs, la miniaturisation des équipements informatiques et le développement des réseaux et services de transmission de données sans fil contribuent à développer les services de mobilité.

Dans un réseau IP, chaque nœud connecté est localisé par rapport à son point d'attachement avec son adresse IP. Lorsqu'un tel nœud est doté de fonctions de mobilité et donc qu'il se déplace en changeant de point d'attachement, il doit impérativement, soit changer d'adresse IP, ce qui l'oblige à interrompre les connexions en cours ; soit informer tous les routeurs du réseau Internet que son adresse IP localise un nouveau point d'attachement. La première solution ne fournit qu'une solution de nomadisme permettant à un terminal de se raccorder successivement au réseau IP à partir de points d'accès distincts sans fournir de continuité de service entre ces points. Elle peut être mise en œuvre à travers les protocoles d'auto configuration permettant notamment l'allocation d'adresse dynamique comme le protocole DHCP très couramment utilisé. Par contre, la seconde solution est irréaliste à l'échelle de l'Internet. L'IETF a donc standardisé Mobile IP [Per01] pour intégrer le support de services de mobilité dans les réseaux IP étendus. Mobile IP permet un routage transparent des paquets IP destinés aux nœuds mobiles et permet d'assurer par conséquent la continuité des communications en cours. Mobile IP est utilisé pour gérer la "macro-mobilité". Les délais de diffusion de la nouvelle localisation ainsi que le trafic de contrôle généré pour joindre le réseau principal sont pénalisants. En effet, ils entraînent des interruptions de services de plusieurs secondes, des pertes de paquets de données correspondant aux connexions en cours et des échanges de signalisations qui chargent inutilement le cœur de réseau.

Ainsi, plusieurs protocoles de macro-mobilité ont été développés [chapitre précédent] pour remédier à ces problèmes. Néanmoins, les temps d'interruption de service dus à la procédure de handover restent importants du fait que ces protocoles n'incluent pas de mécanisme de prédiction de mouvement. De plus, le réseau n'étant pas impliqué dans la procédure de décision de la cellule cible lors d'un handover. Il est possible de mettre en place des mécanismes globaux de gestion de la mobilité et réseaux nécessaires basés sur le modèle générique proposé précédemment (chapitre7), afin de lever ces limites. Nous

proposons dans ce chapitre un nouveau protocole de gestion de la macro-mobilité IP basé sur le modèle générique et une architecture hiérarchique appelé **GHMM** (**Generic Hierarchical Mobility Management**). Il s'appuie sur un contrôle combiné de la mobilité (NM, agent de mobilité) effectué dans le réseau et inclut la phase de prise de décision de l'exécution du handover et les techniques de smoothing et d'anticipation du Handover. Cette solution a pour but d'optimiser les trois axes définis dans la figure 6.5 ; ie : moins de perte, moins de retard et une charge de signalisation réduite pour permettre une meilleure transmission des paquets vidéo sur un WLAN 802.11. Pour cela nous allons proposer dans le cadre de cette solution l'optimisation des composants du modèle générique (chapitre7). En plus, nous proposons que le nœud mobile soit l'initiateur du handover dès que la puissance de signal atteint un seuil S'

8.2. Macro-mobilité IP

La macro mobilité, mobilité restreinte limitée à un même site, est caractérisée par des déplacements locaux, fréquents et rapides. Lorsque Mobile IP est utilisé pour la gestion de ce type de mouvement, le nœud mobile doit émettre à chaque déplacement, même local, des messages de mise à jour d'associations vers son agent mère et ses nœuds correspondants qui peuvent être très éloignés de son réseau visité. Le cœur du réseau IP est alors chargé par les messages de mise à jour et la prise en compte de la nouvelle localisation peut être considérablement longue. Les sessions en cours subissent alors des interruptions de service de plusieurs secondes ainsi que d'importantes pertes de paquets ce qui influe d'une façon directe sur la qualité de la vidéo à la réception. Un protocole de macro mobilité est nécessaire pour pallier à ces inconvénients. Ce qui fait l'objet du protocole GHMM.

8.3. Protocoles avec une architecture basée sur une hiérarchie d'agents de mobilité

Ces protocoles étendent l'idée de Mobile IP en utilisant une hiérarchie d'agents de mobilité. Un nœud mobile s'enregistre auprès de son agent local du plus bas niveau de la hiérarchie : puis l'agent local s'enregistre auprès de l'agent supérieur dans la hiérarchie et les enregistrements remontent la hiérarchie jusqu'à atteindre l'agent mère. Lorsque le nœud mobile se déplace, il n'émet qu'une requête d'enregistrement localisée dans un domaine en fonction de la nature de son déplacement. Les paquets qui lui sont destinés suivent alors la hiérarchie des agents de mobilité et sont livrés d'un agent à un autre via un tunnel. Les protocoles les plus connus sont l'enregistrement Régional Mobile IP [Bel04] et la Gestion hiérarchique de la mobilité MIPv6 [Fen04].

8.4. Limites des protocoles de macro-mobilité

Ces protocoles de macro-mobilité offrent des solutions performantes pour le support de la

mobilité localisée. Ils présentent l'avantage de résoudre les principales limites de Mobile IP sans complexifier le mécanisme de gestion de la mobilité. Cependant, des améliorations sont encore possibles surtout en ce qui concerne la phase de décision d'exécution d'un handover, la procédure de détection du mouvement et le temps d'exécution du handover. Un nœud mobile subit simplement le handover décidé au niveau de l'agent de mobilité. Par exemple, dans un réseau local sans fil, à chaque fois que la qualité de la liaison radio avec son point d'accès courant se dégrade, le mobile change de point d'accès. Au niveau IP, le nœud mobile doit attendre la réception d'une annonce de routeur « router advertisement » pour détecter son mouvement et entamer la phase de mise à jour de sa localisation. Il y a donc un intervalle de temps non négligeable entre le déplacement physique du nœud mobile (la couche 1 et 2 du modèle OSI) et la **détection** du mouvement IP (couche 3). Ce qui accroît le temps d'exécution du handover.

Dans la partie **suivante**, nous proposons un nouveau protocole GHMM qui est une évolution du protocole HMIPv6 basé sur le modèle générique présenté précédemment. Nous supposons dans le protocole GHMM que le NM prend la charge d'initier le handover. GHMM propose de résoudre les problèmes décrits ci-dessus en incluant deux phases :

- 1- La phase pré-handover ; c'est la phase de prise de décision de l'exécution du handover qui comporte les techniques de découvertes des RA et la détection de mouvement et les mécanismes d'anticipation.
- 2- La phase du handover : incluent les techniques d'optimisation du handover et d'enregistrement

8.5. GHMM (Generic Hierarchical Management Mobility)

8.5.1. Architecture considérée

GHMM comporte quatre entités fonctionnelles (figure 8.1) : le routeur d'accès RA (RA+ le point d'accès PA), le MAP (un gestionnaire de mobilité local), la base de données BD et le nœud mobile NM. RA est le routeur d'accès du nœud mobile au réseau IPv6. Chaque routeur d'accès est associé à un gestionnaire de mobilité (MAP). PA assure la liaison radio avec le nœud mobile. Chaque point d'accès est attaché à un unique routeur d'accès. Le MAP intercepte et livre les paquets destinés aux nœuds mobiles du domaine (fonctionnalité identique à celle du serveur de mobilité dans MIPv6). Il intervient, entre autres, dans la prise de décision de l'exécution d'un handover ainsi que dans le choix du point d'accès cible élu du handover. La BD héberge des informations sur les points d'accès du domaine et des paramètres utiles à la gestion des handovers.

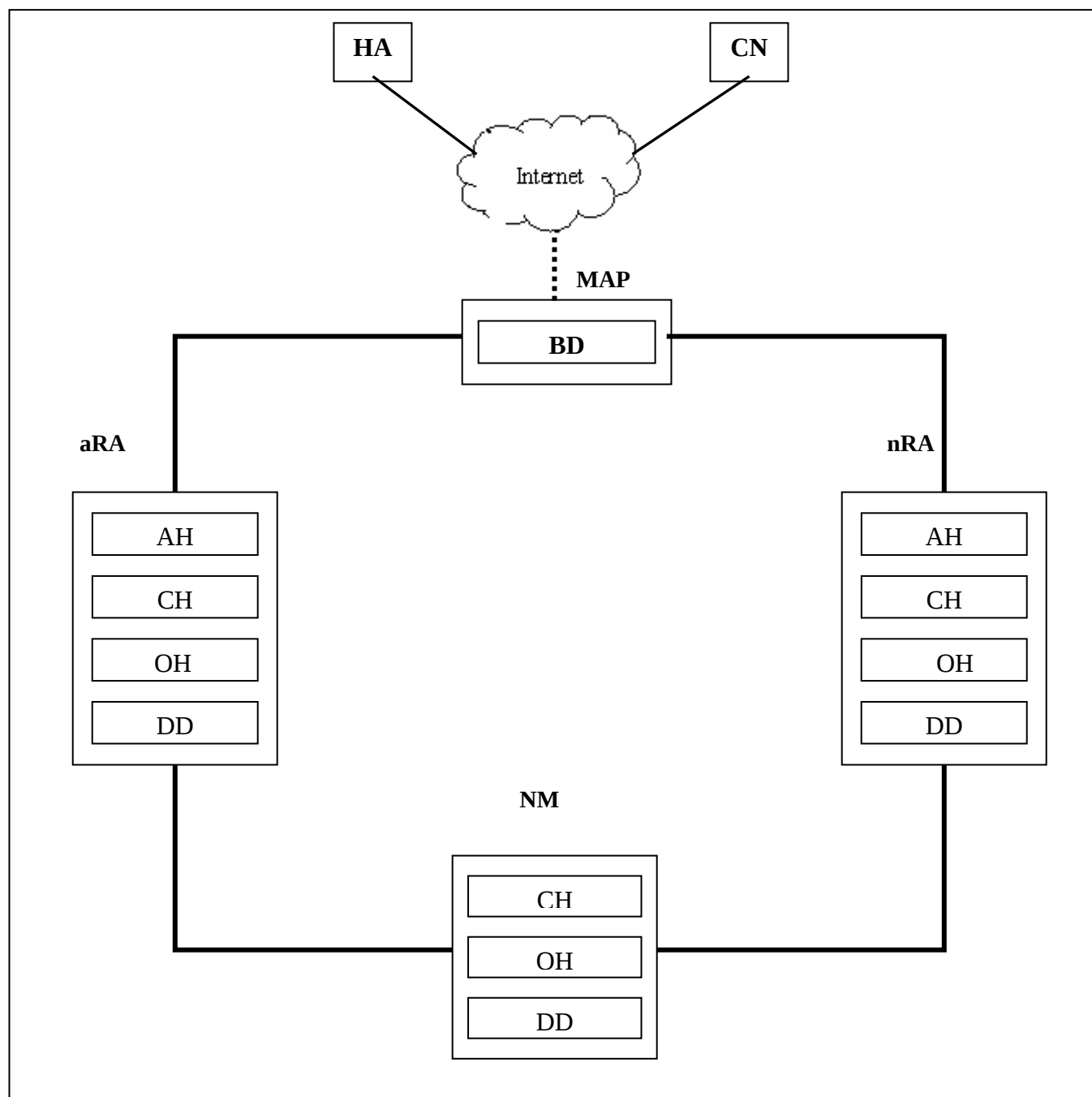


Fig. 8.1 : Architecture proposée

MAP : Mobility Anchor Point

ARA : Ancien RA.

AH : Anticipation du handover

DD : Détection du mouvement et sélection et découverte des RA

CN : Nœud Correspondant

OH : Optimisation du handover

CH : Contrôle du handover

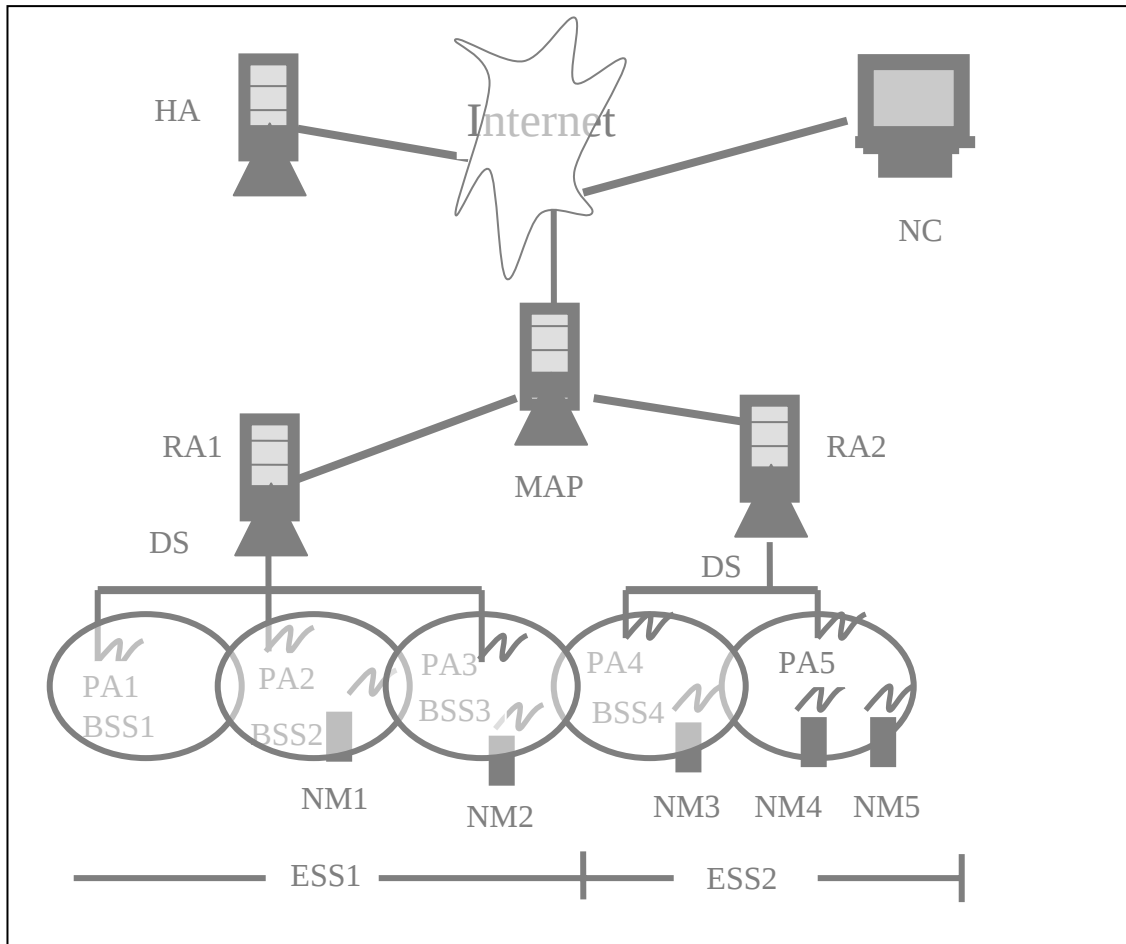


Fig. 8.2 : Architecture de HMIPv6.

BSS : base station set

MAP : Mobility Anchor Point

ESS : extended service set

8.5.2. Description des techniques et des composants du GHMM

8.5.2.1 HMIPv6

HMIPv6 est basé sur Mobile IPv6 qui introduit une nouvelle entité qui est le MAP (Mobility Anchor Point), en plus de petites extensions concernant l'opération du noeud mobile (communication MAP-NM). Les opérations correspondant à l'agent mère ne seront pas affectées. Le MAP est un routeur local dans un réseau visité par le NM. Le MAP est utilisé par le NM comme un HA local. Un ou plusieurs MAPs peuvent exister dans un réseau visité. L'introduction du MAP résout les problèmes de la charge de signalisation des communications mobiles en réduisant le nombre de messages qui circulent dans le réseau. Les deux principes de base d'une communication HMIPv6 sont :

- Le NM envoie les Mises à jour au MAP local plutôt qu'aux HA et CNs;
- Seulement un message de Mise à jour (binding update) est transmis par le NM au MAP qui se

charge d'informer l'agent mère et les correspondants (C'est indépendant du nombre de CNs avec qui le NM communique). Donc un MAP fonctionne comme Agent Mère local.

L'opération HMIPv6

Nous montrons une architecture de HMIPv6 dans la figure 8.2. Avant que nous étudions l'opération de HMIPv6, deux concepts, devraient être introduits en premier lieu l'adresse ; RCOA *Régional Care Of Address* qui est obtenue par le NM dans le réseau visité.

Un RCoA est une adresse sur le sous réseau du MAP. Le deuxième concept c'est (LCoA) *local Care Of Address*. Le LCoA est configuré sur l'interface d'un NM basée sur les *Annonces du routeurs* envoyés par son routeur. Dans le support de la Mobilité IPv6 [Fen04] c'est connu simplement sous le nom de *care - of - Address*. Cependant, LCoA est utilisé pour le distinguer du RCoA.

Un NM qui entre dans un domaine du MAP recevra des *Annonces du routeur* contenant les renseignements sur le MAP local. Par exemple, dans la Figure 8.2, NM2 reçoit des annonces du RA1, et l'annonce contient les renseignements du MAP. Les renseignements incluent les adresses RCoA. Cette procédure est appelée la découverte des MAPs. La phase de la découverte informera aussi le NM de la distance (MAP du NM). Le NM peut lier son emplacement courant (CoA) avec une adresse sur le sous-réseau du MAP (RCoA) et enregistre le RCoA avec son HA et CNs. Maintenant, le MAP agit comme un HA local. Les paquets qui sont envoyés à ou par le NM sont envoyés par le MAP, les autres mécanismes de la communication sont conservés normalement comme dans IP Mobile. Si le NM change de LCoA dans le domaine du MAP, par exemple (dans la Figure8.2) NM2 se déplace entre BSS3, BSS2 et BSS1, il a seulement besoin d'enregistrer la nouvelle adresse avec le MAP. Donc, NM2 a besoin d'une seule mise à jour (nécessaire) avec son HA et ses CNs. En conséquence, aussi longtemps que NM2 est en mouvement dans le domaine du MAP, ce mouvement est transparent au HA et aux CNs avec qui il communique. Les limites d'un domaine du MAP sont définies par les Routeurs (RAs) qui envoient des renseignements du MAP aux NMs attachés. La Figure8.2, montre le domaine du MAP : RA1 et RA2 qui envoient des *annonces* dans ESS1 et ESS2. Le NM continuera à détecter s'il est encore dans le même domaine du MAP. Chaque fois qu'un NM détecte un changement dans le MAP (*les routers advertisement*) cela veut dire que le NM s'est déplacé à un nouveau domaine du MAP, le NM doit agir sur le changement en envoyant des Mises à jour à son HA et NCs. En conséquence, si on prend NM2 comme exemple, quand NM2 entre dans le BSS 3, il a besoin d'exécuter une Mise à jour Obligatoire une seule fois avec son HA et NC aussi long temps qu'il se déplace entre ESS1 et ESS2. Mais quand le NM2 se déplace de BSS3 (appartient à ESS1) à BSS4 (appartient à ESS2), il n'a pas besoin d'enregistrer avec HA ou NC. Par conséquent, sur ce plan IPv6 Mobile hiérarchique accomplit le handover plus rapidement en réduisant les mises à jour entre un NM et son HA et NCs, en plus ça permet de réduire la charge de signalisation.

8.5.2.2. Protocole d'anticipation

Les extensions du protocole Mobile IPv6 [Mon02] sont proposées à l'IETF afin de réduire le temps de coupure occasionné par un changement de sous-réseau suite à un déplacement d'une station mobile (NM). Nous optons dans notre solution GHMM pour une anticipation basée sur l'état de liaison de niveau 2 pour commencer le «handover» de la couche 3 avant que le «handover» de la couche 2 ne soit fini (fast handover).

A. Fast Handover

Le Fast Handover propose d'optimiser le déplacement des NM par l'exécution d'un *Handover Anticipé*. Dans ce *Handover Anticipé*, le NM ou le routeur d'accès (RA) auquel est rattaché le NM, reçoit un déclencheur de niveau 2 lui indiquant que le NM est sur le point de faire un «handover» (étapes 1 et 2 dans la Figure 8.3). Ce déclencheur doit contenir des informations permettant au RA courant du NM d'identifier le nouveau sous-réseau (préfixe, masque, adresse IPv6 du routeur d'accès). Le routeur d'accès courant envoie alors à la fois une adresse IPv6 topologiquement correcte pour le nouveau sous-réseau au NM (étape 3) et une requête au RA cible pour qu'il valide cette adresse (étape 4). Le nouveau RA doit alors vérifier que l'adresse est unique dans son sous-réseau. Ensuite le routeur d'accès cible envoie le résultat de la validation au RA courant (étape 5). Si l'adresse est valide, le RA courant du NM lui transmet l'autorisation pour utiliser la nouvelle adresse IPv6 (étape 6). Ainsi, lorsque le NM établit la connexion avec le nouveau point d'attachement, il peut directement utiliser sa nouvelle adresse IPv6 temporaire comme adresse source dans ses paquets sortant et envoyer une demande de mise à jour vers son agent mère et tous ses correspondants. Afin de minimiser la perte de paquets, l'ancien routeur d'accès pourra aussi faire suivre tous les paquets qu'il reçoit pour ce NM au nouveau RA.

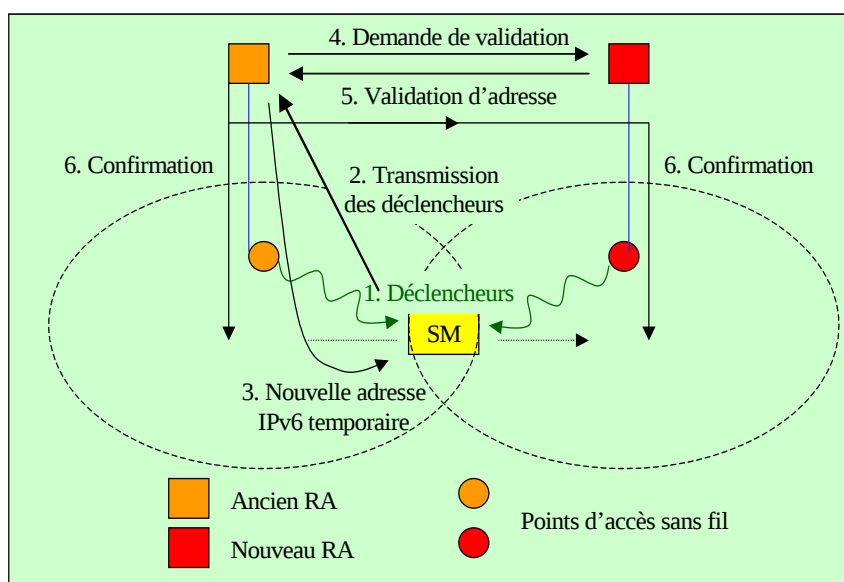


Fig. 8.3 : Procédure d'anticipation

B. Déclencheur de niveau 2

Les déclencheurs de niveau 2 [Mon02] sont l'ensemble des informations nécessaires permettant d'anticiper le déplacement des NM. Un déclencheur de niveau 2 est une abstraction d'un événement du protocole de la couche 2 utilisé par la couche supérieure (IPv6) pour anticiper les opérations de «handover». On parle d'abstraction car ces déclencheurs doivent être génériques, c'est-à-dire indépendant de toute technologie radio. Ainsi, pour chaque technologie, il s'agit d'établir une correspondance entre un événement spécifique du protocole de la couche liaison et le déclencheur attendu par la couche 3.

Un déclencheur de niveau 2 contient divers types d'information : l'événement qui l'a provoqué, l'entité IP destinataire de ce déclencheur. Nous considérons les événements suivants :

- Début de «handover» de niveau 2 : le NM commence à chercher un nouveau point d'accès avec un meilleur signal
- Fin de «handover» de niveau 2 : le NM vient d'établir la connexion avec un nouveau point d'accès.
- Le type de «handover» : le nouveau point d'accès utilise ou non la même technologie que l'ancien point d'accès du NM.
- Changement dans l'état du lien avec le NM : perte de connexion, réduction de l'intensité du signal...etc.

L'entité IP réceptrice est un paramètre important pour les protocoles de niveau 3, puisque c'est elle qui initiera le protocole. On parle de déclencheur source lorsque le RA courant collecte les informations, et de déclencheur cible si c'est le nouveau RA qui est le destinataire des déclencheurs. Dans notre cas c'est le NM qui utilise ces déclencheurs pour initialiser le «handover» de la couche 3.

B.1. Handover L2

Le mécanisme de changement de point d'accès nécessite la découverte d'un point d'accès cible, la synchronisation avec celui-ci, l'envoi d'une authentification et enfin l'établissement de l'association (chapitre5).

Quand un NM détecte que son signal avec son point d'accès courant devient trop faible, il cherche un autre point d'accès qui lui offrirait un meilleur signal. Il ne peut plus envoyer de paquets de données. Il envoie alors un *Probe Request* à une adresse de broadcast de niveau 2. Le point d'accès cible lui répond alors avec un *Probe Response* contenant des informations de synchronisation et divers paramètres. Ensuite le NM et le point d'accès continuent à s'échanger des *Probe Request* et *Probe Response* jusqu'à ce qu'ils soient synchronisés.

Après synchronisation, le NM envoie un message d'authentification qui est acquitté par le point d'accès. Puis le NM demande une association en envoyant un *Association Request* auquel le point d'accès répond avec un *Association Reply*.

B.2. Handover L 3

Le «handover» de niveau 3 tel qu'il est décrit dans Mobile IPv6 [Mon02] comprend la détection du nouveau préfixe, la configuration d'adresse et le temps d'aller-retour pour joindre les correspondants. Dans notre solution, le NM avertit uniquement le MAP (son agent mère locale). Ceci correspond au cas où le NM utilise le MAP pour toutes ses communications. Si on considère des correspondants distants, le temps d'aller-retour pour les joindre devra encore être ajouté.

L'utilisation de l'anticipation peut donc s'avérer fort intéressante puisque dans certains cas le «handover» de niveau 3 pourra finir avant le «handover» de niveau 2 (cas optimal).

Nous présentons dans le paragraphe suivant la correspondance entre les messages utilisés par le protocole MAC de la norme IEEE 802.11b et les déclencheurs génériques.

B.3. Relation entre les déclencheurs de niveau 2 et les messages de contrôle de la norme IEEE802.11b

Nous avons présenté les déclencheurs de niveau 2 précédemment qui sont nécessaires à la couche réseau pour anticiper le «handover» de niveau 3. Dans cette partie nous allons établir leur correspondance avec les messages du protocole MAC de la norme IEEE 802.11b.

Pour traiter ces messages, les RA doivent connaître les RA qui leur sont voisins, par rapport aux zones de couverture de leurs points d'accès respectifs (protocoles de découverte de routeur d'accès). Nous avons présenté dans le chapitre précédent quelque protocole de découverte dynamique qui sont en cours de normalisation [Fen04] [Tea03]. Cette découverte devra inclure un moyen d'identifier l'appartenance d'un point d'accès à un sous-réseau particulier. Les points d'accès pourront être différenciés entre eux par leur adresse MAC. Dans la suite, nous intégrons le protocole de découverte de routeur CARD (Candidate Access Router Discovery Protocol). Ce protocole permet aux RA non seulement la connaissance des RA voisins, mais qu'ils possèdent également une liste des points d'accès rattachés à ces RA (sauvegardée dans une base de données dans le MAP).

8.5.2.3. Détection du mouvement sélection et découverte des RAs (DD)

Le handover IP mobile se produit comme une conséquence du handover couche liaison entre les différentes ESS (changement d'adresse IP). L'événement d'un handover couche liaison ne peut pas être communiqué à IP Mobile, d'où la nécessité d'établir des techniques pour déterminer ces renseignements. Pour ceci, les Agents de mobilité de Mobiles IP (HA, AR) diffusent périodiquement des annonces pour

indiquer leur existence. Donc un NM peut découvrir ces agents en recevant leurs annonces. De plus, un NM détermine son emplacement par l'évaluation de ces annonces. S'il a reçu une nouvelle annonce du RA " non découverte " l'agent est perçu comme une indication de mouvement dans un nouveau réseau. De la même façon, la perte de contact avec un agent " découvert "(ancien RA) est perçu comme une indication de mouvement hors d'un réseau. Après avoir évalué les annonces reçues, le NM veut exécuter une technique de Sélection du meilleur RA en utilisant les informations de voisinage stockées dans le MAP (BD) en utilisant le protocole CARD [Fen04] pour permettre un bon fast handover.

Les deux actions, la découverte du mouvement et la sélection des RAs, consomment un temps considérable [Fen04]. Par exemple, la vitesse de découverte du mouvement dépend directement de l'intervalle des annonces. Un plus petit intervalle des *annonces* mène à une optimisation de la découverte du mouvement et une sélection rapide du nouveau RA, pour cela, l'utilisation du protocole CARD nous paraît adéquate pour une bonne détection et sélection des nouveaux RAS.

8.5.2.4. Optimisation du handover (OH)

Nous avons présenté dans le chapitre précédent un schéma de buffering, qui a pour but la réduction de la perte des paquets. Nous proposons dans le paragraphe suivant un schéma optimal du buffering. Nous essayons de résoudre les problèmes de saturation des buffers et la duplication de paquets lors de l'exécution du buffering durant un handover.

A. Le buffering Optimal

Pour maximiser le bénéfice de l'utilisation du buffering il faut concevoir des mécanismes de buffering optimal qui incluent le NM, ancien RA, le nouveau RA et le HA ou MAP (l'utilisation de HMIPv6). La figure suivante montre un schéma optimal du buffering.

Pour transporter le contrôle du buffering entre les entités impliquées, nous introduisons un nouveau format de message. Le message **RCB** (*Requête de Contrôle du Buffer*) ; ce message est utilisé pour initier le stockage des paquets dans le buffer. Le principe de base de ce schéma optimal de buffering est le suivant :

Quand la demande du buffering est envoyée par le NM à l'ancien RA via le message **RCB**, l'envoi direct de paquets continue aussi longtemps que la communication directe (NM- ancien RA) est possible, ce qui peut causer la perte de paquets en cas de saturation du buffer (due à sa taille limitée).

Un deuxième message de contrôle du buffering est introduit c'est Le message RCB (envoi). ce message est envoyé juste après l'envoi de RRQ (registration request), cela assure un délai minimum des paquets

bufférisés délivrés au NM. Pour éviter la duplication des paquets due aux envois simultanées de paquets bufférisés et les paquets envoyés via l'agent d'envoi, le message RCB (envoi) inclut un paramètre (l'identificateur du dernier paquet) qui permet au NM de connaître le dernier paquet reçu avant de commencer le buffering (indique le début de la phase du buffering).

Quand l'envoi des paquets bufférisés de l'ancien RA au NM via le nRA est en cours d'exécution, le buffering à l'ancien RA continue, cela aide à éviter la perte des paquets durant la phase d'enregistrement (pendant le temps entre l'émission du RRQ par le NM et l'acceptation du NM par HA (MAP)) ; spécialement quand le temps RRT(temps d'aller retour) entre le NM et le HA est assez grand pour subir la perte des paquets pendant cette période.

Nous définissons aussi le mécanisme qui traite la saturation des buffers. Un facteur de conception est très important ; c'est la capacité du buffer optimal en terme de nombre de paquets bufférisés ; si le nombre de paquets alloué est grand ceci peut causer une consommation excessive de ressources partagées, dans le cas opposé, le nombre de paquets est petit et peut causer une perte excessive des paquets bufférisés tôt et saturé par le paquets reçus plus tard.

Le nombre optimal de paquets bufférisés est le nombre qui évite à la fois la consommation excessive des ressources et la saturation du buffer. Il sera calculé à la base des débits d'envoi et les seuils de déclencher le handover ainsi que l'instant de synchronisation (Handover-smoothing)

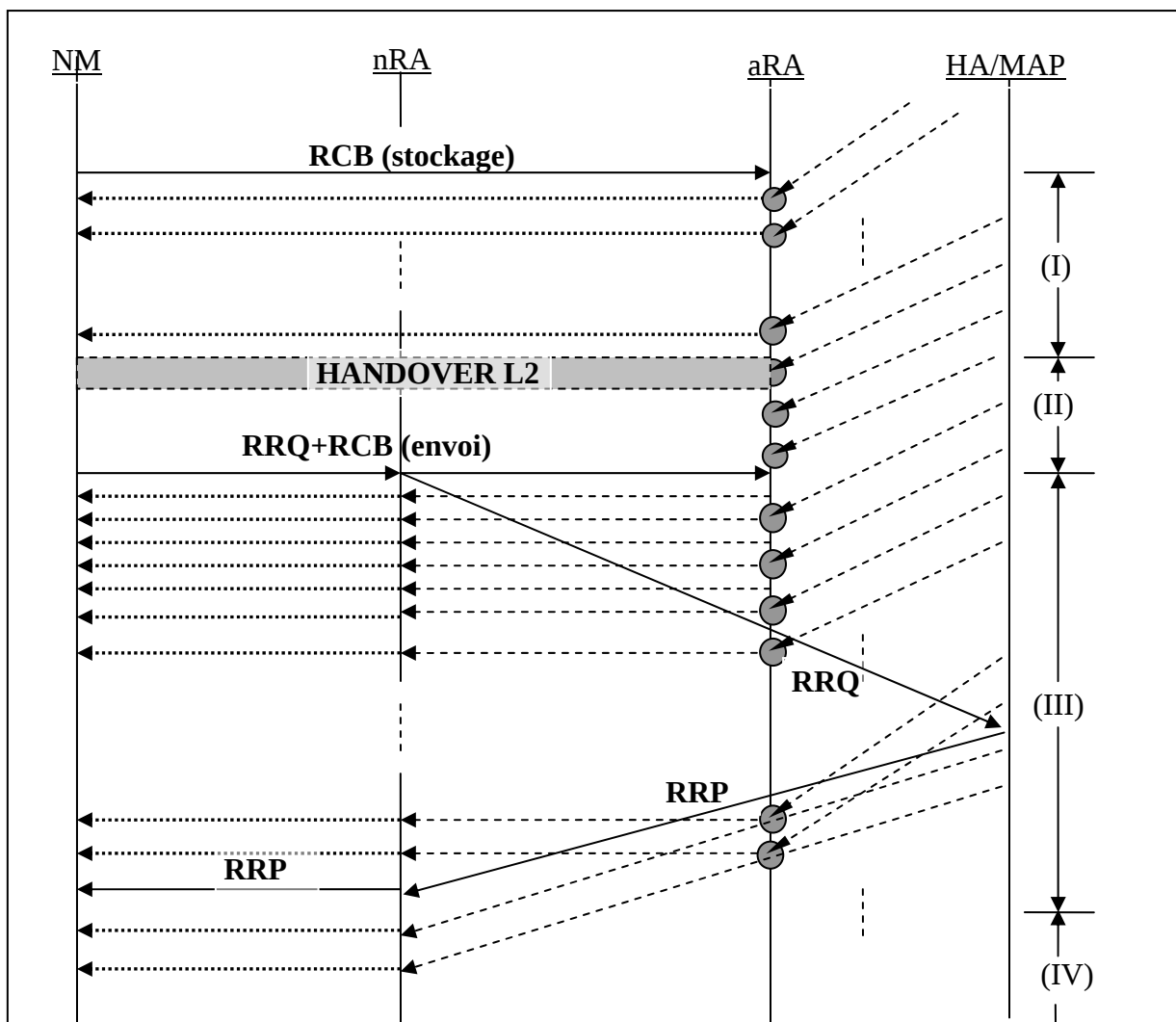


Fig 8.4 : Schéma de buffering optimal

→ : Message de contrôle - - ▸ : Paquets de données (via tunnel)

● : Buffering ▸ : Paquets de données (direct)

I : buffering et envoi direct au NM via l'ancienRA

II : buffering dans l'ancien RA

III : envoi indirect de données bufferer de l'ancien RA au NM via le nouveau RA

IV : envoi direct via le nouveau RA

B. Bicasting optimal

Le scénario de Bicasting que nous proposons est une technique améliorée du Bi casting pour éviter les effets secondaires. Pour cela nous devons combiner la technique du bicasting avec des mécanismes complémentaires proposés déjà dans le Mobile IP. Les techniques les plus importantes à combiner sont le Fast handover et le pre-enregistrement (présentés précédemment)

La table illustre les fonctions de base du pré- enregistrement

- I. Le NM envoie **PrStSol (proxy router solicitation)** et le **aRA** le reçoit.
- II. **aRA** envoie **PrRTAdv(Proxy Router Annonce)** et le NM le reçoit (les informations de voisinage des RA sont toujours disponibles).
- III. NM envoie **RRQ (avec CoA = nRA_addr)** au **nRA** via **aRA**.
- IV. **nRA** échange **RRQ (registration request)** et **RRP registration reply** avec le **MAP**.
- V. **nRA** envoie **RRP** au NM directement ou via **aRA**.

Tab 8.1 Procédure du pré- enregistrement

La technique du pré- enregistrement aide à améliorer le Bicasting : en utilisant le pré enregistrement le NM est capable de recevoir le trafic dès qu'il établit le contact avec le nouveau RA, ce qui aide à minimiser la perte, le retard et la duplication des paquets.

Le NM demande l'adresse du nouveau RA en envoyant un message **PrStSol (Proxy Router Sollicitation)**, à la réception de ce message par l'ancien RA, ce dernier message acquitte le NM en lui envoyant le message **PrRTAdv(Proxy Router Annonce)**. Ce message contient les informations de voisinage et le NM le reçoit. A la réception de ces informations, le NM décide d'initier le handover au nouveau RA et de se déplacer hors de la portée de son ancien RA et il envoie un message de pré-enregistrement à son nouveau RA via l'ancien. En plus le NM aura la possibilité de lancer la synchronisation entre les techniques de smoothing et la procédure de pré-enregistrement.

Les messages de pré- enregistrements (RRQ : registration request/ RRP : registration response) sont échangés via un tunnel entre les deux Routeurs. La figure suivante montre ce concept.

La 1ere phase (b): les paquets envoyés par le nouveau RA sont toujours abandonnés, pour cela il est suggéré de combiner le pré- enregistrement avec le Bicasting voir (b) ça aide à réduire l'envoi simultané des paquets par l'annulation de l' enregistrement avec l'ancien RA, dès que le NM se déplace à la zone de couverture du nouveau RA. L'enregistrement est initié sur l'expiration du temps d'enregistrement, et le Bicasting dans la De- enregistrement est nécessaire pour réduire le temps du De- enregistrement, (c),(d)

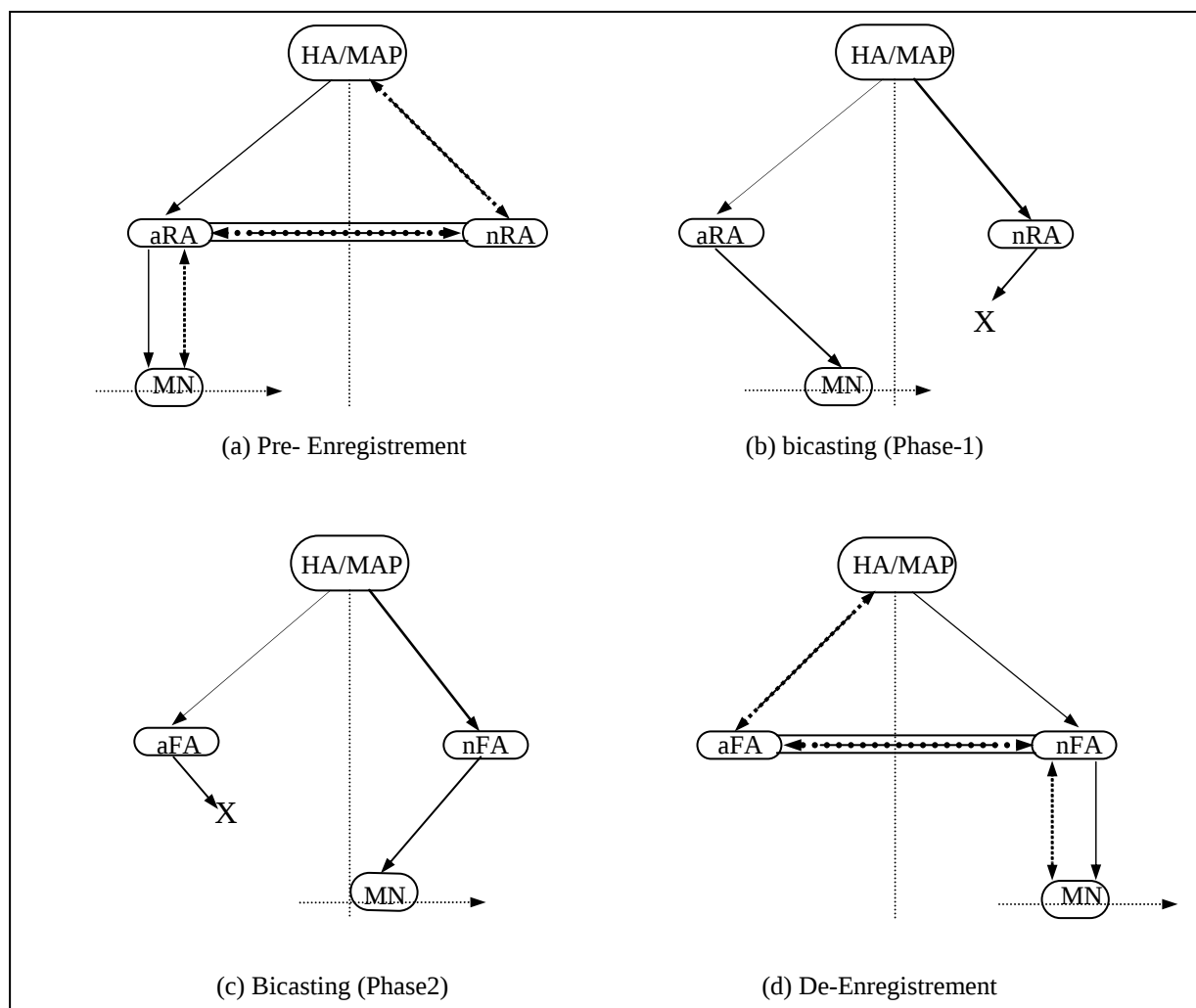


Fig 8.5 : *Le Bicating avec pré- enregistrement*

8.5.3. Comportement du GHMM

Dans un **domaine** GHMM, chaque point d'accès est associé à une interface d'un routeur d'accès et à un gestionnaire de mobilité et possède un ensemble de points **d'accès** voisins. **La liste** des points **d'accès** voisins d'un point d'accès est maintenue dans une **base** de données hébergée par son gestionnaire de mobilité (MAP). Cette base de données maintient aussi des **informations** utiles à la prise de décision de **l'exécution** des handovers. Un point d'accès est identifié avec un identifiant unique, le préfixe réseau de l'interface du routeur d'accès associée et la longueur de ce préfixe.

Le MAP annonce son support de GHMM en diffusant périodiquement à ses routeurs d'accès un paquet d'information contenant son mode de fonctionnement GHMM, son adresse et la longueur de son préfixe réseau. Lorsqu'un nœud mobile entre dans un domaine GHMM, il reçoit une annonce de routeur

de son routeur d'accès avec des options signalant le support de **GHMM** et des informations sur le MAP.

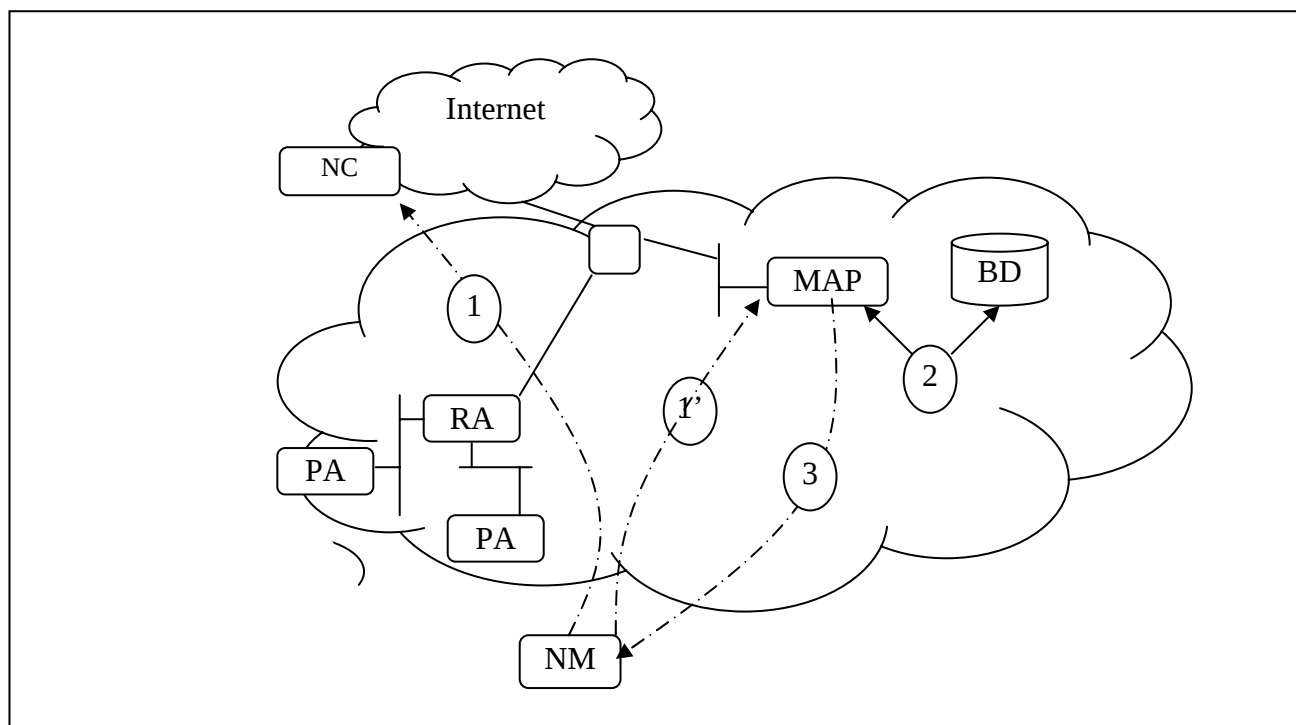


Fig 8.6 : Comportement d'un nœud mobile dans un domaine GHMM

1 : Mise à jour d'association

2: consultation de la base de données

1': Enregistrement de l'association

2' : Accusé de réception et envoi de la liste des PAs

Le nœud mobile acquiert alors deux nouvelles adresses IP temporaires (RCoA, LCoA). La première adresse, adresse temporaire virtuelle @V, est valide au niveau du gestionnaire de mobilité : elle est déterminée en concaténant l'identifiant de l'interface physique du nœud mobile avec le préfixe réseau du gestionnaire de mobilité recueilli dans l'annonce de routeurs. La deuxième adresse, adresse temporaire locale @L, identifie la liaison locale ; elle est obtenue en concaténant l'identifiant de l'interface physique du **nœud** mobile avec le préfixe réseau du routeur d'accès courant.

L'adresse temporaire virtuelle joue le rôle de l'adresse temporaire primaire de Mobile **IPv6**. Le nœud mobile informe alors son agent mère **et** ses nœuds correspondants NC de cette adresse temporaire en **utilisant des messages de mise à jour d'associations Mobile IPv6** (étape 1 dans la Figure 8.6). En outre, pour que le gestionnaire de mobilité soit capable de livrer ces paquets, il doit être informé de la correspondance entre l'adresse virtuelle et l'adresse locale. Pour cela, dès l'acquisition de ces deux adresses temporaires, le nœud mobile informe le gestionnaire de mobilité de la correspondance entre l'adresse **virtuelle** et l'adresse locale (@V, @L) en lui transmettant une **requête d'enregistrement**

HMIPv6. Ainsi, tous les paquets **adressés** au nœud mobile sont routés vers son adresse temporaire virtuelle @V, jusqu'au gestionnaire de mobilité. **Ensuite, le** gestionnaire de mobilité intercepte ces paquets **et** les livre à l'adresse locale @L du nœud mobile. Le gestionnaire de mobilité, après avoir consulté sa base de données (étape 2), acquitte la requête **d'enregistrement en renvoyant** au nœud mobile un **message d'acquittement d'associations Mobile IPv6** qui doit inclure, en outre, la liste des points d'accès **voisins** du point d'accès courant du nœud mobile dans GHMM (étape3)

8.5.4. Gestion du handover dans GHMM

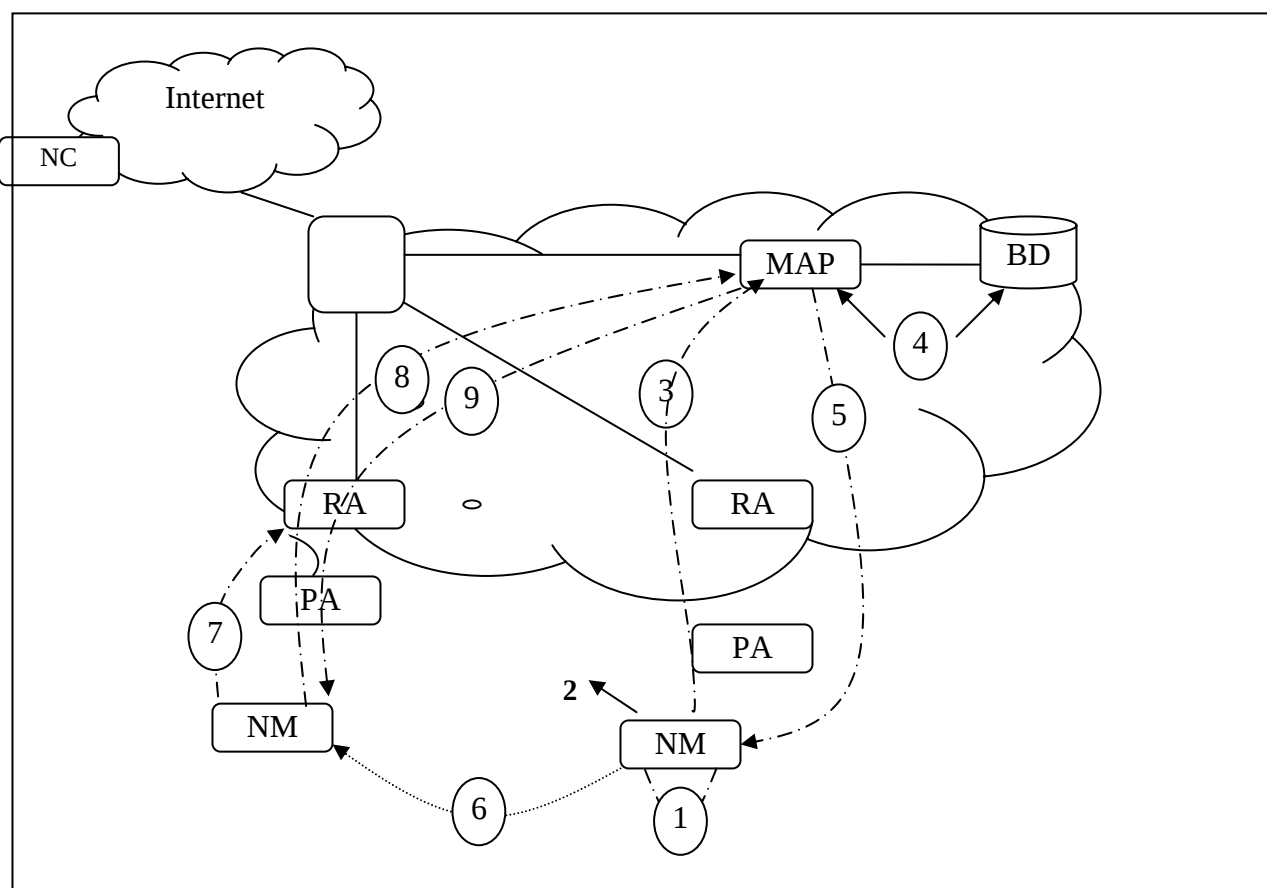


Fig 8.7 : Gestion du handover dans le GHMM.

- 1 Si la qualité de signal <S' (seuil du handover) le composant DD du NM détecte le mouvement
- 2 DD lance le déclencheur L2.
- 3 Initiation de la procédure du handover
 - AH lance la procédure d'anticipation
 - DD lance la découverte des RA s
 - OH initie la procédure d'optimisation (le buffering et le bicasting).

- 4 Consultation de la BD dans le MAP pour la récupération des informations des PAs
- 5 Le MAP envoie la nouvelle @ du RA via l'ancien (grâce au mécanismes d'anticipation)
- 6 Le NM se déplace après la réception du message 5 et après la construction de l'@ temporaire (@L2)
- 7-8 Le NM envoie une sollicitation au nouveau RA et le composant CH se charge de L'enregistrement du NM
- 9 Réception d'acquittement et des listes des PAs mises à jours

8.6. Conclusion

Dans ce chapitre, nous avons proposé notre solution GHMM pour la gestion de la mobilité (handover) dans un WLAN lors d'une transmission vidéo en temps réel, cette solution est basée sur le modèle générique proposé dans le chapitre 7 et le schéma hiérarchique HMIPv6. Dans ce chapitre nous avons amélioré les mécanismes de chaque composant du modèle générique, dans le cadre de notre solution GHMM, basée aussi sur un contrôle combiné du handover(handover initié par le NM et contrôle par le NM et le MAP). Nous résumons dans ce qui suit les améliorations apportées dans la solution GHMM :

Le protocole de découverte CARD adopté pour la solution permet d'améliorer les performances globales du handover, ce protocole fournit les informations de voisinage nécessaires pour une bonne anticipation du handover et pour un smoothing optimal. En plus, nous avons proposé un schéma du buffering qui garantit moins de perte des paquets envoyés en temps réel (flux vidéo). Le schéma offre aussi le meilleur instant pour déclencher le handover et qui permet d'éviter la saturation des buffers et la duplication des paquets. Le deuxième concept de smoothing amélioré, c'est le bicating. Nous avons proposé dans ce chapitre un schéma de bicasting optimal en combinant les mécanismes de bicating et la procédure de pré enregistrement, ce schéma permet d'optimiser la perte des paquets d'où il est capable de délivrer les paquets de données pendant la phase du handover. En effet GHMM a pour but d'améliorer la gestion du handover L3.

Notre solution apporte les points suivants :

- Introduction de nouveaux mécanismes de détection de mouvement
- le contrôle combiné du handover(NM, MAP)
- Améliorer le schéma du buffering et du bicasting
- Permettre une communication des composants durant le handover
- Réaliser moins de perte et de retard pour les applications vidéo en temps réel.

1. LES RÉSEAUX SANS FIL IEEE 802.11

1.1 Introduction

Un nouveau type de LANs (Local Area Network) est apparu depuis quelques années : les LANs sans fil (Wireless LAN ou WLAN). Ils proposent une alternative aux LANs traditionnels à base de paire torsadée, câble coaxial ou fibre optique. Les LANs sans fil ont le même rôle que les LANs filaires : transporter l'information entre des terminaux reliés à un LAN.

Cependant, l'absence de câblage physique rend le réseau beaucoup plus flexible; le déplacement d'un nœud sans fil est facile, comparé au travail nécessaire pour ajouter ou supprimer des câbles dans tout autre type de réseau. Le LAN sans fil est un système de transmission des données conçu pour assurer une liaison indépendante de l'emplacement des périphériques informatiques qui composent le réseau et utilisant les ondes radio plutôt qu'une infrastructure câblée. Dans l'entreprise, les LANs sans fil sont généralement implémentés comme le lien final entre le réseau câblé existant et un groupe d'ordinateurs clients, offrant aux utilisateurs de ces machines un accès sans fil à l'ensemble des ressources et des services du réseau de l'entreprise, sur un ou plusieurs bâtiments. Les WLANs sont en passe de devenir l'une des principales solutions de connexion pour de nombreuses entreprises. Le marché du sans fil se développe rapidement dès lors que les entreprises constatent les gains de productivité qui découlent de la disparition des câbles. Selon Frost et Sullivan, l'industrie du LAN sans fil, qui a dépassé les 300 millions de dollars en 1998, en représentera 1,6 milliard en 2005.

[Lao02]

Le but de ce chapitre est de donner un aperçu technique du standard 802.11, de façon à comprendre les concepts de base, le principe des opérations, et quelques arguments expliquant l'utilité telle fonction ou tel composant du standard IEEE 802.11.

Le standard IEEE 802.11 est actuellement le type de réseau le plus vogue pour les réseaux locaux. (Le standard IEEE 802.11). Il utilise les bandes de fréquences ISM situées aux alentours de 900 MHz, 2.4 GHz, 5.7 GHz. Suivant les pays, différentes fréquences, différentes modulations sont autorisées avec différentes puissances. Par exemple, les Etats-Unis autorisent les canaux (1-13) de la bande 2.4 GHz, le Mexique 10 et 11 seulement, la France 10 à 13, le Japon 1 à 14,...

La puissance est de même limitée suivant les pays (; *100 mw en intérieur en France*), par exemple en France la puissance permise est de 100 mw pour une utilisation interne. Il existe plusieurs groupes de travail : 802.11b pour un débit partagé de 11Mbps, 802.11a pour un débit de 54 Mbps et la bande de fréquences de 5 Ghz, et 802.11g avec un débit de 54 Mbps.

1.2. Problèmes des transmissions radios et des réseaux sans fil

Dans cette section nous présentons les principaux problèmes des transmissions radio et des réseaux sans fil :

- Débit plus faible que celui du filaire : la bande passante est une ressource rare, il faut minimiser la portion utilisée pour la gestion du réseau afin de pouvoir laisser le maximum de bande passante pour les communications données.
- Erreurs de transmission : les erreurs de transmission de radio sont plus fréquentes que dans les réseaux filaires.
- Interférences : les liens radio ne sont pas isolés, deux transmissions simultanées sur une même fréquence ou, utilisant des fréquences proches peuvent interférer. De plus, les interférences peuvent provenir d'autres types de machines non dédiées aux télécommunications. Par exemple, les fréquences utilisées dans les fours à micro-ondes sont dans les fréquences de la bande ISM et peuvent perturber les communications radio en cours.
- Liens versatiles : les transmissions radio sont très sensibles aux conditions de propagation, ce qui les rend versatiles. Un contrôle de la qualité des liens est obligatoire afin de pouvoir les exploiter convenablement pour les communications radio.
- Puissance du signal : la puissance du signal diminue avec la distance, et la puissance utilisée est sévèrement réglementée par les autorités compétentes des pays.
- Portée radio limitée : la propagation des ondes radio dépend de l'environnement de la transmission. Plusieurs facteurs peuvent limiter la portée d'une transmission, comme la faible puissance du signal, les obstacles (exemple : murs) qui empêchent, atténuent (suivant la matière), ou réfléchissent les signaux, rendant la tâche de détection et de synchronisation plus difficile à la réception.
- Energie : parler sur le sans fil c'est parler aussi sur la mobilité et donc autonomie. Pour maximiser la durée de vie des batteries, il faut économiser autant que possible les transmissions inutiles.

- Sécurité : les détecteurs des signaux et les récepteurs passifs peuvent espionner les communications radio si ces dernières ne sont pas protégées.
- Mobilité et topologie changeante : la disparition ou l'apparition d'un nœud peut être le résultat d'un déplacement relatif du nœud considéré ou du nœud voisin, d'un environnement changeant : en raison par exemple d'un obstacle entre les nœuds, d'une batterie épuisée ou simplement d'une panne routeur[Lao02].

1.3. Architecture IEEE 802.11

1.3.1 Description des couches physiques, MAC, Réseau et transport

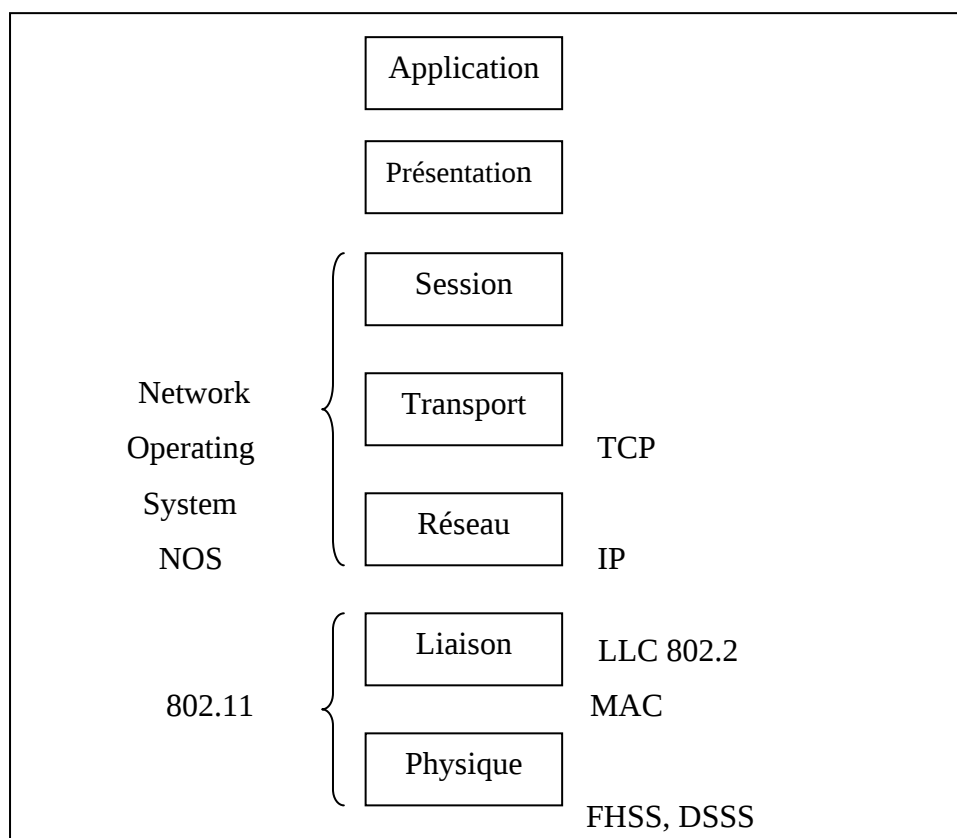


Fig 1.1 Architecture Wlan 802.11

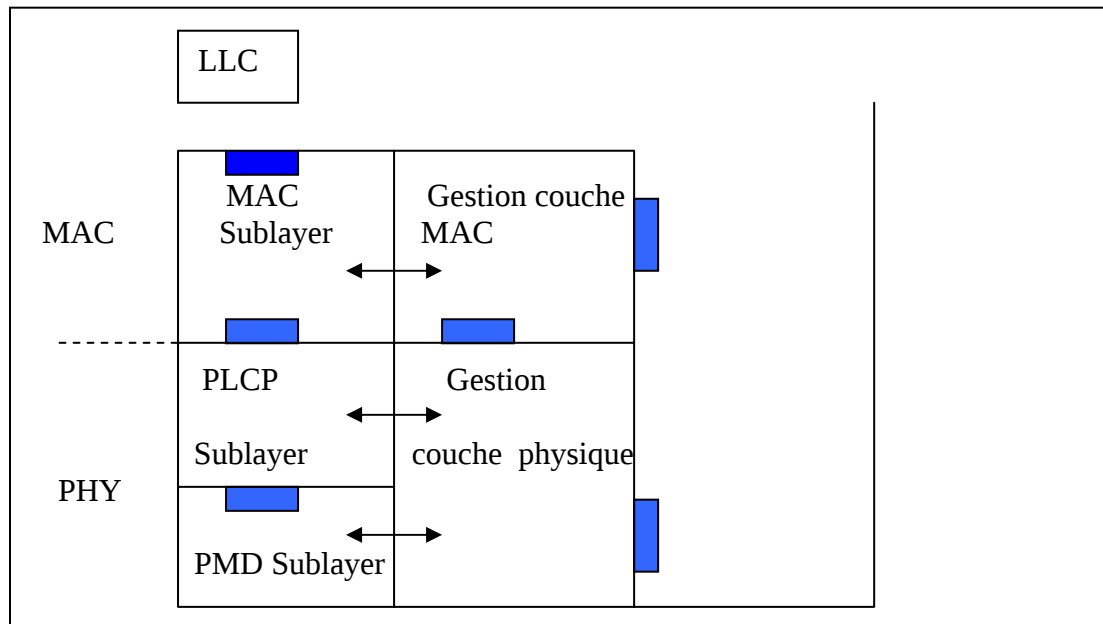


Fig 1.3. Architecture détaillée de la couche 1 et 2

1.3.2. Couche Physique

1.3.2.1. Les techniques radio (DSSS et FHSS) et les débits obtenus.

Les trois couches physiques définies à l'origine par 802.11 incluaient deux techniques radio à étalement de spectre et une spécification d'infrarouge diffus (que nous n'aborderons pas ici).

Les standards radio fonctionnent sur la bande ISM des 2,4 GHz. Ces fréquences sont reconnues par les organismes réglementaires internationaux tels que la FCC (Etats-Unis), l'ETSI (Europe) et le MKK (Japon) pour utilisation sans licence. Dans cette optique, les produits basés sur 802.11 ne nécessitent ni l'obtention d'une licence utilisateur ni une formation spécifique. Les techniques d'étalement du spectre, en plus de satisfaire aux conditions réglementaires, améliorent la fiabilité, accélèrent le débit et permettent à de nombreux produits non concernés de se partager le spectre sans coopération explicite et avec un minimum d'interférences. La version originale du standard 802.11 prévoit des débits de 1 et 2 Mbps sur des ondes radio utilisant une technologie d'étalement de spectre avec sauts de fréquence (FHSS) ou en séquence directe (DSSS). Il est important de remarquer que FHSS et DSSS sont des mécanismes de signalisation fondamentalement différents l'un de l'autre et qu'aucune interopérabilité ne peut être envisagée entre eux [Lao02].

A. FHSS

Par la technique des sauts de fréquence (FHSS), la bande des 2,4 GHz est divisée en 75 sous-canaux de 1 MHz. L'émetteur et le récepteur s'accordent sur un schéma de saut, et les données sont envoyées sur une séquence de sous-canaux. Chaque conversation sur le réseau 802.11 s'effectue suivant un schéma de saut différent, et ces schémas sont définis de manière à minimiser le risque que

deux expéditeurs utilisent simultanément le même sous-canal. Les techniques FHSS simplifient -- relativement -- la conception des liaisons radio, mais elles sont limitées à un débit de 2 Mbps, cette limitation résultant essentiellement des réglementations de la FCC qui restreignent la bande passante des sous-canaux à 1 MHz. Ces contraintes forcent les systèmes FHSS à s'étaler sur l'ensemble de la bande des 2,4 GHz, ce qui signifie que les sauts doivent être fréquents et représentent en fin de compte une charge importante.

B. DSSS

La technique de signalisation en séquence directe divise la bande des 2,4 GHz en 14 canaux de 22 MHz. Les canaux adjacents se recouvrent partiellement, seuls trois canaux sur les 14 étant entièrement isolés. Les données sont transmises intégralement sur l'un de ces canaux de 22 MHz, sans saut. Pour compenser le bruit généré par un canal donné, on a recours à la technique du "*chipping*". Chaque bit de donnée de l'utilisateur est converti en une série de motifs de bits redondants baptisés "chips" (cf la figure ci-dessous). La redondance inhérente à chaque chip associée à l'étalement du signal sur le canal de 22 MHz assure le contrôle et la correction d'erreur : même si une partie du signal est endommagée, il peut dans la plupart des cas être récupéré, ce qui minimise les demandes de retransmission. La version DSSS de la norme IEEE 802.11 est conçue pour parer aux interférences. Il faut cependant noter qu'aucun système LAN sans fil actuel ne vient à bout du brouillage intentionnel. La technologie DSSS envoie en simultanée l'information sur plusieurs canaux parallèles, ce qui donne un taux d'erreur plus faible (donc un débit plus élevé) et une immunité aux perturbations en bande étroite. La technologie FHSS est basée sur le saut de fréquence, ce qui permet d'économiser de la bande passante

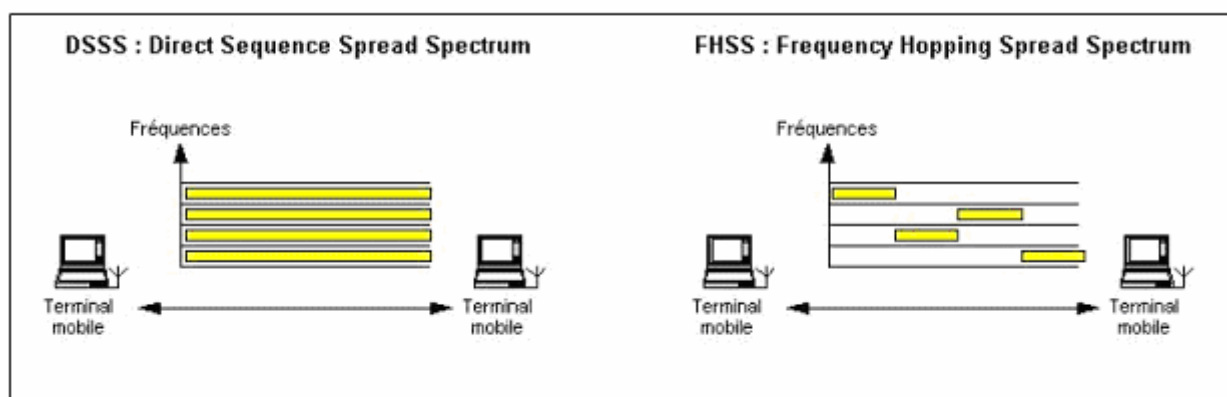


Fig 1.4. Méthodes DSSS et FHSS

1.3.3. Liaison MAC

En plus des fonctions habituellement rendues par la couche MAC, la couche MAC 802.11 offre d'autres fonctions qui sont normalement confiées aux protocoles supérieurs, comme la fragmentation, les retransmissions de paquet et les accusés de réception.

La couche MAC définit deux méthodes d'accès différentes, la « Distributed Coordination Function » et la « Point Coordination Function » (Nous verrons plus de détails dans le paragraphe 1.3.3.4.). L'entité MAC assure les mécanismes d'accès de base, la fragmentation et le cryptage. L'entité MAC Layer Management assure quant à elle la synchronisation, la gestion de la puissance et le roaming [Had02]

1.3.3.1. Fragmentation et réassemblage

Les protocoles de réseaux locaux classiques utilisent des paquets de plusieurs centaines d'octets (eg. les paquets Ethernet peuvent atteindre 1518 octets). Dans un environnement de réseau local sans fil, il y a plusieurs raisons d'utiliser des paquets plus petits :

1. A cause du taux d'erreur par bit qui est plus important sur une liaison radio. La probabilité d'un paquet d'être corrompu augmente avec sa taille.
2. Dans le cas d'un paquet corrompu (à cause d'une collision ou même du bruit), plus le paquet est petit, moins le surdébit engendré par sa retransmission est important.
3. Dans un système à saut de fréquence, le support est interrompu périodiquement pour ce changement de fréquence (dans notre cas, toutes les 20 ms), donc plus le paquet est petit, plus la chance d'avoir une transmission interrompue est faible.

Il n'est pas utile de créer un nouveau protocole LAN incapable de traiter les paquets de 1518 octets utilisés sur Ethernet. Le comité a donc décidé de résoudre ce problème en ajoutant un simple mécanisme de fragmentation et réassemblage au niveau de la couche MAC.

Ce mécanisme se résume à un algorithme simple d'envoi et d'attente de résultat, où la station émettrice n'est pas autorisée à transmettre un nouveau fragment tant qu'un des deux événements suivants n'est pas survenu :

1. Réception d'un ACK pour le fragment correspondant.
2. Décision que le fragment a été retransmis trop souvent et abandon de la transmission de la trame.

Il est à noter que le standard autorise une station à transmettre à des adresses différentes entre les retransmissions d'un certain fragment. Ceci est particulièrement utile quand un Point d'accès a plusieurs paquets en suspens pour plusieurs destinations différentes et qu'une d'entre elles ne répond pas.

1.3.3.2 Espace entre deux trames (Inter Frame Space)

Le standard définit 4 types d'espace entre deux trames, utilisés pour leurs différentes propriétés:

- **SIFS** (Short Inter Frame Space) est utilisé pour séparer les transmissions appartenant à un même dialogue (eg.Fragment – Ack). C'est le plus petit écart entre deux trames et il y a

toujours, au plus, une seule station pour transmettre à cet instant, ayant donc la priorité sur toutes les autres stations. Cette valeur est fixée par la couche physique et calculée de telle façon que la station émettrice sera capable de commuter en mode réception pour pouvoir décoder le paquet entrant. Pour la couche physique FH de 802.11, cette valeur est de 28 microsecondes.

- **PIFS** (Point Coordination Fonction IFS) est utilisé par le Point d'accès (appelé point coordinateur dans ce cas) pour gagner l'accès au support avant n'importe quelle autre station. Cette valeur est SIFS plus un certain temps (Slot Time, défini dans le paragraphe suivant), soit 78 microsecondes.
- **DIFS** (Distributed Coordination Fonction IFS) est l'IFS utilisé par une station voulant commencer une nouvelle transmission et calculé comme étant PIFS plus un temps, soit 128 microsecondes.
- **EIFS** (Extended IFS) est l'IFS le plus long. Il est utilisé par une station quand la couche physique indique à la couche MAC qu'une transmission de trame a débuté et qu'aucune réception de trame MAC avec un FCS correct n'a eu lieu. Ceci est nécessaire pour éviter que la station ne provoque de collision avec un futur paquet du dialogue en cours [Wan02]

1.3.3.3. Algorithme de backoff exponentiel

Le backoff est une méthode bien connue pour résoudre les différends entre plusieurs stations voulant avoir accès au support. Cette méthode demande que chaque station choisisse un nombre aléatoire n entre 0 et un certain nombre, et d'attendre n slots avant d'accéder au support, toujours en vérifiant qu'une autre station n'a pas accédé au support avant elle.

La durée d'un slot (Slot Time) est définie de telle sorte que la station sera toujours capable de déterminer si une autre station a accédé au support au début du slot précédent. Cela divise la probabilité de collision par deux.

Le backoff exponentiel signifie qu'à chaque fois qu'une station choisit un slot et provoque une collision, le nombre maximum pour la sélection aléatoire est augmenté exponentiellement.

Le standard 802.11 définit l'algorithme de backoff exponentiel comme devant être exécuté dans les cas suivants :

- Quand la station écoute le support avant la première transmission d'un paquet et que le support est occupé.
- Après chaque retransmission.
- Après une transmission réussie.

Le seul cas où ce mécanisme n'est pas utilisé est quand la station décide de transmettre un nouveau paquet et que le support a été libéré pour un temps supérieur au DIFS. La figure suivante montre le principe du mécanisme d'accès : Le mécanisme d'accès CSMA/CA le protocole ACK.

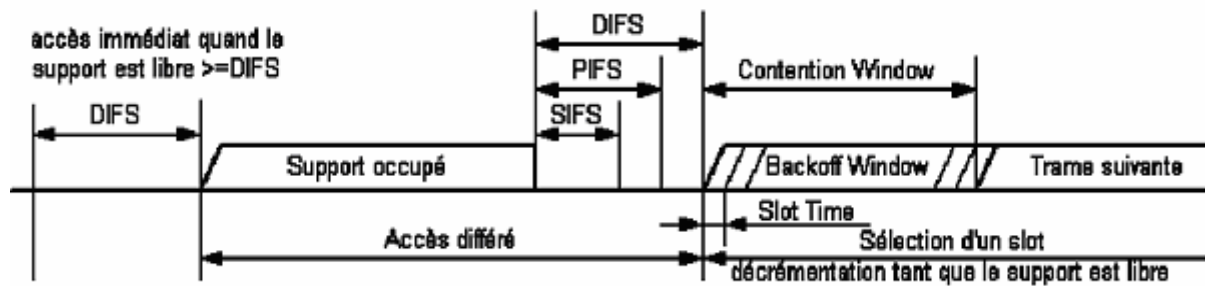


Fig1.5. Explication du Backoff exponentiel (sans prendre en compte les ACK)

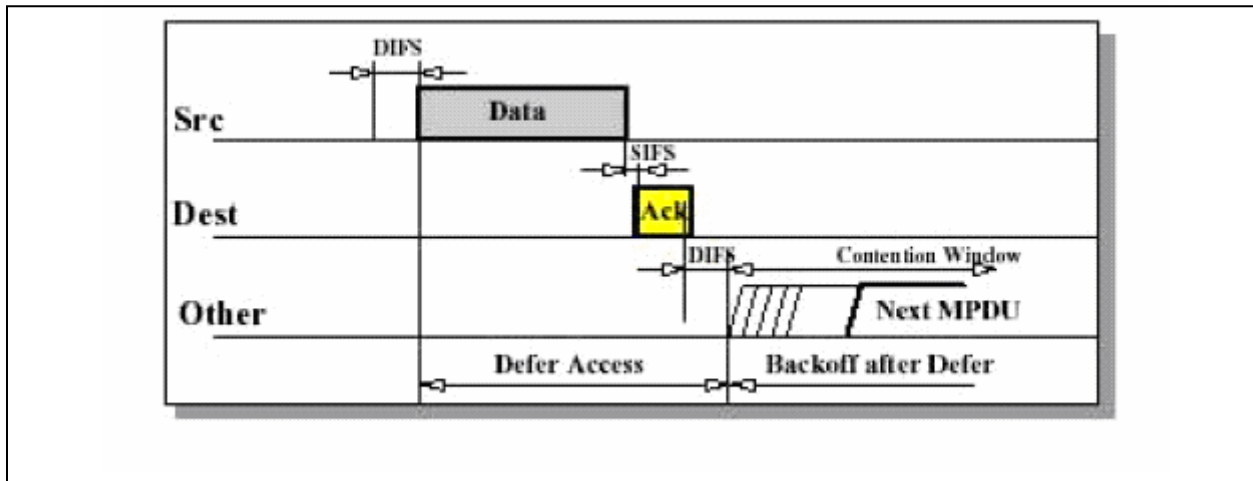


Fig 1.6 Backoff exponentiel + ACK

1.3.3.4. Méthode d'accès de base CSMA/CA

A. Introduction des protocoles CSMA

Le mécanisme d'accès de base, appelé *Distributed Coordination Function* est typiquement le mécanisme CSMA/CA. Les protocoles CSMA sont bien connus dans l'industrie, où le plus célèbre est Ethernet, qui est un protocole CSMA/CD. Un protocole CSMA fonctionne comme suit : une station voulant émettre écoute le support de transmission, et si le support est occupé (ie. une autre station est en train d'émettre), alors la station remet sa transmission à plus tard. Si le support est libre, la station est autorisée à transmettre. Ces types de protocoles sont très efficaces quand le support n'est pas surchargé, puisqu'il autorise les stations à émettre avec un minimum de délai, mais il y a toujours une chance que des stations émettent en même temps (collision). Ceci est dû au fait que les stations écoutent le support, le repèrent libre, et finalement décident de transmettre, parfois en même temps qu'une autre exécutant cette même suite d'opérations.

Ces collisions doivent être détectées, pour que la couche MAC puisse retransmettre le paquet sans avoir à repasser par les couches supérieures, ce qui engendrerait des délais significatifs. Dans le cas d'Ethernet, cette collision est repérée par les stations qui transmettent, celles-ci allant à la phase de retransmission basée sur un algorithme de retour aléatoire exponentiel (exponential random backoff).

Si ces mécanismes de détection de collisions sont bons sur un réseau local câblé, ils ne peuvent pas être utilisés dans un environnement sans fil, pour deux raisons principales :

1. Implémenter un mécanisme de détection de collisions demanderait l'implémentation d'une liaison radio full duplex, capable de transmettre et de recevoir immédiatement, une approche qui en augmenterait significativement le prix.
2. Dans un environnement sans fil, on ne peut être sûr que toutes les stations s'entendent entre elles (ce qui est l'hypothèse de base du principe de détection de collision), et le fait que la station voulant transmettre teste si le support est libre, ne veut pas forcément dire que le support est libre autour du récepteur.

B. DCF : Distributed Coordination Function

Le mécanisme d'accès de base, appelé DCF est typiquement le mécanisme CSMA/CA (*Carrier Sense Multiple Acces with Collision Avoidance*). Le protocole 802.11 utilise le mécanisme d'esquive de collision (Collision Avoidance), ainsi que le principe d'accusé de réception (Positif Acknowledge) comme suit : Une station voulant transmettre écoute le support, et s'il est occupé, la transmission est différée. Si le support est libre pour un temps spécifique (DIFS), alors la station est autorisée à transmettre.

La station réceptrice va vérifier le CRC du paquet reçu et renvoie un accusé de réception ACK. La réception de l'ACK indiquera à l'émetteur qu'aucune collision n'a eu lieu. Si l'émetteur ne reçoit pas l'accusé de réception, alors il retransmet le fragment jusqu'à ce qu'il l'obtienne ou abandonne au bout d'un certain nombre de retransmissions.

Une station voulant émettre transmet d'abord un petit paquet de contrôle appelé RTS (Request To Send), qui donnera la source, la destination, et la durée de la transaction.

Pour résoudre le problème de la « station cachée », le protocole 802.11 a défini le mécanisme de Virtual Carrier Sense. Il est basé sur l'un des premiers protocoles développés pour les WLANs, le MACA (Multiple Access with Collision Avoidance) qui fut mis au point par Karn en 1990[Bra02].

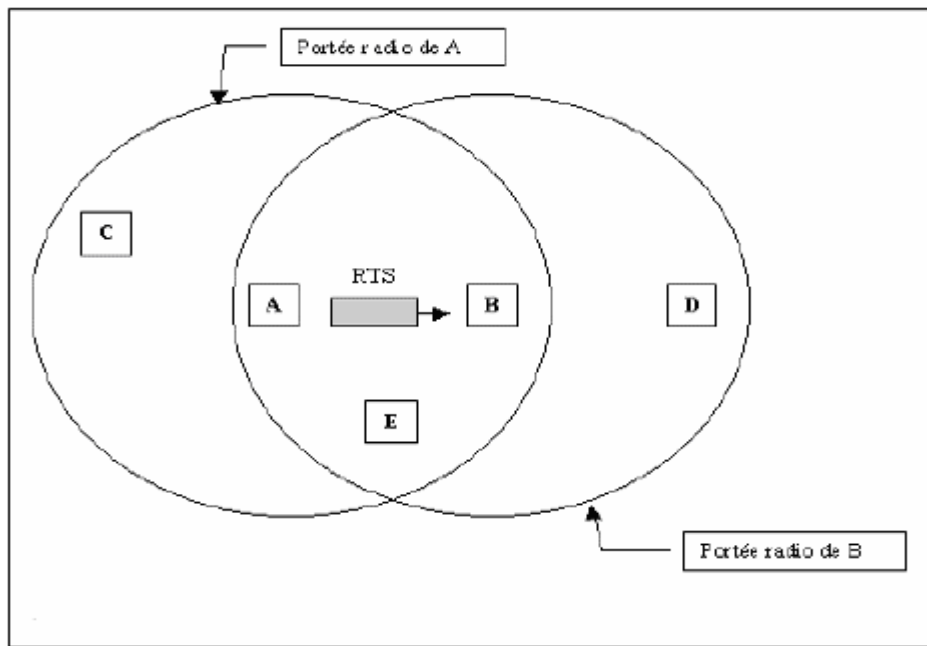


Fig 1.7. La station émettrice (A) envoie un RTS.

Une station voulant émettre transmet d'abord un petit paquet de contrôle appelé RTS (Request To Send), qui donnera la source, la destination, et la durée de la transaction.

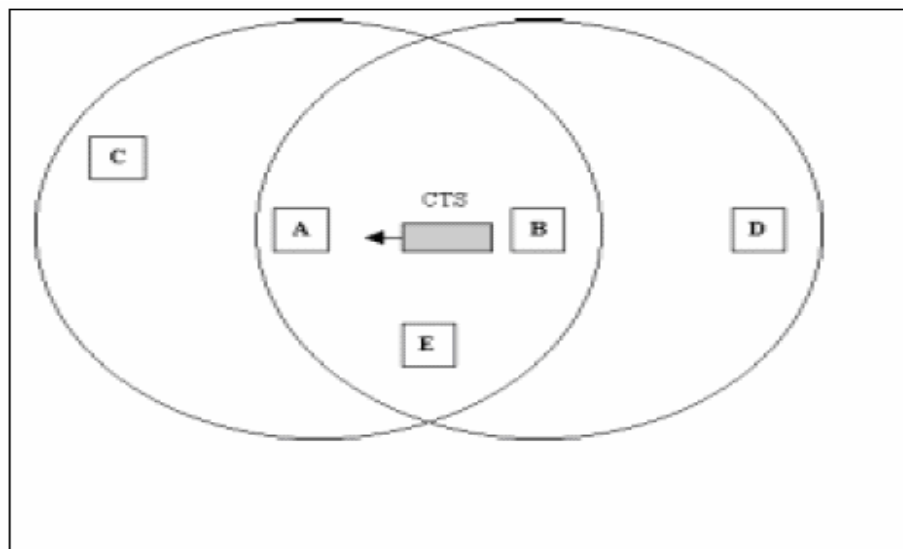


Fig 1.8. Réponse de la station réceptrice: Envoi d'un CTS.

La station destinataire répond (si le support est libre) avec un paquet de contrôle de réponse appelé CTS (*Clear To Send*), qui inclura les mêmes informations sur la durée.

Toutes les stations recevant soit le RTS, soit le CTS, déclencheront leur indicateur de Virtual Carrier Sense (appelé NAV pour *Network Allocation Vector*), pour une certaine durée, et utiliseront cette information avec le Physical Carrier Sense pour écouter le support. Il est également à noter que grâce au fait que le RTS et le CTS sont des trames courtes (30 octets), le nombre de collisions est réduit, puisque ces trames sont reconnues plus rapidement que si tout le paquet devait être transmis. (Ceci est vrai si le paquet est beaucoup plus important que le RTS, donc le standard autorise les paquets courts à

être transmis sans l'échange de RTS/CTS, ceci étant contrôlé pour chaque station grâce au paramètre appelé *RTS Threshold*).

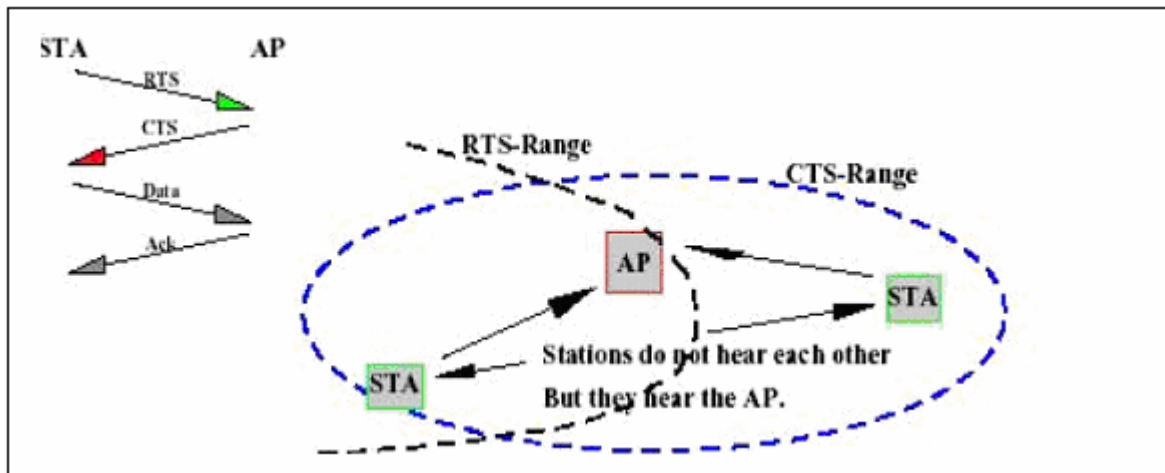


Fig 1.9. Station cachée.

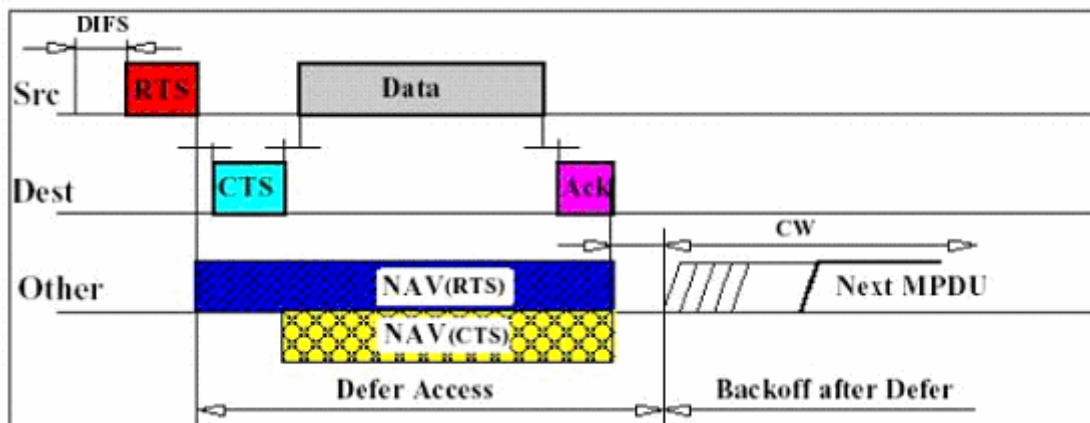


Fig 1.10. NAV.

C. PCF : Point Coordination Function

En plus de la fonction de base de coordination distribuée (DCF), il y a la fonction optimale de coordination par point (PCF) qui peut être utilisée pour implémenter des services temps réel, comme la transmission de voix ou de vidéo.

En utilisant un accès par priorité supérieure, le Point d'Accès peut envoyer des données aux stations en réponse à une Polling Request, tout en contrôlant l'accès au support. Pour permettre aux stations classiques d'avoir accès au support, il y a une condition qui est que le Point d'Accès doit laisser suffisamment de temps DCF par rapport au PCF.

Il est intéressant de remarquer que le DCF et le PCF peuvent coexister avec le même DSS. Les deux méthodes d'accès sont alternées avec des CFP (*Contention Free Period*) et des CP (*Contention Period*).

1.3.4. Synchronisation

Les stations doivent rester synchronisées, ceci est nécessaire pour garder la synchronisation au cours des sauts, ou pour d'autres fonctions comme l'économie d'énergie. Sur une même cellule, ceci est obtenu car toutes les stations synchronisent leur horloge avec l'horloge du Point d'accès en utilisant le mécanisme suivant :

Le Point d'accès transmet périodiquement des trames appelées « trames balise ». Ces trames contiennent la valeur de l'horloge du Point d'accès au moment de la transmission (notons que c'est le moment où la transmission a réellement lieu, et non quand la transmission est mise à la suite des transmissions à faire.

Les stations réceptrices vérifient la valeur de leur horloge au moment de la réception, et la corrige pour rester synchronisées avec l'horloge du Point d'accès. Ceci évite des dérives d'horloge qui pourraient causer la perte de la synchronisation au bout de quelques heures de fonctionnement.

1.3.5. L'économie d'énergie

Les réseaux sans fil sont généralement en relation avec des applications mobiles, et dans ce genre d'applications, l'énergie de la batterie est une ressource importante. C'est pour cette raison que le standard 802.11 donne lui-même des directives pour l'économie d'énergie et définit tout un mécanisme pour permettre aux stations de se mettre en veille pendant de longues périodes sans perdre d'information. L'idée générale, derrière le mécanisme d'économie d'énergie, est que le Point d'accès maintient un enregistrement à jour des stations travaillant en mode d'économie d'énergie, et garde les paquets adressés à ces stations jusqu'à ce que les stations les demandent avec une Polling Request, ou jusqu'à ce qu'elles changent de mode de fonctionnement.

Les Points d'accès transmettent aussi périodiquement (dans les trames balise) des informations spécifiant quelles stations ont des trames stockées par le Point d'accès. Ces stations peuvent ainsi se réveiller pour récupérer ces trames balise, et si elles contiennent une indication sur une trame stockée en attente, la station peut rester éveillée pour demander à récupérer ces trames. Les trames de multicast et de broadcast sont stockées par le Point d'accès et transmises à certains moments (chaque DTIM) où toutes les stations en mode d'économie d'énergie, qui veulent recevoir ce genre de trames, devraient rester éveillées.

1.4. Avantages

Les WLANs libèrent l'utilisateur de sa dépendance à l'égard des accès câblés, lui offrant un accès permanent et omniprésent. Cette liberté de mouvement offre de nombreux avantages dans de nombreux types d'environnements de travail tels que :

- Accès immédiat entre le lit d'hôpital et les informations concernant le patient pour les médecins et le personnel hospitalier.
- Un accès réseau simple et en temps réel pour les consultants et les auditeurs sur site.
- Un accès étendu aux bases de données pour les chefs de service nomades, directeurs de chaîne de fabrication, contrôleurs de gestion ou ingénieurs.
- Une configuration simplifiée du réseau avec un recours minime au personnel informatique pour les installations temporaires telles que stands de foire, d'exposition ou salles des conférences.
- Un accès plus rapide aux informations client pour les fournisseurs de services et détaillants, résultant en un meilleur service et une satisfaction supérieure.
- Un accès omniprésent au réseau pour les administrateurs, pour le support et le dépannage sur site.
- Un accès en temps réel pour les réunions des groupes d'études des liens de recherche pour les étudiants.

1.5. Conclusion

Le but de ce chapitre est de donner un aperçu technique du standard 802.11b, de façon à comprendre les concepts de base, le principe des opérations, et quelques raisons expliquant telle fonction ou tel composant du standard. De toute évidence, ce chapitre couvre l'ensemble des fonctions du standard 802.11b, et à travers les différentes architectures décrites est de pouvoir offrir des services sans fil les plus proches possibles des services filaires. Actuellement, les LANs sans fil offrent des solutions (normalisées) pour les réseaux locaux privés, mais la mobilité qu'ils proposent est relativement réduite, en terme de zone géographique et de vitesse de déplacement. Selon l'environnement et les applications, on optera pour des réseaux radiofréquences.

Conclusion

Depuis l'apparition du problème de transmission vidéo, problème de la mobilité, des liens sans fil et de la compression, plusieurs solutions ont été développées, tant au niveau de la couche MAC qu'au niveau des couche IP et application.

Notre solution, contrairement aux précédentes, agit au niveaux MAC, IP et application, ce qui lui permet d'exploiter au mieux les caractéristiques de ces trois couches.

Elle est basée sur l'utilisation du standard de compression vidéo MPEG et la norme des Wlan 802.11b et MIPv6 conjointement, pour étudier la transmission vidéo sur des liaisons sans fil.

L'objectif de ce travail était d'améliorer la gestion du handover ainsi que la qualité de service et la qualité vidéo pour des applications temps réel. Cependant, pour que de telles applications soient mises sur le WLAN, il faudrait assurer un minimum de conditions, c'est pourquoi nous nous sommes intéressés au problème de la gestion de la mobilité et des liens sans fil.

Pour cela ce mémoire s'est focalisé d'une part sur une étude bibliographique des standards nécessaires pour la transmission vidéo dans des liaisons sans fil (MPEG, 802.11b, MIPv6), et d'autre part sur nos solutions proposées, auxquels nous avons rajouté certaines solutions.

En premier lieu Nous avons étudié l'impact des liens sans fil sur la transmission vidéo, ensuite nous avons proposé un modèle de simulation de la transmission vidéo qui nous a permis d'étudier le comportement de la vidéo MPEG sur un Wlan 802.11. Ce modèle nous a permis aussi de définir un ensemble de paramètres et seuils nécessaires pour une transmission acceptable.

Nous avons présenté aussi la procédure du handover couche liaison, et nous avons essayé de l'adapter et l'améliorer pour un trafic vidéo. Après nous avons établi un test pratique d'une transmission vidéo durant le handover couche liaison.

Nous avons présenté également une étude détaillée de la gestion de la mobilité (handover couche IP). En premier lieu, nous avons présenté une étude comparative des différentes extensions MIPv6. Sur la base de cette étude, nous avons présenté un modèle générique à base de composants pour la gestion du handover pour un trafic vidéo. En fin nous avons définie une solution **GHMM** basé sur le modèle générique et le contrôle combiné du handover (NM, agent de mobilité). L'architecture hiérarchique du protocole GHMM permet de réduire la charge de signalisation, ainsi que les composants d'anticipation et de détection de mouvement du protocole GHMM optimisent la latence du handover. En plus le GHMM fournit une meilleure optimisation de la perte de paquets grâce aux mécanismes de smoothing (buffering, bicating).

Perspectives

La transmission de la vidéo sur des liaisons sans fil est très récente. Plusieurs domaines et axes de recherche intéressants liés à ce problème restent à explorer. Nous pouvons citer quelques axes à explorer :

- Etudier la transmission de vidéo sur des réseaux plus large (internet mobile) est aussi un axe intéressant.
- Etudier les capacités et les performances multimédia d'une cellule.
- La simulation de notre modèle et solution pour la gestion du handover
- L'optimisation de la bande passante et minimiser les taux de perte et les retards Sont des problèmes très importants à étudier malgré les travaux existants.
- Etudier le **vertical handover** (communication : UMTS, 802.11, bluetooth...) est un axe très intéressant aussi
- Etudier l'aspect de sécurité pour une transmission vidéo
- Introduire les mécanismes de prédiction dans la solution GHMM pour optimiser la détection de mouvement et la découverte des RAs.

11. Bibliographie

- [Adj01] C.Adjh. *Multimedia et accès à l'Internet haut débit - l'analyse de la filière câble-*. Thèse de doctorat, Université de Versailles- saint- Quentin en Yvelines, June 2001.
- [Ala01] K.Alagha, G Pujolle , G.vivier, Réseaux mobiles et réseaux sans fils, 2001.
- [Alb96] Albance(A), J.Blomer et J.Edmonds. *Priority encoding transmission*. IEEE Transactions on information Theory, 42(6) : 1737- 1744, 1996.
- [Apo01] Apostolopoulos John, *Reliable video communication over lossy packets networks using multiple state encoding and path diversity*, visual communications and image processing (VCIP), Januray 2001.
- [Apo02] Apostolopoulos John, Tina Wong, Wai-tian Tan, Susie Wee. *On Multiple Description Streaming with Content Delivery Networks* , Proceedings of IEEE INFOCOM, New York, NY. June 2002
- [Bel99] L.Bellier. *heararchical Mobility management freamwork for internet*, in Proceeding of MOMUC'99, san diego, Novembre 1999
- [Bel03] Farouk Belghoul, Yan Moret, Christian Bonnet. *Performance Comparison and Analysis on MIPv6, Fast MIPv6 Bi-casting and EurecomIPv6 Soft Handover over IEEE802.11b WLANs*. Department of Mobile Communications, Institute Eurecom,2003
- [Bel04] Farouk Belghoul ,Yan Moret,Christian Bonnet. *Mécanismes de handover pour les réseaux IP sans-fil*, TSI Journal 2004.
- [Ben01] Racha BenAli, Hossam Afifi. *Etudes de performance des protocoles des réseaux locaux sans fil IEEE 802.11 avec introduction de la QoS*, rapport 2001.
- [Bou03] Imad Bouazizi, Mesut Gunes. *A video streaming Freamwork for transporting MPEG video over Mobile Ad Hoc Networks*, Proccedings of the Med-Hoc NET 2003 Workshop, Mehdiia, Tunisia 25-27, June 2003.
- [Bra02] Brassac Anne François; *Les réseaux sans fil*, DESS Miage, Université Paul Sabatier, Toulouse 2002.
- [Cai97] J.Cai, D.J Goodman, *General Paquet Radio Services In GSM*, IEEE Comm Magazin, Vol 35 N° 10, octobre 1997.
- [Cam02] A Campbell, J Gomez, S. KiM, *Comparison of IP micromobility Protocols*, IEEE wirelss communication .Feburary 2002.
- [Cas02] Castellucia C. *A hierarchical Mobile Ipv6 Proposal*. 4th CTS Mobile Communication Summit,Jun 2002.
- [Cle04] P.cleyn, N.Van dan , *A Smooth handoff scheme using IEEE802.11 triggers*. Computer Networks 45 , P 345-361, April 2004.
- [Cou00] Contantin Coutras, Sanjay Gupta, Ness B.Shroff, *Scheduling of real time traffic in IEEE802.11 Wireless LANs*, Wireless Networks 6, 457-466, 2000.

- [Dee91] Deering S., *ICMP Router Discovery Messages*, Request for Comments, RFC 1256, Internet Engineering Task Force, Spetembre 1991. <http://ftp.lip6.fr/pub/rfc/rfc/rfc1256.txt.gz>
- [Dou03] Douglas Howie, *Conséquence of Using MIPv6 to achieve mobile Ubiquitous multimedia*, Département of Elctrical and Information Engineering University, University of Oulu, Finland 2003.
- [Ens04] Ensoo shim, *Mobility Mangement in Wireless internet*. Phd thesis Columbia university, 2004.
- [Fab01] Fabrice le leannec, *Codage vidéo robuste et hiérarchique pour la transmission sur les réseaux hétérogènes*, Thèse doctorat de l'université de Rennes I , décembre 2001.
- [Fen97] W.C Feng , *Buffering Technique for delivery of compressed video in video on demand Systems*. Kluwer Academic Publishe, 1997.
- [Fen04] Fenping Li .*A study of mobilitu in Wlan*, HUT-y-110.551, seminar on interworking, SjöKulla, April 2004.
- [Fes02] A. Festag. *Optimization of Handover Performance by Link Layer Triggers in IP-Based Networks:Parameters, Protocol Extensions and APIs for Implementation*. TKN Technical Report TKN-02-014, Technical University Berlin Telecommunication Networks Group, Berlin, August 2002,
- [Fit02] F.Fitzek, P. Seeling, and M.Reisslein. *26 k pre- standard evaluation*. Tech. Rep.November 2002.
- [Fit03] F.Fitzek, P Seeling M.Reisslein. *Video streaming in wireless internet*. 2003
- [Gan02] Ganesha Bhaskara. *Multicast based micro-mobility protocol:design and evaluation*. Phd Thesis, Faculty of graduate school university of southern Claifornia ,August 2002.
- [Gen02] Genista, Vdeo quality metrics, Rap.tech, Genista Perceptual Quality Metrics, 2002.
- [Gog02] N.Gogate, D Chung Y. Wang . *Supporting image/ video applications in a mobile multihop radio environment using route diversity and multihop description coding*. IEEE Trans. On Circuits and systems for video technologie, September 2002.
- [Gra01] G.Legrand, *Qualité de service dans les environnements Internet mobile*. Thèse de doctorat de l'université de Pierre et Marie Curie. Juillet 2001.
- [Hai97] Hairuo .l , M.E. Zakari, S.Gupta. *Error control scheme for networks: An overview*. Monet – mobile networks, 1997
- [Had02] D.Hadjari, N.Nouali, *les réseaux Sans fil (Wlan)*, Rap tech, Cerist 2002.
- [Her02] D.Hertrich. *MPEG4 video transmission in wireless LANs: Basic Qos support on the data link layer of 802.11b*. Minor Thesis, Techenical Uinversity Berlin, 2002.
- [Hur03] I.Hurbain, *Optimisation de la bande passante pour le broadcast vidéo sur réseau IP*. Mai 2003
- [Ind98] Indulska J., Cook S., *New RSVP extension to support computer mobility in the internet*, SICON 1998.

- [Ima03] Imad Aad, Qiang Ni, Chadi Barakat, Thierry Turlett. *Enhancing IEEE 802.11 MAC in congested Environments*, INRIA N° 4879 July 2003
- [ITU95] ITU. Recommendation ITU-R BT 500-7. Methodology for subjective Assessment of the quality of Television Pictures. Rep.tech. October 1995.
- [ITU02] ITU-T and ISO/IEC, Draft Text for joint video specification(ITU-T Rec.H264 And ISO/IEC 14496-10 AVC, October 2002.
- [Jon02] Johnson D., Perkins C. *Optimization in Mobile IP*, IETF, Draft-ietf-mobileip-optim-11.txt, September 2001.
- [Jai02] M.Jain C Dovrolis. Pathload : A measurement tool for end to end available Bandwidth . Passive and active Measurements (PAM) Workshop.2002.
- [Jig99] W.Jiang et H.Schulzrinne. Qos measurements of internet Real-time multimedia services. Décembre 1999.
- [Jpg02] ITU-T. { Digital Compression and Continuous-tone Still Images – Requirements and Guidelines. Recommendation T.81, 2002.
- [Joh02] [Joh03] D.Johenson C.Perkins. IETF internet draft : Optimisation in mobile IP, draft-ietf-mobileip-optim-11.txt. Septembre 2002 .
- [Joh03] D.Johenson C.Perkins. IETF internet draft : *Mobility Support in Ipv6*, draft-ietf-mobileip-ipv6-24.txt. June 2003.
- [Koo00] R.Koodli, Perkins C. *A Framework for Smooth Handovers with Mobile IPv6*. IETF, draft - koodli-mobileip-smoothv6-00.txt, Juillet 2000.
- [Koo02] R.Koodli, Perkins c.*Fast Handovers in Mobile IPv6*, IETF, draft-koodli-mobileip-fastv6-06.txt, 2002.
- [Koo03] R.Koodli, Perkins c.*Fast Handovers in Mobile IPv6*, IETF, draft-koodli-mobileip-fastv 6-07.txt, 2002.
- [Kou01] Michael E. Kounavis¹, Andrew T. Campbell¹, Gen Ito, and Giuseppe Bianchi. *Design, Implementation and Evaluation of Programmable Handoff in Mobile Networks*. Comet Group, Center for Telecommunications Research Columbia University, New York,2001
- [Kri03] Santhana Krishnamachari, Mihaela van der Schaar, Sunghyun Choi, Xiaofeng Xu *Video Streaming over Wireless LANs: A Cross-layer Approach*, Packet video, Nantes 2003
- [Lai99] K.lai M.Baker. Measuring Bandwidth. INFCOM, P.235N°245.1999.
- [Lao02] A.laouiti. Unicast et multicast dans les réseaux Ad hoc sans fil.Thèse de doctorat de l'université de Versailles.Juillet 2002.
- [Lia01] Y.Liang ., Eckehard G. Steinbach, and Bernd Girod. *Real-time Voice Communication over the Internet Using Packet Path Diversity*, Information Systems Laboratory, Department of Electrical Engineering Stanford University, Stanford, CA 94305, USA, ACM Multimedia, 2001.
- [Lia03] Y.Liang J.G.Apostolopoulos. *Analyse of packets loss for compressed video: does burst-length matter ?*. ICASSP, April 2003.

- [Lin01] Y.B.Lin, I.Chlamtac. wireless and mobile Network Architecture. John Wiley & Sons,2001.
- [Maj02] A.Madjumdar, D.Sachs, I.Kozinetch , K.Ramchandran, M.Yeung. Multicast and unicast real-time video streaming over wireless LANs. IEEE Trans. On circuits and systems for video technology. June 2002.
- [Mal03] K.ELMalki. Internet task force, internet Draft, Simultaneous bindings for mobile Ipv6 fast handover ,draft-elmalki-mobileip-bicastingv6.03-txt. May 2003.
- [Mal03] K.ELMalki (editor) H.Soliman and S Thalanany, *LOW latency Handoffs in mobile Ipv4*.Internet task force, internet Draft,draft-ietf-mobileip-lowlatency-handoffs-v4-05.txt:work in progress, Julay2003.
- [Mas03] Masala E, C. F. Chiasserini, M. Meo, J. C. De Martin. *REAL-TIME TRANSMISSION OF H.264 VIDEO OVER 802.11-BASED WIRELESS AD HOC NETWORKS*. Dipartimento di Automatica e Informatica, Dipartimento di Elettronica IEIIT-CNR Politecnico di Torino Corso Duca degli Abruzzi, Torino, Italy,2003
- [Mis03] A.Mishra M. Shin W. Arbaugh.*An Empirical Analysis of the IEEE 802.11 MAC Layer Handoff Process*.Arbaugh, in the ACM SIGCOMM Computer Communication Review (ACM CCR), Volume 33, Issue 2, April 2003. Also as a Technical Report, UMD, CS-TR-4395, UMIACS-TR-2002-75.
- [Mon02] N.Montavont, T.Noel. *La mobilité dans les réseaux IP*. Master thesis 2001.
- [Mon02] Nicolas Montavont and Thomas Noël. *Handover Management for Mobile Nodes in IPv6 Networks*. LSIIT - Louis Pasteur University- CNRS, IEEE Communications Magazine, August 2002
- [Mon02] N.Montavont, T.Noel. *Anticipation du handover dans les réseaux sans fil*, IEEE international symposium on advances in wireless communication (ISWC2002) ; Sept 2002
- [Mor02] Moret Y., Bonnet C., Gauthier L., Knopp R., *Process and Apparatus for Improved Communication between a Mobile Node and a Communication Network*. Office Européen des Brevets, n° 02368057.2, Août 2002.
- [Nar98] T.Narten. E. Nordmark. W.Simpson. Neighbor Discovery for Ipv6. RFC 2461.IETF.December 1998.
- [Nit00] H.Nitin, P.Bahl, Seema Gupta. Distributed fair scheduling in a wireless LAN. In proc of 6th ACM Mobicom,Boston, MA, August 2000.
- [NS] Network simulator- NS-2. Available: www.isi.edu/nsnam/ns.
- [Olo02] Jon-Olov Vatn, *Supporting real-time services to mobile Internet hosts*, Dissertation thesis proposal, Dept. of Microelectronics and Information Technology (IMIT), Royal Institute of Technology (KTH), Sweden, June 5, 2002
- [Olo03] Jon-Olov Vatn ,*An experimental study of IEEE 802.11b handover performance and its effect on voice traffic*, Technical Report TRITA-IMIT-TSLAB R 03:01,Telecommunication Systems Laboratory, Department of Microelectronics and Information Technology, KTH, Royal Institute of Technology, Stockholm, Sweden, July 2003.
- [Per96] Perkins C., *IP Encapsulation within IP*, RFC 2003, Octobre 1996.

- [Per99] Perkins C., Wang KY. *Optimized Smooth Handoffs in Mobile IP*. Proceedings of the fourth IEEE Symposium on Computers & Communications, Juillet 1999.
- [Per02] Perkins C. *IP mobility support for IPv4*, IETF, RFC 3220, Janvier 2002.
- [Per02] C.perkins, ED, *IP Mobility support For Ipv4*, internet Engineering Task force, RFC3344, August 2002.
- [Pra91] W.K.Pratt. *Digital Image Processing*, second edition , Chap.3. Wiley-Interscience, 1991.
- [Ric02] Richard J. H. Kok. *Vertical Handover Based Video Streaming*. Bachelor of Information Technology (Honours) School of Information Technology and Electrical Engineering University of Queensland, 2002
- [Sal99] Salamatian Mohemmad. *Transmission multimedia sur internet*. PHD thesis, Universite de Paris – sud, UFR scientifique d’Orsay, 1999.
- [Sch03] Hans Scholten, Pierre Jansen, Ferdy Hanssen, Wietse Mank, Arjan Zwikker. *A Real-Time Multimedia Streaming Protocol for Wireless Networks* .Technical report TR-CTIT-03-37, Centre for Telematics and Information Technology, University of Twente, Enschede, the Netherlands, August 2003.
- [Sol00]Soliman H., Castellucia C., El Malki K., Bellier L., *Hierarchical Mobile IPv6 and Fast Handoffs*, IETF, draft-ietf-mobileip-hmipv6-00.txt, Septembre 2000.
- [Std99] IEEE standard, Part 11: *Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications*, 1999.
- [Tan03] S.Tanenbaum. *Computer Networks 4/e*, Prentice Hall.2003.
- [Tek97] Tektronix.A Guide of Picture quality measurments for modern television Systems, 1997.
- [Tay01] Y.C TAY, K CHUA, *A capacity Analysis for the IEEE 802.11MAC Protocol*, wireless Networks 7,159-171, 2001.
- [Try96] C. Tryfonas. *MPEG-2 Transport over ATM Networks*, 1996
- [Tea 03] Teayreon Park. *Adptive handover control*. Institute for telecommunication reaserch University of south australia.Phd thesis.2003
- [Var02] A.varga, “Omnet++”, IEEE Networks Interactive. Vol.16. No.4. 2002. Availbale : www.hit.bme.hu/phd/vargraa/omnetpp
- [Vel03] Héctor Velayos, *Techniques to Reduce IEEE 802.11b MAC Layer Handover Time* KTH Technical Report TRITA-IMIT-LCN R 03:02, ISSN 1651-7717, ISRN KTH/IMIT/LCN/R-03/02--SE, Stockholm, Sweden. April 2003.
- [Viv04] I.Vivaldi , V.Prakash, *Fast handover Algorithme for hierarchical Mobile Ipv6*, Suranaree J.Sci Technol .Vol.11.No1.january 2004.
- [Wai02] Wai-tian John Apostolopoulos, Tan, Susie Wee. *MODELING PATH DIVERSITY FOR MULTIPLE DESCRIPTION VIDEO COMMUNICATION*, ICASSP 2002, May, 2002
- [Wan99] C.Perkins, K-Y Wang. *Optimised smooth handover in mobile IP*, Proceeding of IEEE international symposium on computers and communication (ISCC’99)

[Wan02] A Wang, J Zhu .*Transmitting video over Wireless LAN*. EE384B, Department of Electrical EngineeringStanford University,2002

[Woo03] Steve Woon, Eric Wu, Ahmet S,ekerciořglu.*A Simulation Model of IEEE802.11b for Performance Analysis of Wireless LAN Protocols*, School of Electrical and Computer Systems EngineeringMonash University, Australia, 2003.

[Xia04] Xiaofeng Xu, Mihaela van der Schaar, Santhana Krishnamachari, Sunghyun Choi, Yao Wang. *Fine-Granular-Scalability Video Streaming over Wireless LANs using Cross Layer Error Control*. IEEE ICASSP'04, Montreal, Canada, May 2004

[You03] Youcef khouadja, Kerine Noel, *Gestion de la mobilité IPv6 contrôle par le réseau* (Annales de télécommunication vol58), 2003.

[Zha01] Zhang T., Chen JC., Agrawal P. *Distributed soft handoff in all-IP wireless networks*, Proceedings of IEEE International Conference on Third Generation Wireless and Beyond (3Gwireless'01), San Francisco, CA, pp. 460-465, mai 2001.

ANNEXE A

1 Network Simulator (NS)

Devant l'évolution de la technologie, de nombreuses solutions sont envisageables pour résoudre un même problème. La simulation permet de tester à moindre coût les nouveaux protocoles et d'anticiper les problèmes qui pourront se poser dans le futur afin d'implémenter la technologie la mieux adaptée aux besoins. NS **[Mon01]** est un simulateur à événements discrets disponible gratuitement sur le site <http://www.isi.edu/nsnam/>. Il permet à l'utilisateur de définir un réseau et de simuler des communications entre les noeuds de ce réseau. NS v2 utilise le langage OTCL (Object Tools Command Language), dérivé objet de TCL. À travers ce langage, l'utilisateur décrit les conditions de la simulation : topologie du réseau, caractéristiques des liens physiques, protocoles utilisés, communications... La simulation doit d'abord être saisie sous forme de fichier texte que NS utilise pour produire un fichier trace contenant les résultats.

NS est fourni avec différents utilitaires dont des générateurs aléatoires et un programme de Visualisation : NAM (Figure 10-1).

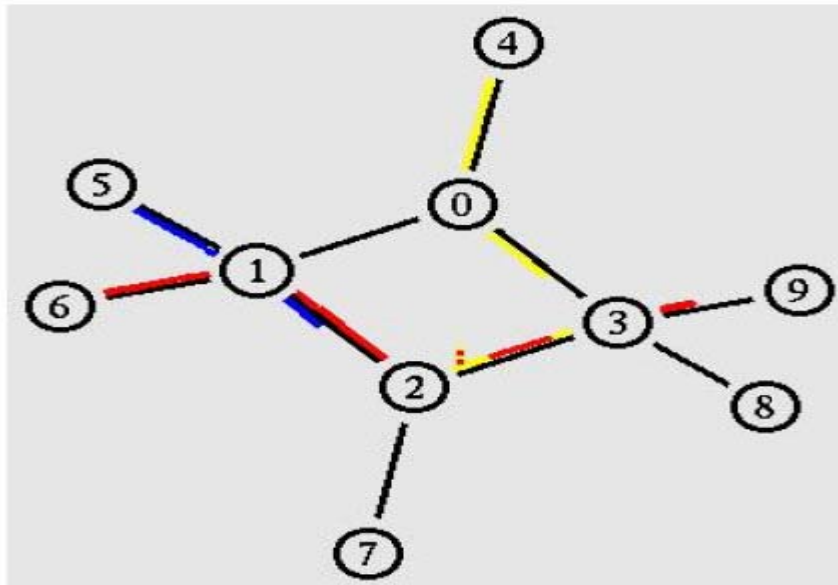


Figure 10-1. NAM

Si OTCL permet de manipuler facilement les données, sa faible vitesse d'exécution ne permet pas de faire tourner des simulations très complexes. Aussi NS utilise un deuxième langage : le C++. Il est possible d'agir sur la partie OTCL en C++ et inversement, les objets étant accessibles aux deux langages, certaines variables pouvant même être communes.

Sous NS, chaque équipement est décrit par un noeud. Les noeuds sont reliés entre eux par des liens. A la base, le noeud n'est composé que de deux éléments : un classifieur d'adresses et un classifieur de ports auquel on ajoute des connexions. Ces connexions peuvent être soit des liens qui manipulent les paquets, soit des agents. Ce sont les agents qui génèrent et traitent les paquets.

Quand un paquet arrive sur un noeud, on commence par regarder son adresse de destination.

Les paquets destinés à ce noeud passent ensuite par le classifieur de port pour être dirigés vers le bon agent. Les autres sont généralement soit jetés soit transmis à un agent de routage qui les envoie ensuite sur un autre lien. Néanmoins ils peuvent aussi être redirigés sur un autre agent comme c'est le cas avec Mobile IP. Lorsqu'un mobile se déplace et quitte la cellule de son Home Agent pour s'enregistrer auprès d'un Foreign Agent, ce dernier crée une association au niveau du classifieur d'adresses entre l'adresse du mobile et l'agent d'encapsulation qui se charge du tunnel IP.

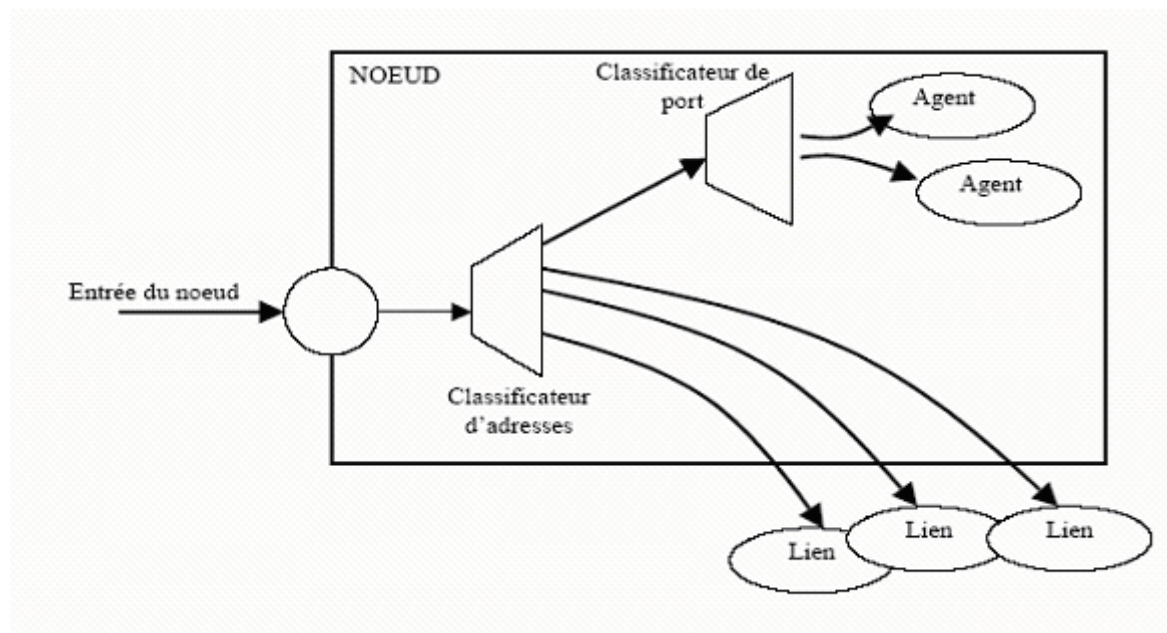


Figure 10.2 Structure d'un noeud mobile

2 Architecture et implémentation

L'architecture réseau de NS-2 est fortement basée sur le modèle des couches OSI. Ce modèle a brièvement été expliqué dans la partie 1. Il s'agit de la décomposition de la pile réseau en couches. Tout au long de cette section on retrouvera donc les éléments de ces couches avec plus ou moins de détails.

La dernière version de NS-2 est sortie le 7 juin 2001 dans sa version 8 (ns-2.1b8). Les sources sont disponibles sur le site ISI dans la section « nsnam » [<http://www.isi.edu/nsnam/ns>]. Les sources se présentent sous deux formes : l'une dite « tout en un » qui contient le code NS-2 et d'autres composants utilisés (comme OTcl, NAM...), soit par morceaux, c'est-à-dire qu'on peut choisir uniquement les composants dont on a besoin. Le package comprend aussi des exemples de script ainsi que des modèles de mouvement pour les noeuds mobiles ou de génération de trafic.

NS-2 est un simulateur à événements à temps discrets orienté objet. Il est développé en C++ et en OTcl21 ; Le package inclus une hiérarchie de classe compilée d'objets écrits en C++ et une hiérarchie de classe interprétée d'objets écrits en OTcl. Ces deux hiérarchies sont étroitement liées ; quand l'utilisateur crée un nouvel objet par l'interpréteur OTcl, un objet correspondant appelé objet reflet est aussi créé dans la hiérarchie compilée. On dit que ces objets sont des « objets fendus ». Bien entendu, les objets peuvent être accédés aussi bien en OTcl qu'en C++ grâce à la mise en place de procédures d'appel entre OTcl et C++.

Le langage OTcl est un langage interprété qui ne demande pas de compilation. Il est principalement utilisé pour concaténer des objets, accéder aux objets à partir de l'interpréteur et configurer des simulations (début et arrêt des événements, perte réseau, rassemblement de statistiques). Son utilisation est rapide et assez conviviale. D'un autre côté, C++ est utilisé pour créer les classes de base et pour traiter un grand nombre de données (tel que calcul des tables de routage, mouvement des mobiles...).

2.1 Implémentation du simulateur

Au plus bas niveau, il y a six classes qui définissent l'ensemble de la structure du programme et fournissent les méthodes élémentaires. Il s'agit des classes *Tcl*, *TclObject*, *TclClass*, *TclCommand*, *EmbeddedTcl*, *InstVar*. Elles définissent entre autres les méthodes utilisées par C++ pour accéder à l'interpréteur, la hiérarchie, les principales commandes de haut niveau et les méthodes pour accéder aux variables C++ et OTcl.

La simulation est configurée, contrôlée et gérée à l'aide des interfaces fournies par la classe OTcl *Simulator*. Cette classe fournit des procédures pour créer et gérer la topologie, initialiser le format des paquets et choisir le planificateur d'événements. Elle stocke intérieurement des références à chaque

élément de la topologie. Un script devra donc toujours commencer par l'instanciation d'une variable de cette classe. L'utilisateur crée ensuite la topologie à travers OTcl en utilisant les classes *node* et *link*, composants essentiels de la topologie. Ces éléments sont décrits dans la sous section suivante

2.1.1 Composants de la topologie

La topologie NS-2 est essentiellement composée de noeuds et de liens. La définition des **noeuds** se fait dans un premier temps à travers l'instance de *Simulator* puis à travers l'instance de la classe *Node*. La fonction d'un noeud est de recevoir des paquets, les examiner et les mapper à ses interfaces sortantes appropriées (voir **Figure 10-2**). Cette classe est composée d'un classificateur et de méthodes pour configurer un noeud. Les méthodes proposées sont des fonctions de contrôle, de gestion d'adresse et de port, de gestion d'agents et de repérage des voisins. Le classificateur est la partie du noeud qui traite chaque segment des paquets reçus. Il en existe donc plusieurs, chacun étant spécifique au champ examiné (ex : le classificateur d'adresse est utilisé pour supporter la propagation des paquets). Nous verrons plus loin l'implémentation de deux noeuds spéciaux (sous-classe de *Node*) qui permette la mise en oeuvre des sous-réseaux locaux de manière simple et la mise en oeuvre de la mobilité.

Les **liens** constituent la deuxième partie de la topologie. Les liens entre les noeuds sont définis dans la classe *Link* et *SimpleLink* plus précisément lorsqu'il s'agit de relier deux noeuds. Plusieurs types de liaisons sont supportés, comme le point à point, le broadcast ou les liaisons sans fil pour la mobilité. Cette classe définit cinq variables clés

Manipulent la file d'attente (enfilement, défilement, destruction et réception). Une liaison est définie par une séquence de connecteurs (classe *Connector*) qui sont rangés dans une file d'attente. Ces connecteurs font suivre les paquets qui leur sont envoyés dans une seule direction ; Le paquet est alors délivré au voisin cible ou il est détruit. A présent, voyons les structures mises en place autour de la topologie pour faire interagir les noeuds entre eux.

2.1.2 La gestion des files d'attente

La gestion des files d'attente et la simulation des délais sur les liens sont implémentés dans les classes *Queue* et *LinkDelay* respectivement. Les files d'attente actuellement disponibles dans NS sont :

- FIFO
- RED buffer management
- CBQ (priorité et circulaire)

- Plusieurs variantes de file d'attente juste (Fair Queue)

Pour simuler un quelconque délai dans la réception ou l'émission d'un paquet, la file d'attente correspondante est simplement bloquée.

2.1.3 Les agents

A un niveau supérieur, on retrouve les agents (*classe Agent*) qui jouent un rôle important dans les simulations. Les utilisateurs créent de nouvelles sources ou récepteurs à partir de la classe *Agent*. Ils font partie intégrante d'un noeud et sont les points terminaux vis à vis des paquets couche réseau ; leur rôle est de générer et réceptionner des paquets suivant les protocoles de transport. La **Figure 10-2** montre les interactions entre ces différents composants dans NS.

La génération de trafic dans NS peut se faire de deux manières différentes et est décrit dans la classe *Application*. Il est possible de générer des paquets par un générateur de trafic (classe *TrafficGenerator*) ou par la simulation d'applications existantes (classe prenant le nom de l'application), toutes ces classes étant dérivées dans la classe *Application*. Les générateurs de trafic peuvent être de quatre types :

- Classe *Exponential* : génère un trafic ON/OFF²² à intervalle de temps régulier.
- Classe *Pareto* : génère un trafic ON/OFF à intervalle de temps aléatoire.
- Classe *CBQ* : débit de bit constant.
- Classe *Trace* : permet de lire la génération de trafic dans un fichier.

Au niveau de la couche physique, l'objet *Channel* simule le médium partagé et supporte des mécanismes d'accès au médium des objets *Mac* (phase d'émission). L'objet *Classifier/Mac* est responsable de la livraison et de la réplique des paquets pour des objets *Mac* (phase réception). Les détections de collisions se font au niveau de la couche MAC où sont implémentés les protocoles d'accès au médium (CSMA...). La couche liaison est composée de deux objets : *Queue* qui simule l'interface de file d'attente et *LinkLayer* qui implémente un protocole de couche de données. Un LAN contient en plus un seul objet *LanRouter* qui est créé automatiquement que le noeud LAN est initialisé. Tous les noeuds d'un LAN ont un pointeur sur cet objet qui joue le rôle de routeur.

2.1.4. La mobilité dans NS

Dans un premier temps, la mobilité a été introduite dans NS-2 par les chercheurs de l'université Carnegie Mellon²⁴ de Pittsburgh (CMU) dans la volonté de simuler des réseaux ad hoc. Le lecteur doit particulièrement faire attention à la terminologie. Dans les deux premières parties, on a parlé de point d'accès comme étant un équipement de niveau 2 fournissant l'accès Internet aux noeuds mobiles. Dans la terminologie NS-2, on parle de point d'accès non seulement pour spécifier un tel noeud mais aussi pour parler d'agent mère ou d'agent visité.

L'apport de la mobilité passe par l'ajout d'un nouveau type de noeuds définis dans la classe *MobileNode*, qui ne sont pas connectés entre eux. Les caractéristiques de la mobilité telles que le mouvement des noeuds, les mises à jour de localisation ou les limites de la topologie sont implémentées en C++. Par contre, les composants réseaux comme le noeud mobile lui-même (classificateur, couche liaison...) sont implémentés en OTcl.

Comme l'objectif était de simuler des réseaux entièrement mobiles, il a fallu mettre en place des protocoles de routage. Actuellement, il y a quatre protocoles de routage mis en oeuvre dans NS-2:

- Séquence de destination à vecteur de distance (DSDV) : des messages de routage sont échangés entre noeuds mobiles voisins. Les mises à jour peuvent être déclenchées (par une information sur un noeud voisin qui change la table de routage) ou périodique. Tout paquet à destination d'un noeud inconnu est mis en buffer pendant que des *Query* sont envoyés.
- Routage par source dynamique (DSR) : utilise un objet particulier de la classe *SRNode*. Toutes les entrées de l'objet *SRNode* pointent sur cet agent de routage. Ce modèle s'avère intéressant pour une future implémentation de protocole utilisant le piggy-backing.
- Algorithme de routage ordonné temporaire (TORA) : protocole de routage distribué basé sur l'algorithme « Link Reversal ». A chaque noeud, une copie séparée de TORA tourne pour chaque destination. TORA opère au-dessus de IMEP²⁵ qui procure la liaison des messages de routage et informe le protocole de routage des changements des liens aux voisins à travers l'émission de beacons.
- Vecteur de distance sur demande (AODV) : combinaison des protocoles DSR pour la découverte et la maintenance des routes et DSDV pour le routage saut par saut, numéro de séquence et l'algorithme des beacons. Lorsqu'un noeud mobile est créé dans une simulation, le simulateur crée un objet *MobileNode*, un agent de routage et la pile réseau qui sera décrite plus loin.

ANNEXE B

1. Les spécifications du modèle de simulation

1) La mise on œuvre des spécifications sans fil couche MAC / PHYSIQUE

```
set val(chan) Channel/WirelessChannel           ;# channel type
set val(prop) Propagation/TwoRayGround          ;# radio-propagation model
set val(netif) Phy/WirelessPhy                  ;# network interface type
set val(mac) Mac/802_11                         ;# MAC type
set val(ifq) Queue/DropTail/PriQueue            ;# interface queue type
set val(ll) LL                                   ;# link layer type
set val(ant) Antenna/OmniAntenna                ;# antenna model
set val(ifqlen 50                               ;# max packet in ifq
clients
set val(rp) DSDV                                ;# routing protocol
$val(mac) set bandwidth_ $val(datarate)         ;# Data rate (to be read from the
Command line arguments)
```

2) La création d'un nouveau simulateur et le fichier trace

```
set ns_ [new Simulator]
$ns_ use-newtrace
set tracefd [open $orig_tr w]
$ns_ trace-all $tracefd
```

3) La création des PAs et la mise des spécifications (le mode PCF).

```
$ns_ node-config -adhocRouting $val(rp) \
-llType $val(ll) \
-macType $val(mac) \
-ifqType $val(ifq) \
-ifqLen $val(ifqlen) \
-antType $val(ant) \
-propType $val(prop) \
-phyType $val(netif) \
-channelType $val(chan) \
```

```

-topoInstance $topo \
-agentTrace ON \
-routerTrace OFF \
-macTrace OFF \
-movementTrace OFF
set AP [$ns_ node]
$AP random-motion 0
$ns_ initial_node_pos $AP 10
# This turns AP into a Point Coordinator
[$AP set mac_(0)] make-pc
# And here we set how long the time between beacons should be
[$AP set mac_(0)] beaconperiod $val(superframe)

```

4) La création du nœud vidéo (Notons que le mode PCF est utilisé et que le nœud vidéo doit être ajouté au « polling list » du PA.)

```

# val(nvideo) is to be read from the command line arguments
for {set i 0} {$i < $val(nvideo)} {incr i} {
  set videonode_($i) [$ns_ node]
  $videonode_($i) random-motion 0 ;# disable random motion
  $ns_ initial_node_pos $videonode_($i) 20
  [$AP set mac_(0)] addSTA [expr $i+1] 1 0
}
# val(ndata) is to be read from the command line arguments
for {set i 0} {$i < $val(ndata)} {incr i} {
  set datanode_($i) [$ns_ node]
  $datanode_($i) random-motion 0 ;# disable random motion
  $ns_ initial_node_pos $datanode_($i) 20
}

```

5) **Traffic vidéo** : Génération de la trace vidéo (séquencece.tr)¹ et son attachement au nœud vidéo.

```

for {set i 0} {$i < $val(nvideo)} {incr i} {
  set videosink($i) [new Agent/Null]
  $ns_ attach-agent $AP $videosink($i)
}

```

```

set videoudp($i) [new Agent/UDP]
$ns_ attach-agent $videonode_($i) $videoudp($i)
$videoudp($i) set packetSize_ 8000
$videoudp($i) set class_ [expr 101+$i]
set tfile [new Tracefile]
$tfile filename "séquence.trc"
set trace_($i) [new Application/Traffic/Trace]
$trace_($i) attach-tracefile $tfile
$trace_($i) attach-agent $videoudp($i)
$ns_ connect $videoudp($i) $videosink($i)
$ns_ at [expr $val(starttime)+0.01*$i] "$trace_($i) start"
}

```

2. Le format d'une séquence YUV

Les schémas suivants présente les formats vidéo utilisés pour le modèle de simulation :

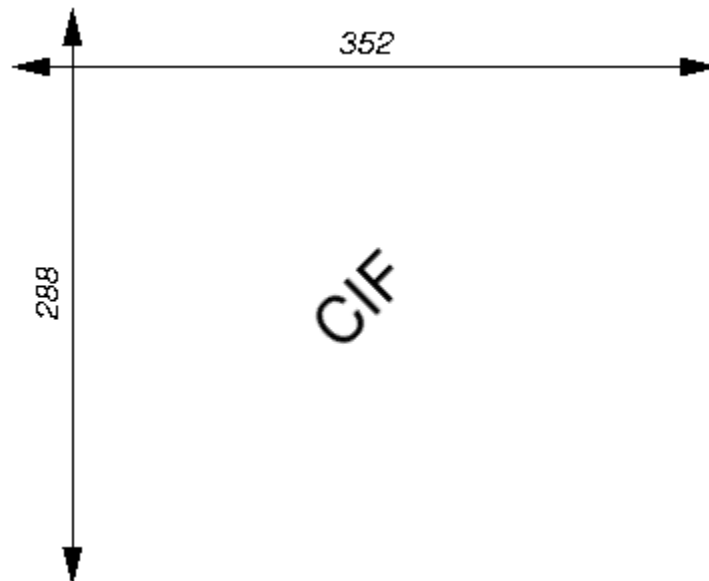


Figure 11.1 .Format CIF(octet)

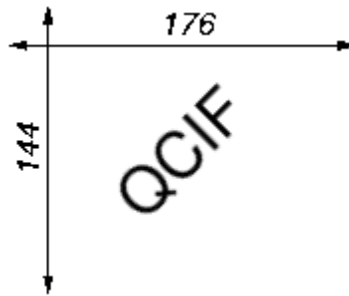


Figure 11.1 .Format QCIF(octet)

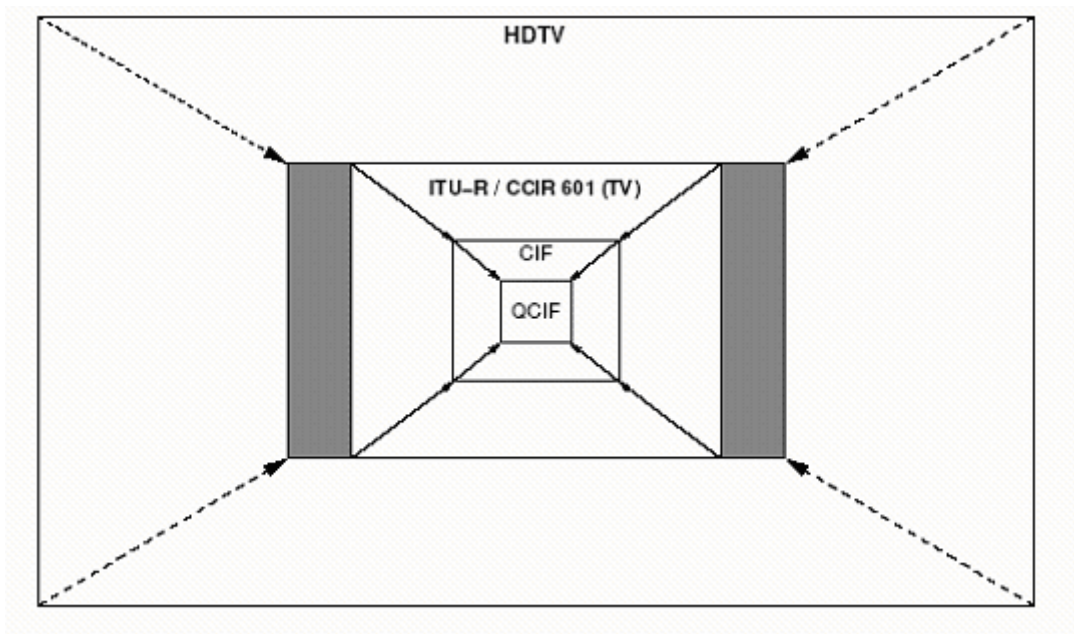


Figure 11.3. Illustration des formats des images (CIF, QCIF)

- YUV 4:2:0 Sequences vidéo

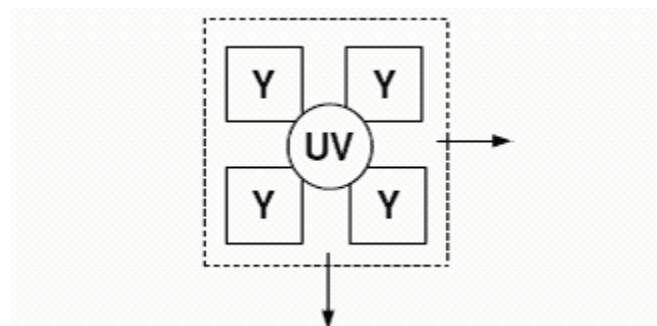


Figure 11.4.YUV 4 :2 :0