

N° d'ordre: 3665

THÈSE

Présentée à

L'UNIVERSITÉ BORDEAUX I
ÉCOLE DOCTORALE DE MATHÉMATIQUES ET
D'INFORMATIQUE

Par **Ismail Adel Djama**

POUR OBTENIR LE GRADE DE

DOCTEUR

SPÉCIALITÉ : INFORMATIQUE

**Adaptations *inter-couches* pour la diffusion des services
vidéo sans fil**

Cross-Layer Adaptations for wireless video streaming services

Soutenue le : 10 Novembre 2008

Devant la commission d'examen composée de :

Richard Castanet	Professeur, ENSEIRB	Président du jury
Samuel Pierre	Professeur, Université de Montréal	Rapporteur
Nazim Agoulmine	Professeur, Université d'Evry Val d'Essonne	Rapporteur
Pascal Lorenz	Professeur, Université of Haute Alsace	Examineur
Francine Krief	Professeur, ENSEIRB	Directrice de Thèse
Toufik Ahmed	Maître de conférences, ENSEIRB	Co-Directeur de Thèse
Dusson Christophe	Ingénieur de Recherche, Orange Labs	Invité

Abstract

One of the big challenges in the convergence of networks and services to the IP technology is to maintain the Quality of service (QoS) for audio/video streams transmitted over wireless networks to heterogeneous mobiles users. In this environment, the multimedia services should face many shortcomings caused mainly by the wireless channel unreliability and its sharing among many users. These shortcomings are increased by the terminals heterogeneity (i.e. decoding capability, memory storage, display resolution, etc.) which should receive, decode and display the multimedia streams.

In order to allow universal access to multimedia services anywhere, anytime and using any kind of terminal, the new generation of multimedia applications have to interact with their environment, on the one hand, to inform the underling network about their need in term of QoS, and on the other hand, to dynamically adapt their services according to the receiver terminal and the changing in network conditions.

In this context, we have proposed a new Cross-layer based streaming system to transmit audio/video streams over 802.11 networks. This new system, called XLAVS (Cross layer Adaptive video streaming), actively communicates with all network layers and the end receiver to determine the optimal adaptation that optimize the QoS of audio/video streams.

Our contributions focus mainly on the Cross-layer adaptations implemented on the XLAVS. These contributions are classified into two major categories: the bottom-up adaptations performed at the application level and the top-down adaptations performed at the 802.11 MAC level.

In the first category, our proposals have revolved around: (1) a dynamic adaptation of video throughput according to the physical rate available in the 802.11 network and (2) a join FEC and video throughput adaptation steered by the signal strength and the loss ratio.

In the second category, we have proposed two Cross-layer mechanisms at the 802.11 MAC level: (1) an adaptive 802.11 MAC fragmentation to find an optimal trade-off between the packet loss and the overhead introduced by the MAC layer and (2) a video frame grouping at MAC level that allows video stream to get access to the 802.11 channel proportionally to its throughput.

Keywords: *Cross-layer* adaptation, Quality of Service, 802.11 wireless networks, DVB-T networks, IPTV service, video adaptation, adaptive FEC.

Résumé

L'un des défis majeurs dans la convergence des réseaux et des services vers la technologie IP est le maintien de la qualité de service (QoS) des flux audio/vidéo transmis sur des réseaux sans fil pour des utilisateurs mobiles et hétérogènes. Dans cet environnement, les services multimédia doivent faire face à plusieurs inconvénients engendrés par le manque de fiabilité d'un canal sans fil et son partage par plusieurs utilisateurs. Ces inconvénients sont accentués par l'hétérogénéité des terminaux de réception (capacité de décodage, espace de stockage, résolution d'affichage, etc.) qui doivent recevoir, décoder et afficher les flux multimédia.

Afin d'assurer un accès universel aux services n'importe où, n'importe quand et en utilisant n'importe quel terminal d'accès, les applications multimédia de nouvelle génération doivent interagir avec leur environnement pour, d'une part, informer les réseaux sous-jacents de leur besoins en QoS, et d'autre part, adapter dynamiquement leurs services en fonction des terminaux de réception et des variations intempestives des conditions de transmission.

Dans ce contexte, nous proposons un nouveau système pour la transmission des flux audio/vidéo sur les réseaux 802.11 basé sur l'approche *Cross-layer*. Ce nouveau système, appelé XLAVS (Cross Layer Adaptive Video Streaming), communique activement avec l'ensemble des couches réseaux ainsi que le récepteur final pour déterminer l'adaptation optimale qui permet d'optimiser la QoS des flux audio/vidéo.

Nos contributions se focalisent principalement sur les adaptations *Cross-layer* mises en œuvre par le XLAVS. Ces contributions sont organisées en deux grandes catégories : les adaptations ascendantes exécutées au niveau applicatif et les adaptations descendantes exécutées au niveau MAC 802.11.

Dans la première catégorie, notre apport s'articule au tour de : (1) l'adaptation dynamique du débit vidéo en fonction du débit physique disponible dans le réseau 802.11 et (2) l'adaptation conjointe du taux de redondance FEC et du débit vidéo contrôlée par la puissance du signal et les taux de perte.

Dans la deuxième catégorie, nous proposons deux mécanismes *Cross-layer* au niveau MAC 802.11 : (1) une fragmentation 802.11 adaptative pour trouver un compromis entre les pertes de paquets et l'*overhead* introduit par les couches 802.11 et (2) un groupage des images vidéo au niveau MAC pour permettre au flux vidéo d'avoir un accès au canal 802.11 proportionnel à son débit.

Mot clés: L'adaptation *Cross-layer*, la qualité de service, les réseaux sans fil 802.11, les réseaux DVB-T, la convergence des réseaux, le service IPTV, l'adaptation vidéo, la FEC adaptative.

Tables des matières

Abstract	iii
Résumé	iv
Tables des matières	v
Liste des figures.....	ix
Liste des tableaux.....	xii
Remerciements	xiv
Liste des Publications	xv
Chapitre 1 Introduction	1
Chapitre 2 L'état de l'art.....	5
2.1 Introduction.....	5
2.2 Les limites de l'architecture en couches pour intégrer les réseaux IEEE 802.11 et transporter les flux multimédia.....	6
2.2.1 La couche Physique IEEE 802.11	6
2.2.1.1 Les standards IEEE 802.11	6
2.2.1.2 L'évanouissement du signal dans un canal 802.11	8
2.2.1.3 Les mécanismes pour palier l'évanouissement du signal	9
2.2.1.4 Discussion	10
2.2.2 La couche d'accès IEEE 802.11	10
2.2.2.1 L'accès au canal dans la couche MAC 802.11.....	10
2.2.2.2 La livraison fiable dans la couche MAC 802.11	12
2.2.2.3 La QoS au niveau de la couche MAC 802.11	13
2.2.2.4 Le standard IEEE 802.11e	15
2.2.2.5 Discussion	16
2.2.3 La couche réseau	17
2.2.3.1 Le service garanti (IntServ : Integrated services).....	18
2.2.3.2 La différenciation de service au niveau IP (Differentiated services) :	19
2.2.3.3 L'architecture MPLS	21
2.2.3.4 La négociation de la QoS de bout-en-bout.....	22
2.2.3.5 Discussion	24
2.2.4 La couche transport	24
2.2.4.1 Les algorithmes de contrôle de congestion pour les flux multimédia	26
2.2.4.2 Le protocole SCTP	28
2.2.4.3 Le protocole DCCP.....	29
2.2.4.4 Discussion	30

2.2.5	La couche Application.....	31
2.2.5.1	Les protocoles multimédia.....	32
2.2.5.2	Les standards de codage vidéo et le codage vidéo hiérarchique.....	33
2.2.5.3	La QoS applicative pour la transmission vidéo	34
2.2.5.4	Discussion	38
2.3	Les architectures <i>Cross-layer</i> pour les réseaux sans fil.....	39
2.3.1	Le concept du <i>Cross-layer</i>	39
2.3.2	La communication dans les architectures <i>Cross-layer</i>	40
2.3.2.1	Communication directe entre les couches.....	40
2.3.2.2	Communication via une base de données partagée	41
2.3.3	Les approches du <i>Cross-layer</i> dans les réseaux sans fil.....	41
2.3.3.1	Les approches ascendantes (Bottom-up).....	42
2.3.3.2	Les approches descendantes (Top-down).....	44
2.3.3.3	Les approches mixtes (Integrated).....	45
2.3.4	Discussion.....	49
2.3.5	Les projets Européens traitant la problématique <i>Cross-layer</i>	50
2.3.5.1	Le projet 4MORE	50
2.3.5.2	Le projet PHOENIX.....	50
2.3.5.3	Le projet ENTHRONE II	50
2.4	Conclusion	51
Chapitre 3 Proposition d'une architecture <i>Cross-layer</i> pour les services IPTV sans fil		52
3.1	Introduction.....	52
3.2	Contexte général.....	53
3.2.1	La convergence des réseaux de diffusion, de télécommunications et de données ..	53
3.2.1.1	Les réseaux de nouvelle génération (NGN : Next Generation Network)	54
3.2.1.2	L'architecture IMS (Internet Multimedia Subsystem)	55
3.2.2	Le concept UMA (Universal Multimedia Access).....	55
3.2.2.1	Les techniques d'adaptation.....	57
3.2.2.2	Les approches d'adaptation	57
3.2.2.3	L'adaptation MPEG-21	58
3.3	Motivations pour interconnecter les réseaux DVB-T et les réseaux IP/802.11	61
3.4	L'architecture d'interconnexion du réseau DVB-T aux réseaux IP/802.11 WLAN ...	63
3.4.1	Les services IPTV.....	64
3.4.2	Les réseaux DVB-T (Digital Video Broadcast- Terrestrial)	66
3.4.2.1	Le MPEG-2 Transport Stream	67
3.4.2.2	Les tables de signalisation DVB.....	67
3.4.3	L'interconnexion entre le réseau DVB-T et le cœur du réseau IP	68
3.4.3.1	Le TVSP (TV Stream Provider).....	68
3.4.3.2	La transmission multicast dans le cœur du réseau	70
3.4.4	L'interconnexion entre le cœur du réseau IP et le réseau d'accès 802.11.....	71
3.4.4.1	La passerelle d'adaptation XLAG.....	71
3.4.4.2	L'architecture du XLAVS (Cross Layer Adaptive Video Streaming)	73

3.4.4.3	L'architecture du XLDP (Cross Layer Decision Point).....	75
3.4.4.4	Le fonctionnement du XLAVS et les adaptations <i>Cross-layer</i>	76
3.4.4.5	L'adaptation au début de la session.....	77
3.4.4.6	L'adaptation au cours de la session	79
3.4.5	Mise en oeuvre de l'architecture.....	79
3.5	Conclusion	81
Chapitre 4 Les adaptations <i>Cross-layer</i> ascendantes exécutées au niveau applicatif.....		82
4.1	Introduction.....	82
4.2	L'adaptation du débit vidéo en fonction du débit physique 802.11	83
4.2.1	Motivation	83
4.2.2	Les algorithmes de contrôle du débit physique 802.11	84
4.2.2.1	Les RCAs basés sur des statistiques (Statistics-based)	84
4.2.2.2	Les RCAs basés sur le SNR (SNR-based).....	85
4.2.2.3	Les RCAs hybrides (Hybrid)	86
4.2.3	L'adaptation dynamique du débit vidéo au cours de la transmission.....	87
4.2.3.1	Le codage vidéo MPEG.....	87
4.2.3.2	Techniques d'adaptation dynamique du débit vidéo	88
4.2.4	L'adaptation du débit vidéo en fonction du débit Physique 802.11.....	91
4.2.4.1	L'estimation du débit applicatif à partir du débit physique	92
4.2.4.2	Mise en œuvre de l'adaptation du débit au niveau du XLAG.....	95
4.2.5	Tests et évaluations de performance	96
4.2.5.1	L'adaptation en utilisant différents RCAs	97
4.2.5.2	L'adaptation en utilisant plusieurs stations.....	101
4.3	Adaptation conjointe de la FEC et du débit vidéo en fonction de la puissance du signal et des taux de perte.....	104
4.3.1	Motivation	104
4.3.2	Le mécanisme FEC au niveau paquet.....	105
4.3.2.1	Le code XOR (eXclusive OR)	106
4.3.2.2	Le code RS (Reed Salomon)	106
4.3.2.3	Le code LDPC (Low Density Parity Check)	107
4.3.3	L'adaptation conjointe du débit vidéo et du taux de redondance FEC.....	107
4.3.3.1	Mise en œuvre de l'adaptation au niveau du XLAG	110
4.3.4	Tests et évaluations de performance	111
4.3.4.1	Évaluation du mécanisme FEC adaptatif.....	113
4.4	Conclusion	116
Chapitre 5 Les adaptations <i>Cross-layer</i> descendantes exécutées au niveau MAC		
802.11		118
5.1	Introduction.....	118
5.2	La fragmentation adaptative au niveau MAC 802.11	119
5.2.1	Motivation	119
5.2.2	Comparaison entre les niveaux de fragmentation au niveau du XLAG.....	120

5.2.2.1	La fragmentation au niveau applicatif.....	120
5.2.2.2	La fragmentation au niveau réseau.....	121
5.2.2.3	La fragmentation au niveau MAC 802.11	121
5.2.2.4	La plateforme d'expérimentation.....	122
5.2.2.5	Tests et évaluations de performance.....	123
5.2.3	La fragmentation adaptative <i>Cross-layer</i> au niveau MAC 802.11	127
5.2.3.1	Déterminer la taille idéale d'une trame	127
5.2.3.2	Déterminer la limite inférieure pour la taille d'une trame.....	129
5.2.3.3	Le fonctionnement de la fragmentation adaptative au niveau du XLAVS.....	130
5.2.3.4	Tests et évaluations de performance	131
5.3	Le groupage de trames basé sur l'image vidéo	135
5.3.1	Motivation	135
5.3.2	Le principe du groupage de trames (frame grouping)	135
5.3.3	Le principe du groupage de trames MAC basé sur l'image vidéo	137
5.3.3.1	Mise en œuvre du VFG au niveau du XLAG.....	139
5.3.4	Tests et évaluations de performance	140
5.3.4.1	Évaluation du VFG par rapport à un groupage fixe	140
5.3.4.2	La capacité du VFG à assurer la bande passante	146
5.4	Conclusion	148
Chapitre 6	Conclusion et Perspectives.....	150
6.1	Principales contributions	151
6.2	Perspectives et nouveaux défis	152
Références		154
A	Annexe	163
A.1	L'en-tête de la couche PHY 802.11	163
A.2	L'en-tête de la couche MAC 802.11.....	163
A.3	L'en-tête d'un paquet TS	164
A.4	L'en-tête du protocole RTP	165
Liste des abréviations		167

Liste des figures

Figure 1-1 : Les dimensions de la QoS dans les NGNs	2
Figure 2-1 : L'accès distribué au canal DCF	11
Figure 2-2 : L'accès centralisé au canal PCF	12
Figure 2-3 : Taxonomie des mécanismes de QoS au niveau MAC 802.11	13
Figure 2-4 : Les catégories d'accès EDCA	15
Figure 2-5 : L'architecture DiffServ	20
Figure 2-6 : Le conditionnement du trafic DiffServ	20
Figure 2-7 : Taxonomie des algorithmes de contrôle de congestion	26
Figure 2-8 : Taxonomie de la QoS au niveau applicatif	35
Figure 2-9 : Les modèles de communication <i>Cross-layer</i>	40
Figure 2-10 : Les approches du <i>Cross-layer</i>	41
Figure 3-1 : L'architecture NGN	54
Figure 3-2 : L'architecture IMS	55
Figure 3-3 : L'architecture considérée pour le concept UMA	56
Figure 3-4 : La localisation de l'adaptation	58
Figure 3-5 : L'architecture de MPEG-21 DIA	59
Figure 3-6 : L'architecture d'interconnexion des réseaux DVB-T et des réseaux IP/802.11	63
Figure 3-7 : L'architecture fonctionnelle du TVSP	69
Figure 3-8 : La communication multicast entre le TVSP et le XLAG	70
Figure 3-9 : L'architecture du XLAG	72
Figure 3-10 : L'architecture du XLAVS	73
Figure 3-11 : L'architecture du XLDP	75
Figure 3-12 : Format des caractéristiques utilisateurs	77
Figure 3-13 : Format des capacités du terminal	77
Figure 3-14 : Format des caractéristiques réseaux	78
Figure 3-15 : La signalisation RTSP entre le XLAG et le récepteur	79
Figure 3-16 : Les entités du prototype	80
Figure 4-1 : Taxonomie des RCAs	84
Figure 4-2 : Le débit physique en fonction du SNR pour le standard IEEE 802.11a et IEEE 802.11b	86
Figure 4-3 : Les types d'images dans le codage MPEG	87
Figure 4-4 : Le codage spatial dans la norme MPEG	88
Figure 4-5 : La variation du débit vidéo suivant la variation de la résolution temporelle	89
Figure 4-6 : Le débit vidéo moyen en fonction du nombre d'images par seconde	90
Figure 4-7 : La variation du débit vidéo en fonction de la résolution SNR	91
Figure 4-8 : Captures d'images vidéo codées avec différents débits	91
Figure 4-9 : L'interaction entre les modules au niveau du XLAG pour l'adaptation du débit vidéo en fonction du débit physique	95

Figure 4-10 : La plateforme d'expérimentation	97
Figure 4-11 : la variation des débits physiques dans les scénarios 1 et 2	98
Figure 4-12 : Le débit d'émission des flux vidéo dans les scénarios 1 et 2	99
Figure 4-13 : Les taux de perte des flux vidéo pour les scénarios 1 et 2	99
Figure 4-14 : Le débit de réception des flux vidéo dans les scénarios 1 et 2	99
Figure 4-15 : Les taux moyens de perte des flux vidéo pour dix tests dans les scénarios 1 et 2	100
Figure 4-16 : PSNR des flux vidéo dans les scénarios 1 et 2	101
Figure 4-17 : SSIM des flux vidéo dans les scénarios 1 et 2	101
Figure 4-18 : Le débit d'émission des flux vidéo dans les scénarios 1 et 2	102
Figure 4-19 : Les taux de perte des flux vidéo pour les scénarios 1 et 2	102
Figure 4-20 : Le débit de réception des flux vidéo dans les scénarios 1 et 2	102
Figure 4-21 : PSNR des flux vidéo dans les scénarios 1 et 2	103
Figure 4-22 : SSIM des flux vidéo dans les scénarios 1 et 2	103
Figure 4-23 : Le code correcteur XOR à une dimension	106
Figure 4-24 : Le code correcteur XOR à deux dimensions	106
Figure 4-25 : Le code correcteur RS	106
Figure 4-26 : Le code correcteur LDPC	107
Figure 4-27 : Les différents états du système FEC adaptatif	109
Figure 4-28 : L'adaptation conjointe du taux de redondance FEC et du débit vidéo	110
Figure 4-29 : L'interaction entre les modules au niveau du XLAG pour l'adaptation conjointe du mécanisme FEC et du débit vidéo	110
Figure 4-30 : La puissance du signal perçue simultanément au niveau du XLAG et au niveau de la station 1	112
Figure 4-31 : La puissance et la qualité du signal perçues au niveau du XLAG	112
Figure 4-32 : La variation du RSSI pour les scénarios 1, 2 et 3	114
Figure 4-33 : Les taux de perte pour les scénarios 1, 2 et 3 après le décodage FEC	114
Figure 4-34 : Le débit d'émission des flux vidéo pour les scénarios 1, 2 et 3	114
Figure 4-35 : Les taux moyens de perte des flux vidéo pour dix tests dans les scénarios 1, 2 et 3	115
Figure 4-36 : PSNR des flux vidéo pour les scénarios 1, 2 et 3	116
Figure 4-37 : SSIM des flux vidéo pour les scénarios 1, 2 et 3	116
Figure 5-1 : L'envoi en rafale des fragments dans la fragmentation 802.11	122
Figure 5-2 : La plateforme d'expérimentation	122
Figure 5-3 : Les taux de perte RTCP d'un test dans chaque scénario	124
Figure 5-4 : Le débit d'émission du flux vidéo au niveau MAC d'un test dans chaque scénario	124
Figure 5-5 : Le débit moyen d'émission des dix tests dans chaque scénario	125
Figure 5-6 : Le taux moyen de perte des dix tests dans chaque scénario	125
Figure 5-7 : Le BER en fonction du SNR	127
Figure 5-8 : Le taux de perte et le taux d' <i>overhead</i> en fonction de la taille d'une trame	128
Figure 5-9 : L'échange d'informations pour la fragmentation 802.11 Adaptive	130
Figure 5-10 : Les taux de perte RTCP d'un test dans les scénarios 1 et 2	132
Figure 5-11 : Le débit d'émission du flux vidéo au niveau MAC d'un test dans les scénarios 1 et 2	133
Figure 5-12 : La variation de la taille des trames d'un test dans le scénario 2	133

Figure 5-13 : PSNR pour les scénarios 1 et 2 entre l'image #1000 et l'image #5000	133
Figure 5-14 : SSIM pour les scénarios 1 et 2 entre l'image #1000 et l'image #5000.....	133
Figure 5-15 : Le débit moyen d'émission des dix tests dans les scénarios 1 et 2	134
Figure 5-16 : Le taux moyen de perte des dix tests dans les scénarios 1 et 2	134
Figure 5-17 : Le groupage de trames comparé à une transmission standard	136
Figure 5-18 : Le groupage des images vidéo	139
Figure 5-19 : La plateforme d'expérimentation pour le groupage des images vidéo.....	140
Figure 5-20 : Les débits moyens de réception du flux vidéo1 pour les tests 1, 2, 3, 4, 5.....	143
Figure 5-21 : Les débits moyens de réception du flux UDP pour les tests 1, 2, 3, 4, 5	143
Figure 5-22 : Les taux moyens de perte du flux vidéo1 pour les tests 1, 2, 3, 4, 5	143
Figure 5-23 : Les taux moyens de perte du flux UDP pour les tests 1, 2, 3, 4, 5	143
Figure 5-24 : Les débits moyens de réception du flux video2 pour les tests 6, 7, 8, 9.....	144
Figure 5-25 : Les débits moyens de réception du flux UDP pour les tests 6, 7, 8, 9.....	144
Figure 5-26 : Les taux moyens de perte du flux video2 pour les tests 6, 7, 8, 9.....	144
Figure 5-27 : Les taux moyens de perte du flux UDP pour les tests 6, 7, 8, 9	144
Figure 5-28 : Les débits moyens de réception du flux vidéo1 pour les tests 1, 2, 3, 4, 5.....	145
Figure 5-29 : Les débits moyens de réception du flux TCP pour les tests 1, 2, 3, 4, 5	145
Figure 5-30 : Les taux moyens de perte du flux vidéo1 pour les tests 1, 2, 3, 4, 5	145
Figure 5-31 : Les débits moyens de réception du flux vidéo1 pour les tests 6, 7, 8, 9	145
Figure 5-32 : Les débits moyens de réception du flux TCP pour les tests 6, 7, 8, 9.....	145
Figure 5-33 : Les taux moyens de perte du flux vidéo1 pour les tests 6, 7, 8, 9.....	145
Figure 5-34 : Les débits moyens de réception du flux vidéo pour les scénarios 1 et 2.....	147
Figure 5-35 : Les débits moyens de réception du flux UDP pour les scénarios 1 et 2	147
Figure 5-36 : Les taux moyens de perte du flux vidéo pour les scénarios 1 et 2.....	147
Figure 5-37 : les taux moyens de perte du flux UDP	147
Figure A-1 : L'en-tête physique 802.11	163
Figure A-2 : L'en-tête MAC 802.11	163
Figure A-3 : L'en-tête d'un paquet TS.....	164
Figure A-4 : L'en-tête du protocole RTP	165

Liste des tableaux

Tableau 2-1 : Les débits fournis par les standards IEEE 802.11	7
Tableau 4-1 : Les caractéristiques techniques des flux vidéo	89
Tableau 4-2 : Les débits applicatifs réalisés avec différents débits physiques IEEE 802.11	92
Tableau 4-3 : Les débits applicatifs réalisés avec différentes combinaisons de débits physiques	93
Tableau 4-4 : Les débits physiques effectifs calculés pour différentes combinaisons de débits physiques.....	94
Tableau 4-5 : La politique d'adaptation dans le scénario 2	113
Tableau 4-6 : La politique d'adaptation dans le scénario 3	113
Tableau 5-1 : Les caractéristiques techniques des réseaux 802.11	123
Tableau 5-2 : Les caractéristiques de la vidéo de référence	123
Tableau 5-3 : La politique d'adaptation pour la fragmentation adaptative.....	132
Tableau 5-4 : Les valeurs par défaut du TXOP dans le IEEE 802.11e.....	137
Tableau 5-5 : Le nombre de paquets RTP moyen par image vidéo	138
Tableau 5-6 : les caractéristiques des flux vidéo	141
Tableau 5-7 : Les caractéristiques des tests	142
Tableau 5-8 : Les caractéristiques du flux vidéo	146
Tableau 5-9 : Le débit du flux UDP pour les différents tests	146
Tableau 5-10 : Le débit moyen de réception et le taux moyen de perte avec un flux concurrent TCP	148
Tableau A-1 : L'utilisation des champs « Address 1/2/3/4 » dans l'en-tête MAC 802.11	164

À mes chers et tendres parents, Mohamed et Khadidja, sans vous, rien n'aurait pu être possible, que dieu vous garde pour moi et vous prête une longue vie pleine de santé et de prospérité.

À mes sœurs, Sibem et Amel, et mes frères, Oussem et Walid, merci de m'avoir soutenu tout au long de cette aventure, que dieu vous préserve.

À mes grands parents, Hocine, que dieu te prête longue vie, Hcene et Zina, que dieu vous accepte dans son paradis éternel.

À tous les membres de ma famille et à tous mes amis.

Remerciements

Je tiens à exprimer mes plus sincères remerciements et ma profonde gratitude à Madame Francine Krief et à Monsieur Toufik Ahmed pour m'avoir proposé cette thèse et m'avoir accompagné tout au long de ces trois années. Leurs aides, leurs encouragements et leurs disponibilités ont été déterminants pour la réalisation de cette thèse.

Je voudrais remercier aussi tous les membres de l'équipe COMET et tous le personnel administratif du LaBRI.

J'adresse également mes plus sincères remerciements aux professeurs Richard Castanet, Samuel Pierre, Nazim Agoulmine, Pascal Lorenz qui m'ont fait l'honneur d'évaluer les travaux de cette thèse et de participer au jury de soutenance.

Cette liste n'étant pas exhaustive, je remercie tous ceux qui, de près ou de loin, ont contribué à la réalisation de ce travail, par leur aide, leur encouragement ou par leur simple présence.

Liste des Publications

Journaux et chapitres de livre

- **Ismail DJAMA**, « *Chapitre 7. L'adaptation Cross-Layer pour les services multimédias dans les systèmes embarqués communicants de type 802.11* », Dans le livre « *Les systèmes embarqués communicants : mobilité, sécurité, autonomie* », Auteur : Francine Krief, Traité IC2, série réseaux et télécoms, Hermès, septembre 2008.
- **Ismail DJAMA**, T. Ahmed, H. Nafaa and R. Boutaba, “*Meet In the Middle Cross-Layer Adaptation for Audiovisual Content Delivery*” in **the IEEE Transactions on Multimedia**, ISSN: 1520-9210, Volume: 10, Issue: 1, On page(s): 105 - 120, January 2008.
- **Ismail DJAMA** and T. Ahmed, “*A Cross-Layer Interworking of DVB-T and WLAN for Mobile IPTV Service Delivery*”, in **IEEE Transactions on Broadcasting**, ISSN 0018-9316, pp. 382-390, Volume: 53 Issue: 1 Part: 2, March 2007.

Conférences internationales avec actes

- **Ismail DJAMA**, Toufik Ahmed, "MAC-Level Video Frame Grouping using Cross-Layer Architecture", In Proc of **the 3rd International Symposium on Wireless Pervasive Computing, ISWPC 2008**, On page(s): 261-265, santorini, greece, May 2008.
- **Ismail DJAMA**, Toufik Ahmed, Daniel Négro, "Adaptive Cross-Layer Fragmentation for QoS-based Wireless IPTV Services", In proc of **the 6th IEEE/ACS International Conference on Computer Systems and Applications**, on page(s) : 993-998, Doha Qatar, avril 2008.
- T. Ahmed, **Ismail DJAMA**, A. Nafaa, P. Gélard, C. Donny, and M. Jérôme, “*Towards Efficient IPTV Multicast Support over Satellite Networks Using an IMS-based Architecture*”, In Proc of **the 57th Annual IEEE Broadcast Symposium**, Washington, D.C., USA, October 2007.
- T.Ahmed, **Ismail DJAMA**, A.Nafaa, P.Gélard, C.Donny, M.Jérôme, “*IMS-based IP Multicast Service Delivery over Satellite Network*”, In Procs of **the 18th Annual IEEE International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC'07)**, ISBN: 978-1-4244-1144-3, On page(s): 1-5, Athens Greece, September 2007.
- **Ismail DJAMA**, T.Ahmed, “*Adaptive Cross-layer Fragmentation for Reliable Wireless IPTV Services*”, In Procs of **the IEEE Global Information Infrastructure Symposium (GIIS 2007)**, ISBN: 978-1-4244-1376-8, On page(s): 203-206, Marrakech Morocco, July 2007.
- **Ismail DJAMA**, T.Ahmed, “*MPEG-21 Cross-Layer QoS Adaptation for Mobile IPTV Services*”, in Procs of **the 2nd international Conference on automated production of cross media content for multi-channel distribution (AXMEDIS 2006)**, ISBN: 888453526-3, pp.215-222, Leeds England, December 2006.
- **Ismail DJAMA**, T.Ahmed, D.Negro, and N.Zangar, “*An MPEG-21-enabled Video Adaptation Engine for Universal IPTV Acces*”, In Proc of **the IEEE International Symposium on Broadband Multimedia Systems and Broadcasting 2006**, Las Vegas, NVUSA, April 2006.

- T.Ahmed, **Ismail DJAMA**, “*Delivering audiovisual content with MPEG-21-enabled cross layer QoS adaptation*”, in Journal Zhejiang Univ SCIENCE A, ISSN 1862-1775, pp. 784 793, Vol. 7 No. 5, the **15th International Packet Video Workshop (PV2006)**, April 2006.

Rapports de recherches

- Ingo Kofler, Christian Timmerer, Hamid Asgari, **Ismail DJAMA**, Toufik Ahmed, "*D05 : MPEG-21-based Cross-Layer Adaptation Decision-Taking Engine*", Livrable WP3 Tasks 3.3, ENTHRONÉ 2, Jan 2008.
- **Ismail DJAMA**, T.Ahmed, A.Nafaa, "*D2 : Implémentation, test et évaluation du démonstrateur de services*", consultation DCT/RF/ST 2007-10894, Livrable, février 2008.
- T.Ahmed, **Ismail DJAMA**, A.Nafaa, "*D1 : Architecture Cross-Layer (XL) sur satellite*", consultation DCT/RF/ST 2007-10894, Livrable, Novembre 2007.
- **Ismail DJAMA**, T.Ahmed, A.Nafaa, "*D3 : Implémentation, Test et Évaluation du Démonstrateur de Services*", consultation CNES N° 06 – 15621, Livrable, Juin 2007.
- T.Ahmed, **Ismail DJAMA**, A.Nafaa, "*D2 : Définition du Démonstrateur de Services*", consultation CNES N° 06 – 15621, Livrable, Février 2007.
- T.Ahmed, **Ismail DJAMA**, A.Nafaa, "*D1 : Nouvelle Architecture pour les Communications de Groupes par Satellite*", consultation CNES N° 06 – 15621, Livrable, decembre 2006.

Chapitre 1

Introduction

Après la révolution agricole et la révolution industrielle, l'humanité est à l'aube de la révolution informationnelle qui est en train de transformer radicalement la vie des êtres humains dans toutes ses dimensions. Chaque révolution se caractérise par un instrument de pouvoir : la terre pour la première, le capital pour la seconde et l'information pour la dernière. La révolution informationnelle est due principalement au développement fulgurant des technologies de l'information qui permettent actuellement le traitement, le stockage et la transmission d'énormes quantités de données quasiment en temps réel. Cette nouvelle puissance bouleverse l'ordre mondial dans plusieurs domaines : social, économique, culturel, etc. Les différents moyens de télécommunications ont permis la création de nouveaux marchés mondiaux pour les marchandises, les services et les monnaies, sans considérer les frontières des nations ni les distances qui les séparent. Ces moyens correspondent principalement aux trois grands réseaux déployés à l'échelle mondiale : le réseau de données Internet, le réseau téléphonique et le réseau de diffusion TV. Jusqu'à présent, la spécialisation était la caractéristique principale de ces moyens de télécommunications puisque chaque réseau permettait l'accès à un service particulier à travers une infrastructure particulière. Ceci obligeait les utilisateurs à s'abonner aux trois réseaux et à utiliser un terminal spécifique pour accéder aux services transmis par chacun d'eux. Cette situation commence à changer graduellement grâce aux évolutions technologiques réalisées durant ces dernières années. En effet, la numérisation des données et l'élaboration de nouveaux standards et de nouveaux composants électroniques de plus en plus petits et de plus en plus performants, ont permis aux réseaux de devenir numériques, sans fil et mobiles. Aussi, l'augmentation des capacités de transmission a permis la multiplication et la duplication des services sur toutes les infrastructures. Par conséquent, les barrières qui séparaient auparavant les réseaux de télécommunications commencent à céder les uns après les autres et les terminaux d'accès deviennent multimédia et multiservices intégrant plusieurs interfaces de communication leur permettant de se connecter à différents réseaux. Actuellement, les acteurs des télécommunications s'acheminent vers la notion de convergence qui regroupera tous les réseaux et tous les services sous une seule infrastructure censée représenter le réseau de nouvelle génération NGN. Cette nouvelle architecture permettra aux fournisseurs de services de réduire les coûts de développement en réutilisant des composants communs aux services et elle permettra aux fournisseurs réseaux de réduire les coûts d'exploitation en offrant plusieurs services avec une seule infrastructure. Les différents acteurs s'accordent à dire que la technologie IP sera la brique de base pour bâtir les NGNs. En effet, la simplicité et la

puissance du protocole IP, démontrées dans les réseaux Internet, fait de ce dernier la technologie de prédilection qui offre un compromis entre le coût de déploiement et l'efficacité de fonctionnement. La technologie IP est confortée par le concept Tout-IP dont l'objectif est de faire migrer tous les services traditionnels vers la technologie IP. Toutefois, la réalisation des NGNs engendre plusieurs défis techniques qu'il faudra relever pour basculer définitivement d'un concept théorique vers une architecture réellement exploitable. La problématique de la qualité de service, héritée de la technologie IP, représente l'un des plus grand défis, mais l'hétérogénéité (réseaux d'accès, terminaux, etc.), la sécurité (transmissions, droits d'auteur, etc.) et le handover (horizontal entre plusieurs domaines et vertical entre plusieurs technologies), représentent aussi des verrous à lever pour le succès des NGNs.

Nous sommes convaincus que la QoS dans les NGNs doit être assurée dans un espace à trois dimensions représentées dans la Figure 1-1.

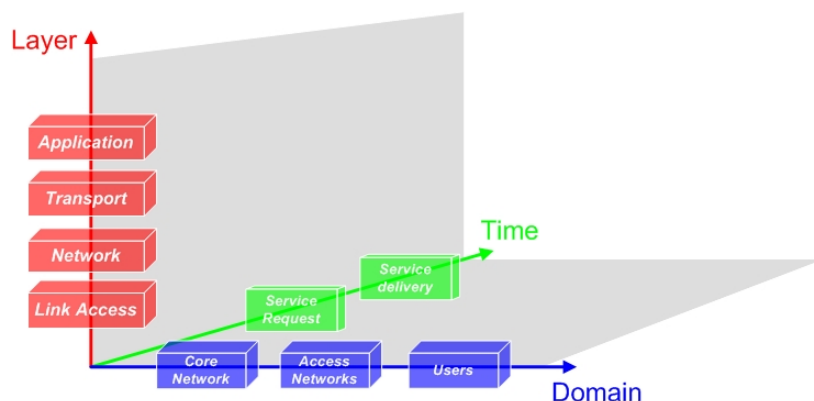


Figure 1-1 : Les dimensions de la QoS dans les NGNs

La première dimension correspond aux domaines qui doivent être traversés par les services afin qu'ils puissent être transmis du fournisseur aux consommateurs finaux. Dans cette dimension, la QoS doit être maintenue horizontalement dans le cœur du réseau et dans les réseaux d'accès. La deuxième dimension correspond au temps de vie d'un service. En effet, la QoS doit être négociée dynamiquement suivant les besoins du service demandé par l'utilisateur et le niveau de QoS contracté doit être maintenu durant la transmission d'un service puisque les conditions de transmission dans les réseaux peuvent varier durant le temps. La dernière dimension correspond aux couches réseaux à travers lesquelles la QoS doit être assurée verticalement en coordonnant les différents mécanismes de QoS qui existent sur toutes les couches. Pour garantir la QoS dans cette dernière dimension, un nouveau paradigme a émergé récemment sous l'appellation *Cross-layer* ou communication inter-couches. Ce nouveau concept remet en question l'isolation des couches qui représente le principe de base des architectures réseaux. Bien que l'efficacité de l'architecture en couches ait été démontrée dans le passé pour les réseaux filaires, le large déploiement des réseaux sans fil commence à illustrer les limites de cette architecture. La communication inter-couches devient un impératif pour améliorer les performances globales des systèmes de transmission et assurer la QoS des services multimédia.

Dans le cadre de cette thèse, nous nous sommes intéressés à la QoS à travers les couches réseau, principalement les réseaux sans fil, et nous avons étudié les interactions inter-couches *Cross-layer* afin d'identifier l'apport de ce nouveau concept dans l'amélioration de la QoS des services vidéo.

Cette étude s'inscrit dans le cadre d'une architecture globale de convergence des réseaux et des services qui permet d'interconnecter deux réseaux largement déployés durant ces dernières années, à savoir les réseaux de diffusion DVB-T et les réseaux d'accès IP/802.11. La fonctionnalité principale de cette architecture est la conversion des services DVB-T en services IPTV accessibles à partir des réseaux IP/802.11 pour des utilisateurs mobiles et hétérogènes. Pour permettre un accès universel aux flux IPTV, nous proposons dans cette thèse un nouveau système de transmission audio/vidéo adaptatif, appelé XLAVS (Cross Layer Adaptive Video Streaming). Ce nouveau système de transmission est embarqué dans des passerelles d'adaptation XLAG (Cross Layer Adaptation Gateway) déployées entre le cœur du réseau IP et les réseaux d'accès 802.11. Il se base sur le concept *Cross-layer* en centralisant toutes les informations concernant l'état des couches et toutes les décisions concernant les adaptations en un seul module, appelé XLDP (Cross Layer Decision Point). Ce dernier représente un superviseur central dont la tâche principale est d'optimiser la QoS du service IPTV. En considérant les dimensions présentées dans la Figure 1-1, le XLAVS permet d'assurer la QoS sur la dimension des couches et la dimension du temps et s'intéresse principalement aux réseaux d'accès 802.11 sur la dimension des domaines. Pour se faire, le XLAVS interagit fortement avec les terminaux de réception et les couches réseaux pour déterminer, à travers le module XLDP, la meilleure adaptation *Cross-layer* qui préserve la QoS des flux IPTV. Le débit de transmission et les taux de perte sont les principaux paramètres de QoS considérés puisque les services de streaming audio/vidéo sont très sensibles au délai de transmission et à la gigue (service non conversationnel). Les adaptations mises en œuvre par le XLAVS se basent sur les outils du standard MPEG-21 et elles sont exécutées en deux phases distinctes : lors de la demande du service IPTV et durant sa transmission. La première phase d'adaptation permet de personnaliser le flux IPTV aux caractéristiques du consommateur (terminal et utilisateur). Tandis que la deuxième phase permet de préserver la QoS des flux vidéo suivant les fluctuations du réseau sans fil durant la transmission. Dans cette dernière phase, Les adaptations sont classées en deux grandes catégories : les adaptations *Cross-layer* ascendantes exécutées au niveau applicatif et les adaptations *Cross-layer* descendantes exécutées au niveau MAC.

Afin de détailler d'avantage les concepts, les problématiques et les solutions introduits ci-dessus, cette thèse est organisée comme suit :

- Chapitre 2 «L'état de l'art» : Nous présentons dans ce chapitre un état de l'art sur les mécanismes de QoS qui ont été définis durant ces dernières années au niveau de toutes les couches du modèle TCP/IP. Au niveau physique et au niveau MAC, nous nous focalisons particulièrement sur les réseaux sans fil 802.11 et au niveau applicatif, nous nous limitons aux applications multimédia. Suite à cette présentation, nous introduisons le concept *Cross-layer* et nous détaillons les travaux récemment réalisés dans ce domaine. Ceci nous servira de base à la définition de notre nouveau système XLAVS et à l'identification des adaptations *Cross-layer* qui peuvent être implémentées dans ce nouveau système.

- Chapitre 3 « Proposition d'une architecture *Cross-layer* pour les services IPTV sans fil » : Dans ce chapitre, nous proposons une architecture complète qui permet d'interconnecter les réseaux DVB-T et les réseaux IP/802.11. Pour cela, nous commençons par détailler tous les segments constituant cette architecture : le réseau DVB-T, le cœur du réseau IP et les réseaux d'accès 802.11. Aussi, nous détaillons toutes les entités réseaux qui permettent de connecter ces segments. Par la suite, nous présentons l'architecture et le fonctionnement du nouveau système de transmission XLAVS en détaillant ses différents modules, principalement le module XLDP qui contrôle les adaptations *Cross-layer* exécutées au niveau de la passerelle XLAG.
- Chapitre 4 « Les adaptations *Cross-layer* ascendantes exécutées au niveau applicatif » : Dans ce chapitre, nous proposons deux interactions *Cross-layer* qui peuvent être mises en œuvre au niveau des passerelles XLAG pour maintenir la QoS des flux vidéo durant la transmission. Les interactions proposées adaptent les mécanismes de QoS applicative en fonction de l'état des couches inférieures : 802.11 MAC et 802.11 PHY. La première interaction assure une corrélation entre le débit applicatif et le débit physique réellement disponible dans les réseaux 802.11. Pour cela, nous discutons les techniques qui permettent d'adapter dynamiquement le débit vidéo et nous proposons un modèle mathématique qui permet d'estimer le débit physique réellement disponible pour chaque utilisateur. La deuxième interaction permet une protection efficace des flux vidéo contre les pertes en adaptant conjointement le débit vidéo et le taux de redondance FEC (Forward Error Correction) au niveau applicatif. Cette adaptation se base sur deux indicateurs de performances, à savoir les taux de perte applicatifs et la puissance du signal physique.
- Chapitre 5 « Les adaptations *Cross-layer* descendantes exécutées au niveau MAC 802.11 » : Dans ce chapitre, nous nous intéressons aux interactions *Cross-layer* qui permettent d'adapter le fonctionnement des couches inférieures, principalement la couche MAC 802.11, en fonction des besoins des couches supérieures (les flux vidéo). La première adaptation proposée autorise la couche applicative à configurer dynamiquement la fragmentation 802.11 au niveau MAC afin de réduire les pertes de paquets causées par les interférences. Au début, nous comparons les différents niveaux de fragmentation possibles : application, réseau et MAC. Ensuite, nous proposons un modèle mathématique qui détermine la taille minimale d'une trame à partir de l'*overhead* pouvant être introduit. La deuxième interaction *Cross-layer* que nous proposons, permet au flux vidéo d'occuper le canal de transmission en fonction de son débit. Pour cela, nous suggérons des opportunités de transmission au niveau MAC qui s'adaptent au débit vidéo en se basant sur des informations applicatives.
- Chapitre 6 « Conclusion et Perspectives » : Enfin, nous présentons dans ce chapitre une conclusion générale pour les travaux présentés et nous proposons quelques perspectives pour de futurs travaux de recherche qui s'inscrivent dans la continuité de cette thèse.

Chapitre 2

L'état de l'art

2.1 Introduction

Actuellement, les architectures de communication sont basées principalement sur le modèle TCP/IP. Ce modèle est caractérisé par sa structure modulaire mettant en œuvre plusieurs couches distinctes totalement isolées qui communiquent à travers des interfaces bien définies. Jusqu'à présent, l'architecture TCP/IP a démontrée sa puissance pour la transmission de données sur des réseaux filaires. Cependant, l'évolution des réseaux d'accès (WiFi, WiMax, Satellite, 3G/2G) et l'augmentation des flux multimédia (audio, vidéo, jeux interactifs) ont dévoilé rapidement les limites de cette architecture. Afin, de permettre à l'architecture TCP/IP d'intégrer d'une manière transparente ces évolutions, des mécanismes d'adaptation, de correction d'erreur, de retransmission ont été ajoutés toujours en respectant l'isolation des couches. Cette isolation stipule qu'un mécanisme mis en œuvre sur une couche particulière est destiné uniquement à améliorer les performances sur cette couche sans aucune considération pour les autres. Ceci a engendré plusieurs effets négatifs comme la redondance des mécanismes, l'annulation d'un mécanisme sur une couche sous-jacente et même par fois la contradiction entre deux mécanismes présents sur deux couches.

Pour faire face à tous ces effets négatifs, nous avons assisté récemment à l'apparition du concept *Cross-layer* qui autorise une communication et une collaboration inter-couches pour améliorer les performances de transmission. Le *Cross-layer* permet d'outre passer l'isolation des couches qui représente un inconvénient pour la transmission de flux multimédia et les transmissions sans fil.

Dans ce chapitre, nous commençons par parcourir les couches du modèle TCP/IP en détaillant, pour chaque couche, son principe de fonctionnement, les mécanismes et les techniques de QoS qui ont été ajoutés, les limites de ces mécanismes et enfin l'avantage d'autoriser une communication inter-couches. Par la suite, nous présentons le nouveau paradigme *Cross-layer* en détaillant son principe de fonctionnement, ces modèles de communication et ces approches. Nous terminerons ce chapitre en présentant les différents travaux ainsi que les différents projets qui traitent du *Cross-layer*.

2.2 Les limites de l'architecture en couches pour intégrer les réseaux IEEE 802.11 et transporter les flux multimédia

2.2.1 La couche Physique IEEE 802.11

Le niveau physique regroupe les règles et procédures qui permettent d'acheminer des éléments binaires non structurés sur un support physique entre l'émetteur et le récepteur. Ce niveau prend en charge les aspects de codage et de modulation de signaux numériques électriques et optiques. La couche physique des réseaux sans fil est définie par la famille des normes IEEE 802.11, 802.11a, 802.11b, 802.11g et bientôt 802.11n. Ces différents standards définissent des techniques d'occupation du spectre et des techniques de modulation du signal sur des bandes de fréquences bien définies. La technique de modulation définit le symbole qui sera utilisé au niveau du signal pour représenter les données, elle définit aussi le nombre de bits qui sont codés par un symbole. Une description de ces normes, suivant la chronologie de leur publication, est donnée ci-dessous.

2.2.1.1 Les standards IEEE 802.11

2.2.1.1.1 La norme IEEE 802.11 initiale

La norme initiale du IEEE 802.11 [3] a défini trois couches physiques distinctes : L'IR (Infrared) pour l'infrarouge, et le FHSS (Frequency Hopping Spread Spectrum) / DSSS (Direct Sequence Spread Spectrum) pour les transmissions radio dans la bande de fréquences IMS (Industrial, Scientific, and Medical) 2,4 GHz. Parmi ces 3 couches, le DSSS a connu un grand succès et a été largement déployé grâce à sa simplicité et aussi à sa capacité à offrir de plus grands débits. Avec le DSSS, la bande 2.412 - 2.484 GHz est partagée en 14 canaux avec une largeur de 5 MHz et un décalage de 12 MHz pour le dernier canal. Les 14 canaux ne peuvent être utilisés simultanément à cause de leur rapprochement. Le standard IEEE 802.11 suggère un espacement de 25 MHz (5 canaux) pour éviter les interférences inter-canaux, ce qui ramène à trois le nombre de canaux qui peuvent être utilisés simultanément.

La largeur d'un canal a nécessité l'utilisation d'une technique d'étalement de spectre qui permet d'étaler le signal transmis sur une bande de fréquence plus large le rendant ainsi plus résistant aux interférences. Pour cela, chaque bit du signal d'origine est remplacé par plusieurs bits dans le signal transmis en utilisant un code d'étalement. Le standard propose l'utilisation du code de *Barker* sur 11 bits qui possède des propriétés avantageuses pour les transmissions sans fil. Le DSSS utilise deux techniques de modulation le DBPSK (Differential Binary Phase Shift Keying) et le DQPSK (Differential Quadrature PSK) qui sont des modulations de phase permettant de coder des données par des changements de phase du signal transmis. Le DBPSK code un bit sur deux phases (bit 0→0°, bit 1→180°) et offre un débit de 1Mbits/s. Le DQPSK code deux bits sur quatre phases (bit 00→0°, bit 01→90°, bit 11→180°, bit 10→270°) et offre un débit de 2Mbits/s.

2.2.1.1.2 La norme IEEE 802.11a

La norme IEEE 802.11a [4] a été définie dans la bande de fréquences U-NII (Unlicensed National Information Infrastructure) 5 GHz. Elle se base sur la modulation OFDM (Orthogonal Frequency Division Multiplexing), multiplexage fréquentiel orthogonal. À l'inverse de l'étalement

du spectre, l'OFDM découpe la bande de fréquences en plusieurs sous-porteuses qui peuvent transmettre des signaux simultanément. Les sous-porteuses peuvent être distinguées puisque leurs fréquences sont choisies en respectant le principe d'orthogonalité. Le canal OFDM de 20MHz est constitué de 52 sous-porteuses dont 48 sont utilisées pour la transmission de données. L'espacement entre sous-porteuses est de 0.3125 MHz. Le standard emploie plusieurs techniques de modulation sur les sous-porteuses pour transmettre les données. Le BPSK et le QPSK offrent des débits faibles comme illustré par le Tableau 2-1. Pour permettre des débits plus élevés, le standard propose d'utiliser la modulation QAM (Quadrature Amplitude Modulation) qui représente une extension logique de QPSK. Elle se base sur deux porteuses transmises sur la même fréquence avec un décalage de 90°. Chaque sous-porteuse est modulée en amplitude pour le codage des données. Ainsi, la QAM-16 (16 états), la QAM-64 (64 états) et la QAM-256 (256 états) utilisent respectivement 4, 8, 16 amplitudes sur chaque porteuse. Le standard propose aussi un code convolutif pour la correction des erreurs en aval. Les taux de ces codes peuvent être de 1/2, 2/3, 3/4. La combinaison de la technique de modulation et du taux de codage détermine le débit binaire.

standard	Débit Mbits/s	Modulation	Le taux de codage	Bits de code par symbole	Bits/symbole
802.11 802.11g	1	DBPSK	1/11	11	1
	2	DQPSK	1/11	11	1
802.11b 802.11g	5.5	DQPSK	1/2	8	4
	11	DQPSK	1	8	8
802.11a 802.11g	6	BPSK	1/2	48	24
	9	BPSK	3/4	48	36
	12	QPSK	1/2	96	48
	18	QPSK	3/4	96	72
	24	QAM-16	1/3	192	96
	36	QAM-16	3/4	192	144
	48	QAM-64	2/3	288	192
	54	QAM-64	3/4	288	216

Tableau 2-1 : Les débits fournis par les standards IEEE 802.11

2.2.1.1.3 La norme IEEE 802.11b

La norme IEEE 802.11b [5] est une extension de la norme IEEE 802.11 initiale pour atteindre des débits de 5.5 et 11 Mbit/s. La norme propose l'utilisation de la modulation CCK (Complementary Code Keying) qui permet de réduire le code d'étalement à 8 bits contre 11 bits pour la séquence de *Barker*. Le CCK augmente aussi le nombre de bits codés par symbole : 4 bits/symbole pour un débit de 5.5 Mbit/s et 8bits/symbole pour un débit de 11Mbit/s.

2.2.1.1.4 La norme IEEE 802.11g

La norme IEEE 802.11g [6] étend les capacités du 802.11b en offrant des débits supérieurs grâce à de nouvelles techniques de modulation, mais en gardant une rétrocompatibilité avec les

normes 802.11 et 802.11b. Elle reprend la technique OFDM du 802.11a en l'adaptant pour la bande de fréquences 2.4 GHz.

2.2.1.1.5 *Le Draft 802.11n* [7]

L'objectif du groupe de travail N du 802.11 (TGN : Task Group N) est d'étendre la norme 802.11 pour atteindre un débit d'au moins 540 Mbit/s en assurant une rétrocompatibilité avec les anciennes normes. Pour atteindre cet objectif, le TGN a opté pour la technologie MIMO (Multiple-input/Multiple-output) qui consiste à utiliser plusieurs antennes, chacune attachée à un canal de radio fréquence. Ces antennes peuvent émettre et recevoir des données simultanément ce qui permettra d'augmenter considérablement le débit assuré par la couche physique. La version 4.0 du draft 802.11n a été approuvée en mai 2008.

2.2.1.2 **L'évanouissement du signal dans un canal 802.11**

Le Tableau 2-1 résume les débits fournis par les différents standards IEEE 802.11 en détaillant la modulation, le taux de codage, le débit symbole et le nombre de bits par symbole. Le bon fonctionnement de toutes ces normes dépend principalement du niveau du SNR (Signal-to-Noise Ratio) dans le canal sans fil. Le SNR, appelé aussi SIR (Signal-to-Interference Ratio) mesure le ratio entre la puissance du signal de données et l'énergie radio ambiante. Cette dernière correspond à tout signal qui ne provient pas de l'émetteur et qui représente du bruit pour le récepteur.

Durant sa transmission, le signal de données est sujet à des dégradations causées par des caractéristiques intrinsèques à la propagation des ondes radios dans l'espace. Les principales dégradations sont citées ci-dessous :

- **L'affaiblissement ou l'atténuation** : Elle est causée par la dispersion du signal avec la distance et de son absorption par les obstacles. Cette dispersion dépend de l'environnement de propagation du signal.
- **Le bruit** : Il est causé par des signaux aléatoires perturbateurs qui altèrent le signal transmis.
- **La propagation multi-trajet** : Elle est causée par la réflexion du signal transmis sur des obstacles. Ceci génère plusieurs copies du signal original qui ont des délais de propagation variables.

La combinaison de tous ces effets provoque l'évanouissement « fading » du signal dans le canal sans fil. Il existe deux types d'évanouissement [8] :

- **Évanouissement rapide « fast or small-scale fading »** : Il correspond à une baisse importante de la puissance du signal pendant un laps de temps très court durant lequel le décodage du signal est impossible. Ceci engendre une altération des symboles et par conséquent une altération des bits de données. Dans certains cas, le code convolutif peut corriger ces erreurs en aval, si ses derniers ne dépassent pas un certain seuil.
- **Évanouissement lent « slow or large-scale fading »** : Il correspond à une baisse moyenne de la puissance du signal durant un intervalle de temps assez long. Cette baisse est causée par l'atténuation du signal lorsque le récepteur s'éloigne de l'émetteur ou bien par l'absorption du signal causée par des changements physiques de l'environnement (espace clos/ouvert,

soleil/pluie/neige). La capacité du récepteur à décoder le signal dépend principalement du type de modulation utilisé par l'émetteur.

2.2.1.3 Les mécanismes pour palier l'évanouissement du signal

Plusieurs techniques peuvent être utilisées pour palier l'évanouissement du signal. Nous détaillons ci-dessous les principales techniques :

- **La Diversité (diversity)** [9][10][11][12] : Le principe de base de la diversité est d'envoyer plusieurs copies du signal. Le récepteur combine tous les signaux reçus pour un décodage optimal des symboles. Il existe plusieurs types de diversité exploités dans les communications sans fil : la diversité fréquentielle, la diversité temporelle et la diversité spatiale. En utilisant la diversité fréquentielle, le signal est modulé sur plusieurs porteuses en utilisant des fréquences différentes. Pour la diversité temporelle, chaque symbole est envoyé plusieurs fois sur la même porteuse. Le temps qui sépare l'envoi du même symbole doit être cohérent afin que les deux symboles expérimentent un évanouissement différent. Enfin, la diversité spatiale propose d'utiliser plusieurs antennes pour transmettre ou recevoir le même signal. La distance entre les antennes doit être assez importante pour que les évanouissements du signal perçus par les antennes soient indépendants. Cette dernière diversité est plus intéressante puisqu'elle ne nécessite pas de bande passante (diversité fréquentielle) ou de délai de transmission (diversité temporelle) additionnels, par contre, il est pratiquement impossible de déployer ce type de diversité sur des terminaux de petite taille.
- **Les Codes convolutifs** [13] : Il existe deux catégories de code de correction d'erreurs largement utilisés pour la transmission sans fil : les codes en blocs et les codes convolutifs. Un code en bloc (n, k) , avec $n > k$, traite les données par bloc de k bits à la fois, produisant un bloc de n bits en sortie. Par contre, un code convolutif est défini au moyen de trois paramètres n , k et K , avec des valeurs de n , k plus petites. De plus, le bloc de n bits produit en sortie dépend des k bits du bloc traité et de $K-1$ bits du bloc précédent. Cette catégorie a été conçue principalement pour les transmissions de flux en continu afin que le contrôle et la correction d'erreurs puissent, également, se faire en continu.
- **Le support de plusieurs débits (Multi-rate capability)** : les standards IEEE 802.11 proposent plusieurs types de modulation et taux de codage dont la combinaison fait varier le débit physique de 1 Mbits/s à 54 Mbits/s (voir le Tableau 2-1). Les débits élevés, comparés aux débits faibles, utilisent des modulations complexes qui placent plus de bits dans un intervalle temps. Ces modulations complexes sont plus sensibles aux bruits et nécessitent un certain niveau de SNR pour leur décodage. Ceci limite par conséquent leur portée. Par contre, les modulations simples qui fournissent de faibles débits ont une plus grande portée puisqu'elles sont plus robustes et nécessitent un signal moins puissant pour leurs décodages. Ceci permet d'étendre le champ de couverture d'un réseau sans fil en utilisant des débits moins importants. Cependant, les couches physiques des standards se sont contentées de définir les modulations et les codages sans définir des mécanismes et des algorithmes qui permettent de basculer d'un débit vers un autre.

2.2.1.4 Discussion

Dans une architecture TCP/IP classique, les couches supérieures subissent les évanouissements lents et les évanouissements rapides avec des réactions limitées à leurs niveaux, régies par le principe d'isolation des couches. Cependant, il est clair que le partage d'information, sur la variation de l'état du canal de transmission et/ou la puissance du signal reçu, entre la couche physique et les couches supérieures pourrait améliorer les performances de transmission dans les réseaux sans fil. Cette amélioration peut se concrétiser par l'élaboration d'algorithmes sophistiqués au niveau MAC qui permettent de basculer d'un débit physique vers un autre. Ces algorithmes pourront accéder à la couche physique pour récupérer le niveau du signal et décider du débit physique idéal. De plus, le débit utilisé au niveau physique peut être partagé avec les couches applicatives qui doivent adapter leurs débits en conséquence.

2.2.2 La couche d'accès IEEE 802.11

La première fonction de la couche liaison de données est le partage équitable du support physique entre plusieurs stations en mettant en œuvre des mécanismes distribués pour accéder au canal. La deuxième fonction consiste à délimiter les trames et à corriger d'éventuelles erreurs causées par la transmission physique. Dans le modèle de référence IEEE 802, la couche liaison de données est constituée de la couche LLC (Logical Link Control) et de la couche MAC (Media Access Control). La couche MAC IEEE 802.11 couvre trois domaines fonctionnels : le contrôle de l'accès au canal, la livraison fiable de données et la sécurité. Dans la suite de cette section, nous nous intéresserons uniquement aux deux premiers domaines fonctionnels.

2.2.2.1 L'accès au canal dans la couche MAC 802.11

Le standard 802.11 se base principalement sur deux méthodes pour partager l'accès au canal : un accès distribué DCF (Distributed Coordination Function) et un accès centralisé PCF (Point Coordination Function). Pour permettre la cohabitation entre ces deux méthodes d'accès, la norme 802.11 a défini des intervalles de temps, appelés « superframe », qui sont partagés en deux périodes : une période sans contention durant laquelle le PCF est utilisé et une période de contention utilisant le DCF.

2.2.2.1.1 DCF (*Distributed Coordination Function*)

Le DCF utilise la méthode d'accès CSMA/CA (Carrier Sence Multiple Access/ Collision Avoidance) dérivée du CSMA/CD (Collision Detection) qui est utilisée par le protocole Ethernet dans les réseaux filaires. Avec le CSMA/CD, la station émettrice transmet le signal et écoute en même temps le canal pour détecter d'éventuelles collisions avec d'autres signaux présents sur le canal. Cependant, dans un réseau sans fil, les collisions entre signaux ne peuvent être détectées par la station émettrice parce que la puissance du signal émis par cette dernière masque tous les autres signaux présents sur le canal. Ainsi, le CSMA/CA est utilisé pour écouter le canal et vérifier qu'il n'est pas occupé avant d'entamer la transmission.

Pour assurer son fonctionnement, le DCF se base sur un jeu d'intervalle de temps appelé IFS (Inter Frame Spacing). La Figure 2-1 illustre les règles d'accès au canal avec l'emploi des différents IFS. Lorsqu'une station souhaite émettre une trame de données, elle écoute le canal durant un

intervalle de temps appelé DIFS (Distributed IFS). Si celui-ci est inactif durant cette période, la station transmet immédiatement sa trame. Dans le cas contraire, le canal est occupé, la station doit attendre jusqu'à ce que le canal soit inactif durant une période SIFS (Short IFS). À la fin de cette période, la station entre dans la procédure du *Backoff* qui l'oblige à calculer un *Backoff Time* (BT) durant lequel elle s'abstient de transmettre. Le BT est calculé comme suit $BT = Random() \times timeslot$. La fonction *Random()* retourne une valeur aléatoire dans la fenêtre de contention CW (Contention Windows) $[0, CW-1]$ et le *timeslot* représente une unité de temps fixe. Si le canal reste inactif durant cette période, la station transmet sa trame. La taille de la fenêtre de contention est initialisée avec une valeur minimale CW_{min} et augmente exponentiellement en doublant sa taille à chaque échec de transmission. La taille de la fenêtre est réinitialisée à CW_{min} dans les cas suivants : lorsqu'elle atteint la valeur CW_{max} , après un nombre *retry* de retransmission, ou bien après une transmission correcte. La station réceptrice vérifie le CRC (Cyclic Redundancy Check) de la trame reçue et envoie une trame d'acquittement à l'émetteur après un SIFS. Ce dernier est utilisé pour séparer les trames d'un dialogue. La réception de l'acquittement informe l'émetteur du succès de la transmission.

Lorsque le PCF est utilisé, le coordinateur central a la priorité d'accéder au canal. Pour cela, il utilise l'intervalle PIFS (Point IFS) qui est plus court que le DIFS. Une fois le canal saisi, le coordinateur central instaure le PCF.

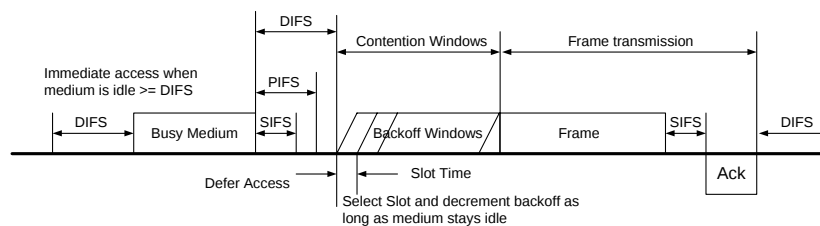


Figure 2-1 : L'accès distribué au canal DCF

2.2.2.1.2 PCF (Point Coordination Function)

Le PCF met en œuvre un accès de type réservation. Il est complémentaire au DCF mais il est optionnel. Au début de la période sans contention, le point coordinateur PC (Point coordination) transmet des balises *Beacons* qui contiennent la durée maximale $CFPMaxDuration$ de la période d'accès sans contention. Ensuite, il commence à interroger les stations associées en envoyant des trames *CF-Poll* afin de savoir si elles possèdent des données à transmettre. Le *CF-Poll* peut être accompagné d'une trame de données si le PC a des données à transmettre pour une station. La station qui est destinataire du *CF-Poll* envoie sa trame en intégrant un acquittement *CF-ACK* qui acquitte le *CF-Poll* après un temps d'attente SIFS. Enfin, le PC acquitte la trame envoyée par la station après un intervalle SIFS. Cet acquittement est généralement accompagné par un *CF-Poll* pour interroger une autre station. Dans le cas où une station interrogée ne répond pas au bout d'un temps PIFS (elle n'a pas de données à transmettre), le PC reprend l'interrogation des stations qui restent. La Figure 2-2 illustre un exemple d'une période sans contention durant laquelle le PC interroge trois stations sta1, sta2 et sta3.

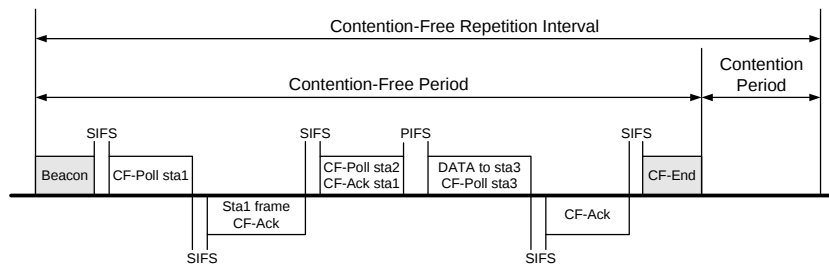


Figure 2-2 : L'accès centralisé au canal PCF

2.2.2.2 La livraison fiable dans la couche MAC 802.11

Le standard IEEE 802.11 a proposé plusieurs mécanismes pour palier au manque de fiabilité de la couche physique présentée dans la section 2.2.1.2. Un bref descriptif de ces mécanismes est donné ci-dessous :

- La retransmission :** La couche MAC 802.11 se base sur le système de correction d'erreurs ARQ (Automatic Repeat reQuest). Ce dernier emploie les codes de détection d'erreurs, les acquittements et les retransmissions pour assurer la fiabilité des données. En effet, la trame MAC intègre une valeur de contrôle CRC qui permet de vérifier l'intégrité des données au niveau du récepteur. Dans le cas où les données sont correctes, le récepteur envoie un acquittement à l'émetteur. Dans le cas contraire, l'émetteur attend l'acquittement durant un certain temps *timeout* et retransmet la trame en considérant que la précédente est perdue. L'émetteur répète cette opération un nombre de fois limité jusqu'à la réception d'un acquittement. La procédure d'envoi d'une trame de données et d'attente de son acquittement constitue une unité atomique de dialogue durant laquelle aucune autre station ne peut accéder au canal.
- L'utilisation de RTS/CTS :** Ce mécanisme a été proposé par la norme 802.11 pour éviter les collisions des trames transmises, en même temps, par deux ou plusieurs stations. Principalement, dans une configuration où deux stations éloignées communiquent avec une station se trouvant au milieu. Cette configuration engendre le problème de « la station cachée » puisque les deux stations éloignées ne perçoivent pas leurs signaux et peuvent transmettre en même temps des trames à la station intermédiaire, ce qui provoque des collisions. Pour éviter ce problème, la requête RTS (Ready To Send) et la réponse CTS (Clear To Send) ont été intégrées au mécanisme CSMA/CA avant un envoi d'une trame de données. Ainsi, lorsqu'une station émettrice détient le canal pour une transmission, elle commence par envoyer une requête RTS à la station réceptrice qui répond par une réponse CTS. Les messages RTS/CTS sont aussi reçus par toutes les autres stations qui doivent différer leurs transmissions en conséquence.
- La fragmentation :** Afin de faire face à l'évanouissement rapide du signal responsable de la corruption des trames durant leur transmission, la norme 802.11 définit la fragmentation au niveau MAC. La fragmentation permet la décomposition des paquets provenant de la couche LLC en plusieurs trames MAC. Ceci permet de réduire la taille des trames transmises sur le canal sans fil et d'augmenter, par la même, le nombre de trames livrées sans erreur. Le

mécanisme de fragmentation proposé par la norme 802.11 est augmenté du mécanisme *burst* qui permet d'envoyer tous les fragments d'un même paquet en rafale une fois le canal détenu.

2.2.2.3 La QoS au niveau de la couche MAC 802.11

Le DCF défini par le standard 802.11 fournit un service au mieux (best-effort) qui n'offre aucune garantie aux stations pour l'accès au canal. Ceci pose inévitablement le problème de la qualité de service exigée par les flux multimédia qui sont très sensibles au délai de transmission et à la variation du débit. Plusieurs travaux ont été menés pour introduire des mécanismes de QoS au niveau de la couche MAC 802.11 en préservant son principe de fonctionnement. Ces mécanismes peuvent être classés en deux catégories : la différenciation de services et la garantie de service. La Figure 2-3 présente une taxonomie des travaux menés dans chaque catégorie. Nous détaillons ces travaux dans la suite de cette section.

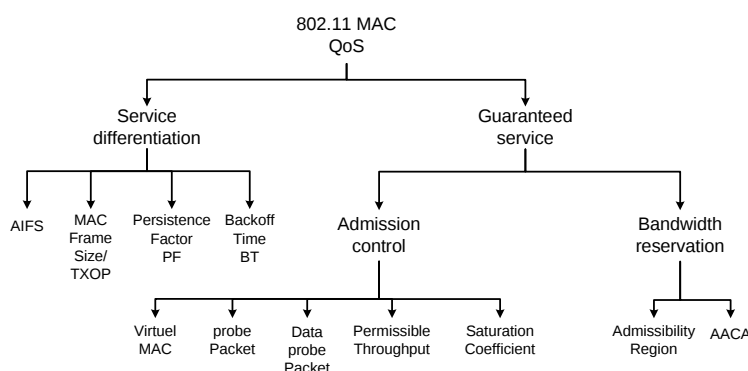


Figure 2-3 : Taxonomie des mécanismes de QoS au niveau MAC 802.11

2.2.2.3.1 La différenciation de service au niveau MAC

Le principe de base de la différenciation de service au niveau de la couche MAC est la définition de plusieurs classes de trafic offrant un accès différencié au canal de transmission. Plusieurs paramètres peuvent être utilisés pour construire ces classes de trafic. Nous présentons, ci-dessous, les principaux paramètres :

- **Arbitrary IFS (AIFS)** [14][15] : Les classes de trafics utilisent des IFS différents. La classe utilisant un petit IFS peut démarrer la transmission ou la procédure du *Backoff* plus rapidement qu'une classe utilisant un grand IFS. Ceci lui procure une priorité pour accéder au canal de transmission.
- **La taille d'une trame MAC/TXOP (Transmission Opportunity)** [16] : Les classes de trafics sont autorisées à émettre une quantité de données différentes une fois le canal de transmission détenu. La quantité de données peut être différenciée en changeant la taille de la trame, ou bien, en envoyant plusieurs trames à la fois durant une opportunité de transmission TXOP. Les classes prioritaires pourront transmettre plus de données, ce qui fait augmenter leur débit de transmission.
- **Persistence Factor PF** : Un facteur P est associé à chaque classe de trafic, les classes prioritaires possèdent un petit facteur P . Dans [17], durant la procédure du *Backoff*, une valeur aléatoire r est générée pour chaque time slot. Le backoff est arrêté pour démarrer la

transmission si $r > P$. Dans [18], pour chaque retransmission, la taille de la fenêtre de contention est multipliée par le facteur P au lieu de la multiplier par deux comme le propose l'algorithme standard DCF.

- **Backoff Time (BT)** : Plusieurs modifications ont été proposées pour l'algorithme du *Backoff* afin de construire des classes de trafics. L'approche, la plus simple, consiste à choisir des CW_{min} et CW_{max} différents pour chaque classe [18]. Ainsi, les classes prioritaires ont de petites valeurs CW_{min} et CW_{max} et par conséquent un BT moins important, comparé aux classes moins prioritaires. D'autres approches proposent la modification de la manière dont le BT évolue [19]. Des algorithmes multiplicatifs ou bien additifs peuvent être choisis [20]. Un algorithme multiplicatif réagit rapidement par des augmentations et des diminutions importantes. Tandis que les algorithmes additifs réagissent lentement avec un comportement lissé (smooth). Enfin, les approches [21][22][23] proposent de calculer le BT en utilisant plusieurs paramètres comme le débit espéré, le débit réalisé et le poids de la classe de trafic.

2.2.2.3.2 La garantie de service au niveau MAC

Les mécanismes de différenciation de service présentés dans la section précédente fonctionnent moins bien lorsque la charge du trafic dépasse un certain seuil [25]. Dans ces cas, l'utilisation de mécanisme de garantie de service est préférable pour mieux contrôler et partager les ressources disponibles. La garantie de service se base sur deux mécanismes importants : le contrôle d'admission et la réservation de la bande passante. L'implémentation de ces mécanismes dans le DCF est très contraignante, due à la nature distribuée du mécanisme d'accès CSMA/CA.

Le contrôle d'admission nécessite des moyens pour déterminer l'état du réseau et estimer les ressources disponibles. Les travaux dans [26] proposent d'utiliser un algorithme MAC virtuel qui effectue un contrôle passif en utilisant des trames MAC virtuelles. Ceci afin d'estimer le niveau de service (débit et délais) et de déterminer si le canal peut supporter d'autres demandes de service. Dans [27], les auteurs proposent d'utiliser des paquets sondes pour estimer les ressources disponibles et décider de nouvelles admissions. Par contre dans [28], l'estimation se base sur les paquets de données pour mesurer la charge du réseau. Une solution heuristique est proposée dans [29] en utilisant le débit acceptable (permissible throughput) pour décider de nouvelles admissions. Une implémentation de cette solution, utilisant le protocole de routage (AODV Ad Hoc On Demand Distance Vector), a été proposée. Enfin, un coefficient de saturation est proposé dans [23] pour permettre à une station de savoir si elle approche des conditions de saturation. Le coefficient de saturation est calculé en utilisant plusieurs informations : nombre de stations actives, le débit de leurs trafics et la longueur moyenne des paquets. Ces informations sont incluses dans des paquets de données qui sont échangés entre les stations.

Concernant la réservation de bande passante, les auteurs dans [21] présentent le principe de la région d'admissibilité AR (Admissibility Region) qui se base sur la technique du saut à jeton pour réserver la bande passante. Lorsqu'un nouveau nœud rejoint le réseau, l'AR réajuste les fenêtres de contention utilisées par tous les nœuds dans le réseau afin d'améliorer les performances. Dans [30], un nouveau protocole d'accès MAC, utilisant plusieurs canaux et supportant des classes de réservation, est présenté sous l'appellation AACA (Adaptive Acquisition Collision Avoidance). AACA utilise les requêtes RTS/CTS pour une réservation adaptative des canaux libres. L'utilisation

de plusieurs canaux permet d'améliorer les performances de transmission, comparé à l'utilisation d'un seul canal, même en partageant la même bande passante.

2.2.2.4 Le standard IEEE 802.11e

Devant la nécessité d'introduire la QoS au niveau de la couche MAC 802.11, le groupe de travail 802.11 a publié en 2005, sous le nom de 802.11e [31], un amendement au standard IEEE 802.11. Cet amendement définit un ensemble d'améliorations de la couche MAC pour le support de la QoS tout en gardant une rétrocompatibilité avec les anciennes normes IEEE 802.11 a/b/g. Les points d'accès et les stations qui implémentent cette nouvelle norme sont appelés respectivement QAP (QoS-enhanced Access Point) et QSTA (QoS-enabled station). Le 802.11e définit une nouvelle fonction de coordination appelée (HFC : Hybrid Coordination Function). Le HFC utilise deux méthodes d'accès concurrentes : EDCA (Enhanced Distributed Channel Access) et HCCA (HFC-Controlled Channel Access). C'est deux méthodes sont détaillées dans les sous-sections suivantes.

2.2.2.4.1 EDCA (Enhanced Distributed Channel Access)

L'EDCA représente une amélioration de la méthode d'accès par contention DCF pour introduire plusieurs mécanismes de différenciation de services. L'EDCA définit quatre catégories d'accès (ACs : Access Categories) : AC_BK (Background), AC_BE (Best Effort), AC_VI (VIdéo), AC_VO (VOIce). Chaque catégorie possède ces propres paramètres pour accéder au canal de transmission. Trois techniques ont été retenues pour la construction de ces catégories : différents intervalles $AIFS[i]$, différentes tailles pour la fenêtre de contention ($CW_{min}[i]$ et $CW_{max}[i]$) et enfin différentes opportunités de transmission ($TXOP[i]$), $i = \{AC_BK, AC_BE, AC_VI, AC_VO\}$. Ceci est illustré par la Figure 2-4. L'algorithme DCF présenté dans la section 2.2.2.1.1 n'a pas été modifié, mais le jeu d'intervalle sur lequel il repose a été personnalisé pour chaque catégorie d'accès. En utilisant des intervalles de temps réduits ($AIFS[i]$, $CW_{min}[i]$ et $CW_{max}[i]$), les classes prioritaires ont un accès plus rapide au canal et elles peuvent assurer un plus grand débit en utilisant un TXOP plus long.

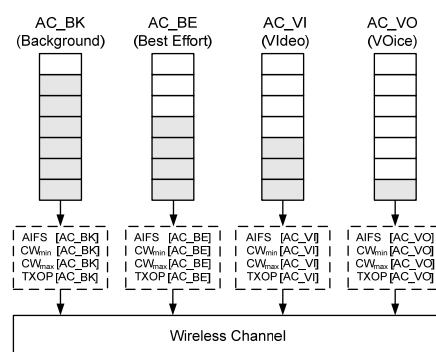


Figure 2-4 : Les catégories d'accès EDCA

2.2.2.4.2 HCCA (HFC-Controlled Channel Access)

L'objectif principal du HCCA est de fournir des mécanismes de garantie de service au niveau de la couche MAC. Le HCCA étend la méthode d'accès PCF pour le support de plusieurs classes de trafics et fonctionne en concurrence avec l'EDCA. Le HCCA définit le HC (Hybrid Coordinator) qui peut être associé au QAP dans une architecture WLAN. D'une manière similaire au PC, le HC utilise le PIFS pour détenir le canal de transmission et alloue des périodes HCCA-TXOPs aux QSTAs pour les autoriser à transmettre. Les TXOPs alloués sont de différentes durées suivant la QoS demandée par les QSTAs. De plus, le HC peut allouer ces TXOPs, même durant la période de contention. Chaque QSTA doit transmettre une demande explicite de ces besoins en QoS au HC en fonction des flux TS (Traffic Stream) qu'elle veut envoyer. Pour décrire ces besoins, la station utilise le TSPEC (Traffic Specification) qui définit un ensemble de paramètres de QoS pour un TS: le délai d'interactivité, le débit moyen, le délai limite, la taille d'une trame MAC. Cependant, il est important de noter que le HCCA-TXOP est alloué par QSTA qui est, à son tour, responsable du partage de HCCA-TXOP entre ces TS.

2.2.2.4.3 Autres mécanismes de QoS

Le standard IEEE 802.11e définit d'autres mécanismes qui permettent d'améliorer les performances dans un réseau WLAN, et par la même, d'améliorer la QoS. Ci-dessous, nous donnons un bref aperçu des principaux mécanismes :

- **Block Acknowledgments (Block Ack) :** Ce mécanisme permet d'améliorer l'exploitation d'un canal de transmission en autorisant un QAP à envoyer plusieurs trames vers une QSTA sans recevoir d'acquiescement immédiat. À la fin des transmissions, la QSTA envoie un bloc d'acquiescements pour informer sur l'état de réception du bloc de trames. Ceci permet de partager le temps nécessaire pour accéder au canal sur plusieurs trames.
- **No Acknowledgments (No Ack) :** Ce mécanisme améliore aussi l'exploitation d'un canal en permettant à une QSTA de ne pas envoyer d'acquiescement. Cependant, la mauvaise réception des trames est signalée par des « No Ack ». Lorsque le « No Ack » est utilisé, la retransmission au niveau MAC n'est pas activée, ce qui réduit considérablement la fiabilité des transmissions.
- **Direct Link Protocol (DLP) :** Ce protocole autorise deux QSTAs dans un réseau WLAN à échanger des trames directement entre elles, sans passer par le QAP. Ceci permet de réduire le délai de transmission en supprimant le temps de relais QAP. De plus, il permet de diminuer la charge du QAP.

2.2.2.5 Discussion

Le standard 802.11 a introduit au niveau MAC plusieurs mécanismes, exposés dans la section 2.2.2.2, afin de fiabiliser les transmissions sur un canal sans fil. Cependant, ces mécanismes ont été définis sans aucune considération de ce qui existe sur les couches supérieures, ni de l'effet que peut avoir ces mécanismes sur les flux transmis. Par exemple, le mécanisme de retransmission introduit un délai supplémentaire pour la transmission des trames. Ce délai peut être contraignant pour les services temps réel interactifs. De plus, la retransmission au niveau MAC engendre un autre inconvénient pour les flux TCP qui utilisent un mécanisme de retransmission au niveau transport.

La duplication de ce service, sur deux couches différentes est contre-productive pour le fonctionnement total du système.

L'utilisation des messages RTS/CTS introduit, lui aussi, un délai de transmission supplémentaire en alourdissant l'algorithme d'accès au canal DCF. Enfin, la fragmentation qui permet de réduire la probabilité de corruption d'un paquet, augmente le pourcentage d'en-tête (*overhead*) puisque l'en-tête de la couche MAC et de la couche physique sont dupliqués sur chaque fragment. Ceci signifie que la taille d'un fragment est un compromis entre les pertes provoquées par les interférences et l'*overhead* introduit par la duplication des en-têtes. Les fonctions d'optimisation de ces deux paramètres (taux de perte et *overhead*) s'opposent, le premier paramètre tend à réduire la taille d'un fragment et le deuxième à l'augmenter.

Une communication inter-couches permettra de faire face à tous ces inconvénients en autorisant un contrôle dynamique de ces mécanismes (nombre de retransmissions, RTS/CTS et fragmentation) en fonction de l'état du canal, des taux de perte et/ou du type de trafic. L'objectif de ces communications inter-couches est d'améliorer les performances de transmission en trouvant un compromis entre tous ces paramètres.

De plus, une collaboration entre les mécanismes proposés par le 802.11e avec les mécanismes de QoS existants sur les couches supérieures est nécessaire. Cette collaboration permettra de mapper les besoins des flux à travers les couches réseaux et d'assurer une continuité verticale de la qualité de service.

2.2.3 La couche réseau

Cette couche représente le cœur de l'architecture TCP/IP. Elle se base principalement sur le protocole IP (Internet Protocol) dont la première fonction est l'acheminement de paquets IP de bout-en-bout indépendamment les uns des autres jusqu'à leur destinataire final. Le protocole IP offre un acheminement au mieux (Best-Effort), non fiable et non orienté connexion. La version largement déployée actuellement est la version IPv4 [32] antérieure à la version IPv6 [33] qui a été conçue principalement pour élargir la plage d'adressage IP. L'adoption de IPv6 s'effectue progressivement notamment par les pays à forte population : la Chine, l'Inde, la Corée, etc.

Le service Best-Effort fourni par le protocole IP n'offre aucun contrôle, ni garantie, sur les conditions de transport de bout-en-bout des paquets IP. Grâce à sa simplicité, ce service a rencontré un grand succès dans le passé pour le transport des données asynchrones. Cependant, ce succès et de plus en plus contesté, actuellement, par la multiplication des flux synchrones, ou temps réel, qui exigent une certaine QoS pour leur transport. Cette QoS, appelée communément QoS réseau, se base sur plusieurs métriques qui caractérisent les performances d'un réseau IP. Les principales métriques sont :

- **La bande passante** : Elle représente le débit de bout-en-bout disponible sur le réseau.
- **Le taux de perte** : Il correspond au rapport du nombre de paquets perdus par le nombre de paquets émis durant la transmission.
- **La gigue** (Variation du délai) : Il s'agit de la variation des délais d'acheminement de bout-en-bout des paquets sur le réseau.

- **Le délai de transmission** (la latence) : Il représente le temps nécessaire pour traverser le réseau de bout-en-bout de l'émetteur jusqu'au récepteur.

Pour répondre aux besoins des flux temps réel, l'IETF (Internet Engineering Task Force) a développé plusieurs modèles et protocoles pour une gestion optimisée de la QoS réseau sans pour autant changer le mode de fonctionnement de base du protocole IP :

- **Le modèle IntServ** : Pour garantir la qualité de service.
- **Le modèle DiffServ** : Pour offrir une qualité de service différenciée.
- **Le protocole MPLS (Multi Protocol Label Switching)** : Pour une meilleure gestion du trafic dans le réseau IP en considérant la QoS.

2.2.3.1 Le service garanti (IntServ : Integrated services)

IntServ [34] est une architecture de QoS qui permet la réservation de ressources sur le réseau pour un flux spécifique. Ce dernier correspond à une séquence de paquets possédant les mêmes adresses IP sources et destinations et les mêmes besoins de QoS. Pour décrire les besoins d'un flux en QoS, IntServ se base sur la description du trafic TSpec [35] qui inclut plusieurs informations : p (débit crête, octets/seconde), b (taille maximale du burst, octets), r (débit moyen, octets/seconde), m (unité minimum contrôlée, octets), M (taille maximum d'un paquet, octets). L'architecture IntServ définit deux types de services principaux, en plus du service Best-Effort :

- Le service de charge contrôlée (CLS: Controlled Load Service) [36]: Ce service fournit approximativement la même QoS pour un flux indépendamment de l'état du réseau : état normal ou état de surcharge. La différence entre ce service et le service Best-effort est décelable uniquement quand le réseau est surchargé.
- Le service garanti (GS : Guaranteed Service) [37]: Ce service assure une bande passante stricte, un délai de transmission de bout-en-bout borné et une perte de paquet nulle pour un flux.

IntServ a défini trois types d'éléments réseaux (NE, Network Element) pour son déploiement : Le NE *QoS-capable* implémente un ou plusieurs services IntServ, le NE *QoS-aware* offre un service QoS équivalent à celui offert par IntServ et enfin le NE *Non-QoS* n'offre aucun service QoS.

Afin d'assurer une QoS spécifique à chaque flux, les NEs *QoS-capable* mettent en œuvre un certain nombre de fonctions pour la gestion des paquets, par exemple, un ordonnanceur, un classificateur et un gestionnaire de files d'attente. De plus, le NE *QoS-capable* doit contenir un plan de contrôle pour la réservation de ressources, le contrôle d'admission et le contrôle d'accès. Ces fonctionnalités se basent principalement sur le protocole de réservation de ressources RSVP (ReSerVation Protocol) [38][39][40] recommandé par l'IETF.

Le protocole de signalisation RSVP permet la réservation dynamique de ressources sur le réseau en mettant en place une connexion logique, appelée session. Cette dernière est identifiée par trois paramètres : l'adresse de destination, l'identifiant du protocole (IP) et le port de destination (optionnel). Dans RSVP, l'initialisation et le maintien de la réservation des ressources sont contrôlés par le récepteur et la réservation est unidirectionnelle, de l'émetteur vers le récepteur. Les

paramètres de contrôle du trafic (QoS) et les politiques de contrôle véhiculés par RSVP sont définis en dehors de ce dernier. Ces paramètres sont documentés dans des spécifications propres à IntServ.

Malgré sa robustesse, l'architecture IntServ n'a pas rencontré le succès escompté pour assurer la QoS au niveau IP. Le premier inconvénient de cette architecture est sa scalabilité. En effet, la nécessité de maintenir l'état des ressources réseau pour chaque flux limite sérieusement le déploiement à grande échelle d'IntServ. Les autres reproches concernent sa complexité de mise en œuvre et l'*overhead* introduit sur le réseau en utilisant un protocole de réservation de ressources.

2.2.3.2 La différenciation de service au niveau IP (Differentiated services) :

Devant les limites de l'architecture IntServ, l'IETF a introduit un autre modèle de QoS plus simple à déployer, appelé DiffServ [41]. Au lieu d'offrir une QoS pour chaque flux comme le propose IntServ, DiffServ classe les paquets de différents flux en classes de trafic en se basant sur une signalisation dans la bande « in-band », à savoir, le champ DSCP (DiffServ Code Point) défini sur 6 bits et présent au niveau de l'en-tête des paquets IP [42]. Les paquets appartenant à la même classe identifiés par le DSCP sont traités suivant un comportement approprié, appelé PHB (Per Hop Behavior).

Un PHB est une description du comportement de transmission d'un nœud DiffServ, face à une classe de trafic particulière, observé de l'extérieur. La mise en œuvre d'un PHB au niveau des nœuds DiffServ se fait grâce à des mécanismes de gestion et d'ordonnancement de files d'attente. L'architecture DiffServ définit deux types de PHB :

- **Expedited Forwarding (EF)** [43] : Le PHB EF fournit un service équivalent à une liaison virtuelle ou à un service premium, en garantissant la bande passante et en assurant un taux de perte, un délai de transmission et une gigue très faible. Pour ce service il est important de bien connaître le débit minimum sortant d'un nœud DiffServ. Ce débit doit être supérieur au débit entrant. Le DSCP correspondant à ce service est : 101110.
- **Assured Forwarding (AF)** [44]: le PHB AF permet la définition de plusieurs classes de trafic en utilisant plusieurs PHBs. Il est possible de définir N classes (AF) indépendantes de M différents niveaux de priorités (drop precedence). Actuellement, quatre classes de traitement (N=4) sont définies et chacune d'elles comprend trois niveaux de priorité (M=3). Les classes AF sont configurées pour consommer davantage de ressources lorsqu'elles sont disponibles. En cas de surcharge, le nœud DiffServ supprime les paquets dans les classes AF en fonction de leur niveau de priorité. La mise en œuvre de l'AF doit minimiser la surcharge permanente de chaque classe tout en autorisant les pics de trafic.

La Figure 2-5 illustre une architecture DiffServ qui définit des domaines, des régions et des nœuds spécifiques. Le domaine DiffServ (DS Domain) représente une zone administrative, avec plusieurs nœuds partageant la même définition de services et de PHB. Un ensemble de domaines DiffServ constitue une région DiffServ (DS region). Le domaine DiffServ est constitué de deux types de nœuds : les nœuds frontières (DS boundary nodes) et les nœuds intérieurs (DS interior nodes).

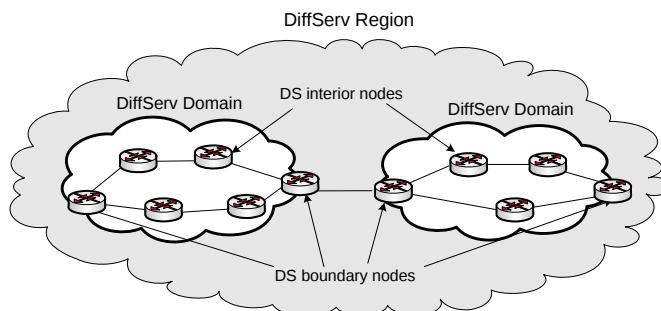


Figure 2-5 : L'architecture DiffServ

Les nœuds frontières mettent en œuvre deux principaux mécanismes : la classification et le conditionnement du trafic. La classification peut s'effectuer suivant deux techniques :

- **BA (Behavior Aggregate)** : La classification est établie en fonction de la valeur DSCP.
- **MF (Multi-Field)** : La classification peut prendre en considération d'autres paramètres comme l'adresse source ou destination.

Le conditionnement du trafic, illustré par la Figure 2-6, contient les composants suivants :

- **Le mètre (Meter)** : Il mesure le trafic afin de s'assurer qu'il est conforme au profil contracté avec l'utilisateur. Les mesures sont communiquées aux autres composants pour appliquer la politique de contrôle appropriée.
- **Le marqueur (Marker)** : Il modifie la valeur du champ DSCP.
- **Le lisseur (Shaper)** : Il se charge de lisser le trafic en retardant l'envoi de certains paquets pour les rendre conforme à leur profil. Pour cela, il stocke les paquets dans un buffer de taille limitée.
- **Le supprimeur (Dropper)** : Il supprime les paquets dans le cas où le débit défini dans le profil est dépassé.

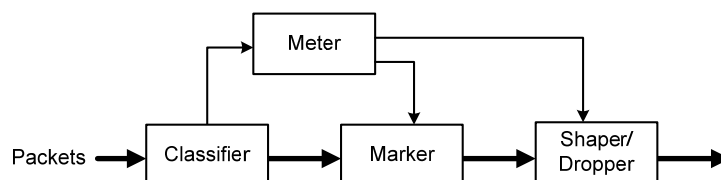


Figure 2-6 : Le conditionnement du trafic DiffServ

2.2.3.3 L'architecture MPLS

Contrairement aux deux architectures précédentes, IntServ et DiffServ, MPLS n'est pas, proprement dit, une architecture de qualité de service mais plutôt une architecture qui appartient à l'ingénierie de trafic. Cette dernière déploie des protocoles et des mécanismes dans un réseau afin d'améliorer son fonctionnement en exploitant plus efficacement les ressources disponibles. Ceci conduit indirectement à améliorer la qualité de service d'un réseau.

L'architecture MPLS (MultiProtocol Label Switching) [45] a été définie par l'IETF pour unifier le monde de la commutation de circuits et le monde de la commutation de paquets en utilisant une commutation de labels. Ceci signifie que dans un réseau MPLS, le routage des paquets se base sur des labels qui sont inclus dans leur en-tête.

MPLS se caractérise par son indépendance par rapport aux protocoles de niveau 2 (ATM, Frame Relay, etc.) et aux protocoles de niveau 3 (IP). Pour son fonctionnement il se base sur trois notions importantes :

- **FEC (Forwarding Equivalence Class)** : Une FEC regroupe un certain nombre de paquets possédant des caractéristiques communes qui déterminent leur cheminement dans le réseau. MPLS permet la constitution des FEC selon de nombreux critères : destination, source, QoS, application, etc. L'assignation d'un paquet à une FEC est effectuée une seule fois à l'entrée du réseau.
- **Le label MPLS** : Il représente un petit en-tête de paquet, appelé MPLS *Shim*. Son format dépend principalement du réseau sur lequel MPLS est mis en œuvre. Le label MPLS contient principalement un champ label de 20 bits, utilisé pour identifier une FEC sur un routeur particulier. Une FEC peut être identifiée par différents labels sur différents routeurs.
- **Le LSP (Label Switch Path)** : Il représente le chemin d'une FEC dans un réseau MPLS. Ce chemin est constitué de succession de labels attribués par chaque routeur pour une FEC.

L'architecture MPLS est constituée de deux types de routeurs : LSR (Label Switch Router) et LER (Label Edge Router). Le LSR est localisé à l'intérieur d'un réseau MPLS, il est capable de commuter les paquets suivant leur label. Tandis que le LER est localisé à l'extrémité du réseau, il se charge de l'assignation des labels aux paquets entrants et de leur suppression sur les paquets sortants.

Chaque routeur doit gérer deux tables : une table de routage et une table de labels. La table de routage est construite en utilisant les protocoles de routage IP standards (RIP, OSPF, BGP, etc.). En ce qui concerne la table de labels, sur laquelle repose la commutation, l'IETF propose deux solutions, soit l'utilisation d'un protocole de distribution de label spécifique LDP (Label Distribution Protocol) [46], soit l'extension des protocoles existant pour assurer la distribution de labels, par exemple RSVP [47].

Le principe de fonctionnement d'un réseau MPLS peut être résumé comme suit : lorsqu'un paquet entre dans un réseau MPLS, un routeur LER l'affecte à une FEC selon la politique déployée dans le réseau en lui assignant un label local de cette FEC. Le LER transmet le paquet au prochain routeur suivant sa table de labels. Les LSRs font suivre le paquet suivant son label en lui affectant à

chaque fois le label local qui identifie la FEC auquel il appartient. Ceci se répète, jusqu'au routeur LER de sortie qui supprime le label sur le paquet.

Une stratégie de QoS peut être mise en œuvre sur un réseau MPLS en séparant les files d'attente des FECs sur les routeurs et en configurant chaque FEC pour qu'elle garantisse certains critères de performance. Plusieurs travaux ont été menés dans cette direction. Les travaux dans [48] présentent une analyse des performances des flux UDP et TCP en utilisant MPLS. Les simulations ont montré une augmentation du débit en utilisant l'ingénierie de trafic de MPLS. Dans [49], un nouveau schéma de routage basé sur MPLS est développé, afin de donner une priorité au trafic multimédia, plus sensible au délai de transmission. Enfin, dans [50], les auteurs proposent un modèle de QoS pour les réseaux de nouvelle génération basé sur MPLS. Ils présentent aussi une implémentation de leur modèle en détaillant les principaux mécanismes de QoS à mettre en œuvre : le contrôle d'admission, le mapping des services, le contrôle du trafic, l'ordonnancement des files d'attente et en fin le routage respectant la QoS.

2.2.3.4 La négociation de la QoS de bout-en-bout

Dans les sections précédentes, nous avons présenté les architectures principales de QoS qui peuvent être déployées au niveau IP. Chaque domaine réseau peut mettre en œuvre n'importe quelle architecture de QoS en fonction de ses équipements et de ses besoins. Cependant, pour assurer la QoS de bout-en-bout d'un trafic qui peut traverser plusieurs domaines indépendants, il est impératif de définir un langage commun entre ces différents domaines [51]. Ce langage correspond à un protocole de négociation, ou de signalisation, qui permet de négocier dynamiquement la QoS de bout-en-bout d'un trafic IP.

Dans la majorité des cas, les protocoles de négociation de QoS emploient un contrat de service, appelé SLA (Service Level Agreement) [52]. Ce dernier définit d'une manière formelle tous les aspects du service négocié. Le SLA est accompagné d'une partie technique, appelée SLS (Service Level Specification), qui définit les paramètres négociés entre les différents domaines IP afin de se mettre d'accord sur un niveau de QoS. Les principaux paramètres utilisés dans un SLA sont :

- **Temps de service** : Spécifie le temps pendant lequel la QoS négociée doit être garantie.
- **Scope** : Point d'entrée et de sortie d'un domaine.
- **Paramètres de QoS** : Délai, gigue, taux de perte, bande passante.
- **Profil du trafic** : Taille des paquets, débit crête, etc.
- **Traitement d'excès** : Spécifie le traitement que le réseau applique aux paquets hors profils.
- **Identification du trafic** : Adresse IP source et destination, port source et destination, identification du protocole...
- **Identification du client** : Utilisée pour les fonctions de AAA (Authentication, Authorization, Accounting).
- **Mode de négociation** : Prédéfini ou non, c'est-à-dire que les paramètres de SLS sont contraints par le fournisseur ou prennent des valeurs quelconques.
- **Intervalle de renégociation** : L'intervalle de temps pendant lequel un SLS négocié ne peut être renégocié.
- **Fiabilité** : Temps d'inaccessibilité et temps de rétablissement d'un service suite à une panne.

Durant ces dernières années, nous avons assisté à l'apparition de plusieurs protocoles de négociations de QoS. Une brève présentation des ces protocoles est donnée ci-dessous :

2.2.3.4.1 *Le protocole COPS-SLS*

COPS-SLS [53] est une extension du protocole de contrôle par politique COPS (*Common Open Policy Service*) [54]. COPS-SLS permet une négociation « out-of-band » de SLS en se basant sur l'architecture de gestion par politiques. Dans l'architecture COPS-SLS, chaque domaine possède deux entités principales : le SLS-PDP et le SLS-PEP. Le SLS-PDP fournit les paramètres à négocier au SLS-PEP qui se charge de les transmettre au SLS-PDP d'un autre domaine. La négociation se déroule en trois étapes : (1) l'envoi d'un message REQ (Request) de proche en proche, du domaine initiateur jusqu'au domaine destinataire, pour demander des ressources, (2) l'envoi de la réponse DEC (Decision) dans le sens inverse pour informer les différents domaines des décisions, enfin, (3) l'envoi du rapport RPT (Report State) pour informer sur le succès ou l'échec de la négociation.

2.2.3.4.2 *Le protocole SrNP*

Le protocole SrNP (Service Negotiation Protocol) a été défini dans le cadre du projet TEQUILLA [55] pour la négociation de SSS (Service Subscription Structure) qui représente un ensemble de SLS. SrNP se base sur une architecture client/serveur et les entités de négociation ne sont que des serveurs SrNP et des clients SrNP. Ces deux dernières disposent d'un certain nombre de requêtes (SessionInit, Proposal, LastProposal, AcceptToHold, Revision, LastRevision, AgreedProposal, ProposalOnHold) qui permettent d'établir un accord, de modifier un accord déjà établi et d'annuler un accord déjà négocié [56].

2.2.3.4.3 *Le protocole DSNP*

Le protocole DSNP (Dynamic Service Negotiation Protocol) [57] est un protocole de négociation de SLS au niveau de la couche IP. La simplicité et la légèreté du protocole DSNP permet son utilisation dans les réseaux sans fil. De plus, il prend en charge des aspects de mobilité pour la négociation de la QoS. L'architecture DSNP repose aussi sur un client DSNP et un serveur DSNP qui s'échangent un certain nombre de requêtes (SLS_List_Request/Response, SLS_Nego_Request/Response, SLS_Stat_Request/Response) pour réaliser la négociation.

2.2.3.4.4 *Le protocole NSLP*

Le protocole NSLP [58] est issu des travaux du groupe de travail NSIS (Next Steps In Signaling) de l'IETF. L'objectif de ce groupe de travail est de standardiser un protocole de niveau IP pour la signalisation de la QoS. Les protocoles qui existent actuellement sont utilisés comme base de départ pour définir ce nouveau standard. Les premiers résultats de ce groupe est un modèle structurel composé de deux couches [59]: une couche inférieure, appelée NTLP (Signaling Transport Layer Protocol), et une couche supérieure, appelée NSLP (Network Signaling Layer Protocol). NTLP est responsable de l'acheminement des messages de signalisation dans le réseau, tandis que NSLP contient des fonctionnalités spécifiques à l'application comme le format et les séquences des messages. NSLP fournit des fonctionnalités similaires à RSVP avec plusieurs extensions. De plus, il est indépendant de l'architecture de QoS déployée sur le réseau.

2.2.3.5 Discussion

Actuellement, le protocole IP est considéré comme la technologie de prédilection pour réaliser la convergence des réseaux. Dans cette perspective, la couche IP est considérée comme une couche commune qui permettra le transport de n'importe quel service (Internet, téléphonie, télévision) sur n'importe quelle technologie réseau.

Le succès commercial de ces réseaux IP de nouvelle génération dépend principalement de la QoS qu'ils peuvent offrir. La continuité horizontale et verticale de la QoS est très importante vu la nature hétérogène de ces nouvelles architectures. Nous avons présenté dans les sections précédentes plusieurs protocoles permettant la négociation de la QoS inter-domaines afin d'assurer la continuité horizontale de la QoS de bout-en-bout. Cependant, les efforts fournis pour la continuité verticale de la QoS IP, entre les couches supérieures et les couches inférieures, sont beaucoup moins importants. Parmi ces efforts, nous pouvons citer le protocole de signalisation SBM (Subnet Bandwidth Manager) [60] qui définit une communication entre les routeurs et les commutateurs d'un réseau afin de permettre un mapping entre la QoS du niveau IP et la QoS du niveau 2 des réseaux IEEE 802. Ce protocole décrit principalement les fonctionnalités des nœuds réseau supportant RSVP et des commutateurs LAN afin de permettre une réservation de ressources au niveau 2 aux flux prioritaires. Pour cela, SBM considère l'utilisation de commutateurs mettant en œuvre des files d'attente prioritaires, comme cela est spécifié dans le standard IEEE 802.1p [61].

Ainsi, SBM représente un premier pas vers une communication inter-couches pour permettre une continuité verticale de la QoS. Ce genre de protocole doit être investi plus profondément pour une plus grande interopérabilité entre la QoS IP et les différentes QoS des couches inférieures (WiFi, WiMax, xDSL, FTT, etc.). Plusieurs exemples peuvent être cités pour démontrer le besoins de ces protocoles, parmi eux : le service garanti au niveau IP (IntServ) et la variation du débit physique dans les réseaux sans fil (voir section 2.2.1.3). En effet, le débit garanti par la couche IP peut ne pas être garanti au niveau physique à cause de la mobilité de l'utilisateur ou de la dégradation de l'environnement physique (voir section 2.2.1.2). La communication avec les couches supérieures est tout aussi importante. Les couches supérieures et principalement les couche applicatives doivent avoir connaissance des mécanismes de QoS déployés sur la couche IP afin d'exécuter les réservations nécessaires, ou, de choisir les classes de service à utiliser suivant leurs besoins.

2.2.4 La couche transport

La fonctionnalité principale de la couche transport est le transfert des messages de bout-en-bout de l'émetteur vers le récepteur. Pour cela, elle emploie des protocoles de transport qui utilisent les services de la couche réseau (IP) pour assurer l'acheminement des messages. Les deux premiers protocoles définis et largement déployés dans les réseaux IP sont : UDP (Usage Datagram Protocol) [62] et TCP (Transmission Control Protocol) [63].

Le protocole TCP offre un service de transport fiable, orienté connexion afin d'assurer un transfert dans l'ordre et sans erreurs des données applicatives. Les erreurs sur un segment TCP sont détectées grâce à l'utilisation d'un *Checksum* sur l'en-tête et les données TCP. Pour les pertes de segments, TCP se base sur un mécanisme ARQ qui emploie les acquittements, les temporisateurs et

les retransmissions en cas de perte. Ce mécanisme assure la bonne réception de chaque octet de données présent dans un segment. Durant la transmission des données, TCP emploie deux mécanismes importants : le contrôle de flux et le contrôle de congestion. Le premier assure une corrélation entre le débit d'émission et le débit de réception. Tandis que le deuxième, assure une corrélation entre le débit d'émission et le débit disponible sur le réseau tout en garantissant une certaine équité pour le partage de ce débit entre plusieurs flux TCP.

Les flux TCP représentent la majorité des flux qui circulent actuellement sur les réseaux IP. Son succès s'explique principalement par sa fiabilité exigée par plusieurs applications réseaux, comme le transfert de fichiers, l'email, le web, etc. En effet, ce type d'application possède une tolérance nulle pour les pertes de données, mais supporte, avec une certaine limite, la gigue et la latence causées par la retransmission des segments perdus.

Ceci n'est pas tout à fait le cas pour les applications multimédia dont les flux possèdent des caractéristiques plus au moins incompatibles avec le mode de fonctionnement de TCP. Cette incompatibilité a été explorée dans plusieurs travaux [64] et se résume en deux points importants :

- **La fiabilité en utilisant la retransmission :** La retransmission des segments perdus introduit une latence (délai de retransmission + délai de réordonnancement) et une gigue insupportables pour les applications multimédia interactives, comme la vidéoconférence et la VOIP. La latence pour ce genre d'application doit être strictement bornée entre 100 ms et 200 ms, mais elle est moins stricte pour les applications de type VoD (Video on Demand) ou LoD (Live on Demand).
- **Le contrôle de congestion :** L'algorithme de contrôle de congestion utilisé par TCP permet de réduire le débit d'émission en cas de perte de segment, interprétée par une congestion dans le réseau, et de l'augmenter dans le cas contraire. Pour faire face à la variation du débit, les applications multimédia doivent utiliser des buffers ou des mécanismes d'adaptation au niveau applicatif.

Le deuxième protocole largement déployé actuellement est le protocole UDP qui, à l'opposé de TCP, n'offre aucun service de fiabilité. UDP fonctionne en mode non connecté, il ne garantit par la bonne réception des messages et n'implémente aucun mécanisme de contrôle de flux ni de contrôle de congestion. Par conséquent, UDP est adapté pour le transport des flux multimédia puisqu'il n'introduit aucun délai ni gigue supplémentaires durant la transmission des paquets. En effet, le protocole UDP a été largement adopté par les applications multimédia pour le transport des flux audio et vidéo en reléguant d'autres fonctionnalités nécessaires au niveau applicatif, comme le séquençement et la synchronisation des paquets.

Cependant, les flux UDP sont caractérisés par leur agressivité vis-à-vis du réseau puisqu'ils n'implémentent aucun mécanisme de contrôle de flux ni de contrôle de congestion. Ceci engendre une iniquité dans le partage du débit réseau entre les flux UDP et TCP et entre les flux UDP eux-mêmes. De plus, les flux multimédia subissent une dégradation de la QoS à cause de la suppression des paquets au niveau des routeurs.

Pour palier à tous ces inconvénients, nous avons assisté ces dernières années à l'apparition de divers algorithmes de contrôle de congestion adaptés au flux multimédia et compatibles avec les algorithmes de contrôle de congestion proposés par TCP.

En parallèle à ce développement, nous avons assisté à l'apparition de deux nouveaux protocoles de transport : DCCP (Datagram Congestion Control Protocol) et SCTP (Stream Control Transmission Protocol) qui fournissent un compromis entre la rigidité de TCP et la simplicité d'UDP pour le transport des flux multimédia.

Dans le reste de cette section, nous présentons une taxonomie des algorithmes de contrôle de congestion en détaillant les algorithmes spécifiques aux flux multimédia. Par la suite, nous décrivons les deux protocoles DCCP et SCTP.

2.2.4.1 Les algorithmes de contrôle de congestion pour les flux multimédia

Un Algorithme de Contrôle de Congestion (ACC) peut être exécuté au niveau transport ou au niveau applicatif pour contrôler le débit d'émission ou le débit de réception d'un flux. Ceci réduit la congestion tout en assurant un partage équitable du débit disponible dans le réseau.

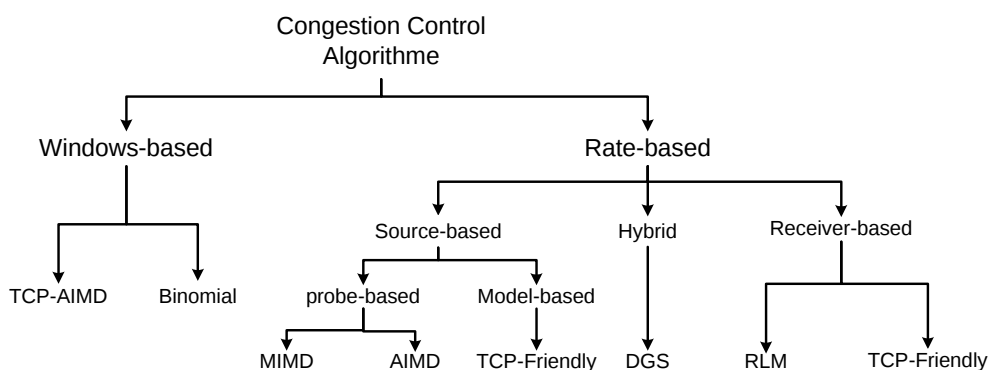


Figure 2-7 : Taxonomie des algorithmes de contrôle de congestion

La Figure 2-7 présente une taxonomie des ACC qui ont été proposés durant ces dernières années. Il existe deux types d'ACC [65][66]:

- **Les ACC basés sur une fenêtre (window-based) :** Cette catégorie d'algorithmes sonde le débit disponible dans le réseau en faisant varier la taille de la fenêtre de congestion en fonction des conditions de transmission. Le plus connu de ces algorithmes est celui implémenté par TCP. La taille de la fenêtre de congestion de TCP suit un algorithme AIMD (Additive-Increase / Multiplicative-Decrease). Cet algorithme est exécuté côté émetteur et il propose de démarrer la transmission avec une fenêtre contenant un segment. À Chaque réception d'un acquittement, la taille de la fenêtre augmente d'un segment jusqu'à une taille limite. Cette phase est appelée démarrage lent. Par la suite, l'algorithme passe dans la phase d'évitement de congestion dans laquelle la taille de la fenêtre augmente par un segment chaque RTT (Round Trip Time). Cependant, dans le cas où l'émetteur ne reçoit pas d'acquittement, la taille de la fenêtre est divisée par deux. Il existe d'autres algorithmes qui appartiennent à cette catégorie, comme les algorithmes binomiaux [67] qui proposent une variation non linéaire de la taille de la fenêtre de congestion mieux adaptée aux flux multimédia.

- **Les ACC basés sur un débit (rate-based) :** Dans cette catégorie d'algorithmes, le débit d'émission se base sur une estimation du débit disponible sur le réseau. Dans la majorité des cas, ils sont implémentés par les applications multimédia directement au niveau applicatif puisque ces applications utilisent le protocole UDP au niveau transport.

Dans ce qui suit nous nous intéressons uniquement à la deuxième catégorie conçue principalement pour les flux multimédia. Cette catégorie est subdivisée en trois grandes classes d'algorithmes suivant l'emplacement de leur exécution [65][66]: (1) le contrôle de débit côté source, (2) le contrôle de débit côté récepteur et (3) le contrôle de débit hybride.

2.2.4.1.1 Le contrôle de débit côté émetteur (source-based)

Dans ce genre d'algorithme, l'émetteur est responsable de l'estimation du débit du réseau en se basant sur des retours d'informations (feedback) sur les conditions de transmission. Ils sont généralement utilisés pour des transmissions unicast. L'estimation du débit peut s'effectuer suivant deux approches : en utilisant des sondes (généralement le taux de perte des paquets) ou bien en utilisant des modèles prédéfinis.

Dans la première approche, les taux de perte sont utilisés comme critère de performances. L'émetteur fait varier le débit d'émission pour maintenir un taux de perte inférieur à un seuil qui maintient la QoS du flux. La variation du débit peut suivre un algorithme de type AIMD [68] ou bien un algorithme MIMD [69] (Multiplicative Increase / Multiplicative Decrease). Dans le système de vidéoconférence de l'INRIA (IVS : INRIA Video-conference System) [70], un algorithme de contrôle de débit, basé sur des sondes, est proposé pour les transmissions multicast. Suivant le taux de perte subi, chaque récepteur détermine l'état du réseau qui peut être soit UNLOADED, LOADED, CONGESTED. Ensuite, l'émetteur récupère ces états sur un échantillon limité pour éviter d'être surchargé par les réponses de tous les récepteurs. À partir de cet échantillon, il décide d'augmenter le débit ou bien de le diminuer suivant un algorithme AIMD.

En ce qui concerne l'approche basée sur un modèle, le débit d'émission est estimé en utilisant un modèle prédéfini, dans la majorité des cas, à partir d'une connexion TCP. Ces modèles se basent sur plusieurs paramètres, comme le taux de perte, le RTT et le MTU (Maximum Transit Unit), pour estimer le débit [71]. Par conséquent, ce type d'algorithme évite la congestion d'une manière similaire à TCP ce qu'il lui a valu le nom de « TCP-friendly ». Le plus célèbre de ces algorithmes est le TFRC (TCP-friendly Rate Control) [72] défini par l'IETF dans le RFC 3448. TFRC propose un modèle mathématique de contrôle de congestion compatible avec TCP, mais dont la variation du débit, durant le temps, est beaucoup moins importante. Pour cela, le débit est calculé suivant l'équation Eq. 2-1 donnée ci-dessous :

$$R_{TCP} \cong \frac{s}{RTT \sqrt{\frac{2bp}{3}} + t_{RTO} (3 \sqrt{\frac{3bp}{8}}) p (1 + 32p^2)} \quad \text{Eq. 2-1}$$

R_{TCP} représente le débit estimé en octets/secondes, s est la taille de paquet en octets, RTT est le temps d'aller retour, p est le taux de perte entre [0, 1], T_{RTO} est le temps du timeout pour effectuer une retransmission, b est le nombre de paquets acquittés par un seul acquittement. Une variante de cette algorithme, appelée TFRC-SP (Small Packet) [73] a été proposée récemment par l'IETF pour

les applications utilisant des paquets de petite taille, comme les applications de VoIP. L'objectif principal de TFRC-SP est de fournir à ces applications un débit similaire à celui qu'offre TCP en utilisant une taille de paquet de 1500 octets.

2.2.4.1.2 *Le contrôle de débit côté récepteur (Receiver-based)*

Ce type d'algorithmes est généralement utilisé pour les transmissions multicast de flux vidéo codés en couche. Le codage en couche d'une vidéo, permet de structurer cette dernière en une couche de base et plusieurs couches d'amélioration. La couche de base peut être décodée sans les couches d'amélioration, mais elle offre une qualité vidéo limitée. L'ajout d'une ou de plusieurs couches d'amélioration durant le décodage permet d'augmenter cette qualité. Pour permettre à des récepteurs hétérogènes (débits de réception différents) de recevoir le flux vidéo, chaque couche est transmise sur un groupe multicast. Le récepteur décide, par la suite, des couches qu'il veut recevoir en s'abonnant ou pas au groupe multicast. Pour cela, le récepteur doit effectuer une estimation judicieuse de son débit de réception.

Plusieurs travaux ont été menés dans cette direction. Le RLM (Receiver-driven layered multicast), présenté dans [74], propose un algorithme simple pour décider de rejoindre ou de quitter un groupe multicast. Si le récepteur ne détecte pas de congestion (pertes de paquets), il rejoint un nouveau groupe. Dans le cas contraire, il quitte un groupe. RLM propose aussi un algorithme d'apprentissage qui permet au récepteur de partager leurs expériences (historiques) sur les abonnements aux groupes qui ont réussi (pas de pertes après avoir rejoint un groupe) et les abonnements qui ont échoué (détection de perte juste après avoir rejoint un groupe). D'une manière similaire, l'algorithme proposé dans [75], est un algorithme « TCP-friendly » qui exploite la forte corrélation entre le débit disponible et les taux de perte dans TCP et qui est simplement modélisée par l'équation Eq. 2-2.

$$T \propto 1/P \quad (T : \text{débit}, P : \text{taux de perte}) \quad \text{Eq. 2-2}$$

Pour permettre une certaine coordination dans la procédure d'abonnement et de désabonnement des récepteurs, les auteurs proposent une procédure de synchronisation basée sur un champ d'information inclus directement dans les paquets de données. Ces informations permettent aux récepteurs de partager leur expérience des abonnements aux groupes.

2.2.4.1.3 *Le contrôle de débit hybride (Hybrid)*

Dans cette catégorie d'algorithmes, le contrôle du débit est effectué par l'émetteur et le récepteur en même temps. Le récepteur décide du nombre de flux qu'il peut recevoir en se basant sur une estimation de son débit de réception. Tandis que l'émetteur ajuste continuellement le débit d'émission de chaque flux indépendamment des autres. Un algorithme de ce type, appelé DGS (Destination Set Grouping) a été présenté dans [76]. DGS se base sur le principe de IVS [70] présenté dans la section 2.2.4.1.1. Cependant, dans DGS le récepteur aussi s'abonne et se désabonne des groupes en fonction de l'état du réseau, déterminé à partir du taux de perte.

2.2.4.2 **Le protocole SCTP**

Le protocole SCTP [77] (Stream Control Transmission Protocol) est un protocole de transport défini à la base, par l'IETF, pour le transfert de signalisation dans un environnement VoIP. Il a été

conçu principalement comme une alternative à TCP qui est caractérisé par sa rigidité (notions de flux, séquençement stricte des données transmises, etc.). SCTP offre un service de transport fiable en mode connecté et emploie le même algorithme de contrôle de congestion que TCP. Cependant, il se distingue par de nouvelles fonctionnalités décrites ci-dessous :

- **L'établissement et la fermeture d'une association** : Une connexion entre deux nœuds SCTP est appelée une association. Une association dans SCTP s'établit en quatre temps afin d'éviter le problème du « SYN Attacks » dans TCP. De plus, SCTP ne supporte pas une association ouverte à moitié, où un seul nœud peut transmettre. Lorsque l'association est fermée entre deux nœuds, aucun des deux n'accepte une réception de données.
- **Multihoming** : La possibilité pour une association SCTP de supporter plusieurs adresses IP. Ceci permet d'augmenter la fiabilité d'une association au cas où une adresse deviendrait inaccessible.
- **Multi-streaming** : Contrairement à TCP, un flux SCTP représente une suite de messages de différentes tailles et non pas une suite d'octets. Une association peut supporter plusieurs flux distingués par des identifiants. Les messages appartenant au même flux sont livrés en séquence. L'identifiant du flux et le numéro de séquence sont inclus dans l'en-tête du message.
- **Fragmentation des messages** : SCTP offre un mécanisme de fragmentation/réassemblage des messages de l'utilisateur en fonction du MTU du réseau.
- **Regroupement des messages** : SCTP permet de regrouper plusieurs messages de l'utilisateur dans un seul paquet SCTP

Une extension de SCTP, appelé PR-SCTP [78] (Partial-Reliable), a été proposée par l'IETF pour le transport de données temps réel dont le temps de vie est limité. PR-SCTP offre un service partiellement fiable qui permet à l'utilisateur de définir la politique de fiabilité pour le transport des données. Par exemple, une fiabilité limitée dans le temps permet à l'utilisateur d'indiquer une durée de temps durant laquelle un message est transmis ou retransmis.

Plusieurs travaux se sont intéressés aux performances de SCTP et de PR-SCTP pour la transmission des flux multimédia. Dans [79], les auteurs proposent une nouvelle architecture pour le streaming unicast de vidéo H.264/AVC basée sur SCTP. Dans cette architecture, chaque type NALU (Network Abstraction Layer Unit) de H.264 est mappé sur un flux SCTP possédant une fiabilité spécifique, à savoir le nombre de retransmissions possibles. Ceci permet de mieux protéger les données importantes. La même idée est reprise dans [80] pour les flux MPEG-4. Enfin, dans [81], les auteurs démontrent que l'utilisation de PR-SCTP pour la transmission de flux IPTV fournit de meilleures performances que l'utilisation de TCP et d'UDP.

2.2.4.3 Le protocole DCCP

Le protocole DCCP [82] (Datagram Congestion Control Protocol) est un nouveau protocole de transport, orienté message, défini par l'IETF. Il offre un service de transport non fiable mais en mode connecté et il implémente plusieurs algorithmes de contrôle de congestion. DCCP représente un compromis entre TCP et UDP pour le transport des flux multimédia puisqu'il n'introduit pas de délai ni de gigue supplémentaires, causés par les retransmissions dans TCP, et il implémente un

algorithme de contrôle de congestion, à l'opposé d'UDP. DCCP se distingue par les mécanismes suivants :

- Etablissement fiable de la connexion avant le début de la transmission.
- Négociation fiable des options qui vont être utilisées durant la transmission. Par exemple, l'algorithme de contrôle de congestion choisi par les deux extrémités.
- Acquiescement systématique des paquets qui sont bien reçus, mais pas de retransmission des paquets qui sont perdus. Les informations des acquiescements sont exploitées par les algorithmes de contrôle de congestion côté émetteur.
- Le contrôle de congestion prend en considération l'ECN (Explicit Congestion notification) qui permet au nœud réseau de signaler explicitement une congestion au protocole de transport. Ceci permet de différencier les types de perte que peuvent subir les paquets (suppression dans un routeur ou bien corruption des données).
- Intégration de plusieurs algorithmes de contrôle de congestion. La spécification de DCCP propose pour l'instant deux algorithmes : Le CCID 2 (Congestion Control IDentifier) qui correspond à l'algorithme AIMD utilisé par TCP et le CCID 3 qui correspond à l'algorithme TFRC [72].

DCCP a été implémenté à partir de la version 2.6.14 du noyau linux, apparue en Octobre 2005, et il est continuellement mis à jour dans les nouvelles versions. DCCP fait actuellement l'objet de plusieurs travaux qui s'intéressent particulièrement à ses performances pour le transport des flux multimédia en comparaison avec les protocoles TCP et UDP, mais aussi à sa cohabitation avec les autres protocoles de transport. Par exemple, les travaux dans [83], montrent une iniquité entre les flux DCCP (avec CCID 2) et les flux TCP causée par les retransmissions de TCP. De même, les auteurs dans [84] comparent les performances de TCP et DCCP (avec CCID 3) pour le streaming vidéo basé sur H.264/SVC. Ce dernier est choisi pour ses capacités d'adaptation spatiale, temporelle et SNR (Signal to Noise Ratio). Les résultats montrent que l'adaptation basée sur DCCP offre une meilleure qualité vidéo avec une petite taille du buffer de réception.

2.2.4.4 Discussion

Malgré l'apparition de nouveaux protocoles mieux adaptés aux transports des flux multimédia, nous constatons que les protocoles TCP et UDP restent majoritairement utilisés. Ceci peut s'expliquer par le fait que la couche transport est implémentée au niveau des systèmes d'exploitation qui se sont limités, durant le passé, aux protocoles TCP et UDP. Maintenant que les nouveaux protocoles commencent à être intégrés dans les systèmes d'exploitation, il faudra attendre une réelle volonté de la part des développeurs d'applications et de services multimédia pour les exploiter. Donc, dans le cadre de notre étude, nous nous limitons aux protocoles TCP et UDP largement déployés actuellement sur les réseaux sans fil 802.11.

En ce qui concerne TCP, l'objectif de son algorithme de contrôle de congestion est de diminuer la congestion du réseau puisque chaque perte de paquet est interprétée par sa suppression au niveau des nœuds réseaux. Cependant, cette interprétation n'est pas exacte dans les réseaux sans

fil [85] où les pertes de paquets peuvent être causées par des interférences. Dans ce cas, la diminution de moitié de la fenêtre de congestion est un peu drastique.

Ainsi, une distinction entre les deux types de perte est nécessaire. De plus, le fonctionnement de l'algorithme de contrôle de congestion doit être revu pour les réseaux sans fil en prenant en considération d'autres paramètres : le débit disponible au niveau physique et le taux d'erreur sur le canal (BER). La même remarque peut être faite pour les algorithmes de contrôle de congestion présentés dans la section 2.2.4.1 qui sont généralement utilisés avec UDP.

D'autre part, la réduction ou l'augmentation du débit décidée par ces algorithmes doit être partagée avec les couches applicatives. Afin de permettre aux applications multimédia d'exécuter des mécanismes d'adaptation qui feront correspondre le débit des flux audio/vidéo au débit réseau.

Enfin, la redondance de la retransmission entre la couche MAC dans le standard 802.11 et la couche transport au niveau du protocole TCP doit être considérée. En effet, le standard 802.11 introduit au niveau MAC un mécanisme ARQ qui permet la retransmission des trames perdues dans le cas où l'acquittement de la trame n'est pas reçu (voir section 2.2.2.2). Dans le cas d'un réseau d'accès 802.11, ce mécanisme de retransmission et la retransmission au niveau TCP représente une redondance de service entre les couches. Il serait intéressant d'avoir une certaine collaboration entre ces deux mécanismes pour optimiser le système et réduire le délai de transmission.

2.2.5 La couche Application

La couche application est la dernière couche du modèle TCP/IP. Elle comprend un ensemble de protocoles de haut niveau développés pour de nombreuses applications : la messagerie électronique (SMTP), le transfert de fichiers (FTP), le Web (HTTP) et le streaming (RTSP, RTP/RTCP).

Dans le cadre de notre étude, nous nous intéressons uniquement aux protocoles et aux applications de streaming vidéo qui, à l'opposé du téléchargement vidéo, permettent la visualisation d'un flux vidéo au fur et à mesure de sa réception. Les applications de streaming permettent la diffusion des flux vidéo temps réel (capture caméra, flux TV) ou bien des flux stockées sur des serveurs de contenus (films, documentaires, etc.). Le streaming sur les réseaux IP se base principalement sur deux modes de transmission : l'unicast et le multicast. Le premier mode nécessite la duplication d'un flux vidéo pour chaque récepteur, il est généralement utilisé pour des services VoD (Video on Demand) ou LoD (Live on Demand). En ce qui concerne le multicast, un flux est transmis pour un groupe de récepteur. Il est utilisé pour la diffusion en continu de flux temps réel et pour les services de communication de groupe (vidéoconférence).

Actuellement, le streaming vidéo sur les réseaux IP, filaires et sans fil, est devenu une réalité. Ceci s'explique par le développement de protocoles multimédia qui permettent la transmission et le contrôle de flux vidéo ainsi que le développement de codecs performants qui fournissent une qualité vidéo satisfaisante avec des débits relativement bas.

Cependant, les applications de streaming souffrent toujours du manque de la QoS dans les réseaux IP qui affecte la transmission des paquets et dégrade, par la même, la qualité de la vidéo

perçue par le récepteur. De plus, l'hétérogénéité grandissante dans les réseaux IP oblige ces applications à s'adapter à plusieurs paramètres, comme le débit du réseau d'accès, les capacités du terminal récepteur, etc.

Dans ce qui suit, nous décrivons, en premier, les principaux protocoles applicatifs qui sont utilisés par les applications multimédia. Ensuite, nous présentons les différents codecs qui ont permis la numérisation de la vidéo et sa transmission sur les réseaux IP. Enfin, nous détaillons les différents mécanismes de QoS déployés au niveau applicatif qui permettent de s'adapter à certaines caractéristiques réseaux, comme le débit disponibles, mais aussi de faire face à d'autres caractéristiques, comme la perte de paquets, la gigue et le délai de transmission.

2.2.5.1 Les protocoles multimédia

Le rôle principal des protocoles multimédia est de fournir des fonctionnalités de base aux applications pour le contrôle, la description et la transmission de flux multimédia. La majorité de ces protocoles sont définis par l'IETF, sauf la suite protocolaire du H.323 qui a été définie par l'ITU-T (International Telecommunication Union - Telecommunications). Une brève description de chaque protocole est présentée ci-dessous :

- H.323 [86][87] : La première version du H.323 a été publiée par l'ITU-T en 1996. Elle définissait un ensemble de protocoles et d'architectures pour la communication audio/vidéo/données sur les réseaux IP en mode centralisé. Parmi ces protocoles, nous avons le protocole H.225.0/Q.931 et le protocole H.225.0/RAS (Registration, Admission, and Status) pour le contrôle d'une session de communication, le protocole H.245 pour le contrôle des flux audio/vidéo entre entités communicantes et enfin la suite de protocoles T.120 pour la transmission de données.
- SIP (Session Initiation Protocol) [88] : Le protocole SIP est un protocole de signalisation défini par l'IETF pour contrer H.323, caractérisé par sa complexité. Il permet l'établissement, la modification et la libération de sessions multimédia. Une session multimédia peut représenter un appel téléphonique sur IP, une conférence multimédia ou une distribution de contenus multimédia. Pour cela, il se base sur un jeu de transactions requête/réponse codé en format texte (INVITE, REGISTER, BYE, CANCEL, OPTIONS). L'architecture SIP est constituée de plusieurs entités : l'agent utilisateur, le serveur proxy, le serveur d'enregistrement/localisation, le serveur de redirection. Chaque entité possède une fonctionnalité particulière dans l'offre de service multimédia.
- RTSP (Real Time Streaming Protocol) [89] : RTSP est un protocole de niveau applicatif qui permet l'établissement et le contrôle d'une session multimédia entre un client et un serveur afin de transmettre un, ou plusieurs, flux audio/vidéo en unicast ou en multicast. À l'instar de SIP, RTSP se base aussi sur un jeu de transactions requête/réponse, codé en format texte, qui offrent des fonctionnalités similaires à celle d'un magnétoscope (OPTION, DESCRIBE, SETUP, PLAY, PAUSE, STOP, TREADOWN). Ces requêtes permettent au client de demander et de contrôler la transmission d'un flux multimédia présent sur un serveur.

- **SDP (Session Description Protocol) [90]:** SDP définit un format pour la description d'une session multimédia. La description inclut plusieurs informations, comme le nom de la session, sa durée, les flux multimédia qu'elle contient et les informations nécessaires pour recevoir ces flux. Le format de cette description est constitué de plusieurs lignes de texte qui suivent le modèle <type>=<valeur>. SDP peut être utilisé par n'importe quel protocole de contrôle ou d'annonce de session multimédia, comme RTSP, SIP ou SAP.
- **SAP (Session Announcement Protocol) [91] :** le protocole SAP est utilisé pour annoncer une session multimédia multicast et pour transmettre la description de cette session aux futurs participants. Pour cela, des annonces SAP sont envoyées périodiquement sur une adresse multicast et un numéro de port bien définis. Les annonces SAP sont constituées d'un en-tête authentique et d'une description SDP et elles sont transmises sur UDP.
- **RTP/RTCP (Real Time Protocol/Real Time Control Protocol) [92]:** RTP est un protocole de niveau applicatif qui offre des fonctions de transport de bout-en-bout pour les flux multimédia. Parmi ces fonctions, nous citons l'identification d'une source multimédia, la numérotation des paquets et la synchronisation des paquets. RTP n'offre aucun service de fiabilité et il est indépendant du protocole de transport. Pour contrôler la transmission des paquets RTP, le protocole RTCP est utilisé. Ce dernier définit cinq types de rapports (SR, RR, SDES, APP, BYE) qui peuvent être échangés périodiquement entre l'émetteur et le récepteur.

2.2.5.2 Les standards de codage vidéo et le codage vidéo hiérarchique

La majorité des standards de codage vidéo ont été développés par deux organismes de standardisation : MPEG (Moving Pictures Expert Group) de ISO/IEC et VCEG (Video Compression Expert Group) de l'ITU-T. Une brève description de ces standards est donnée ci-dessous :

- **H.261 [93]:** Défini en 1990 par le groupe VCEG, il a été utilisé principalement pour la vidéoconférence sur les réseaux ISDN (Integrated Services Digital Network).
- **H.263 [94]:** Défini en 1996 par le groupe VCEG, il se base sur l'architecture H.261 avec de nouveaux algorithmes pour améliorer les performances du codage. La version 2, appelée H.263+, a été ratifiée en 1998. Elle a permis d'élargir le domaine d'application en offrant plus de flexibilité et en améliorant l'efficacité du codage.
- **MPEG-1 [95]:** Publié par le groupe MPEG en 1991, il a été développé principalement pour le stockage des vidéos sur des supports numériques (CD-ROM) avec un débit vidéo de 1.5 Mbits/s
- **MPEG-2 / H.262 [96]:** Publié en 1994, il a été développé conjointement par le groupe MPEG et le groupe VCEG. Il permet une très grande flexibilité des formats et des débits vidéo élevés pour la HDTV (High-Definition Television) et la SDTV (Standard-Definition Television). Il a été adopté par plusieurs standards TV, à savoir l'ATSC (Advanced Television Systems Committee) en Amérique et le DVB (Digital Video Broadcast) en Europe. Il est aussi utilisé pour stocker des vidéos sur le support DVD (Digital Video Disc).

- **MPEG-4 Part-2** [97]: Publié en 2000 par le groupe MPEG, il représente le premier codec vidéo orienté objet développé principalement pour les applications multimédia interactives.
- **H.264 AVC (Advanced Video Coding) / MPEG-4 part-10** [98]: Publié en 2003 par la JVT (Joint Video Team) de MPEG et de VCEG, il offre une plus grande performance de codage, comparé au MPEG-2 et MPEG-4 et il vise diverses applications : la diffusion TV, le HD-DVD, le stockage numérique, et la TV mobile.

Actuellement, les efforts de la JVT sont orientés vers le codage vidéo hiérarchique SVC (Scalable Video Coding) [99] afin de répondre aux besoins des nouvelles applications multimédia qui doivent transmettre des flux vidéo sur des réseaux différents et pour des terminaux hétérogènes. Contrairement aux codecs précédents, qui génèrent un seul flux vidéo avec une seule couche, le SVC génère plusieurs flux correspondant à plusieurs couches hiérarchiques, une couche de base (BL : Base Layer) et une ou plusieurs couches d'amélioration (EL : Enhancement Layer). La couche de base se suffit à elle-même pour le décodage, mais le décodage des couches supérieures nécessite le décodage de la couche de base. La définition du nouveau standard SVC se base sur l'architecture H.264/AVC. Les couches hiérarchiques peuvent être construites sur trois dimensions :

- La hiérarchie temporelle : Correspond à différents nombres d'images par seconde (frame rate). La couche de base est constituée des images I (Intra-coded frame) et P (Predictatively coded frame) et les couches d'amélioration sont constituées d'images B (Bi-directionally predicted frame) insérées entre les images I et P.
- La hiérarchie spatiale : Correspond à différentes dimensions d'image. Les couches supérieures fournissent une plus grande taille d'image.
- La hiérarchie en qualité (SNR) : Correspond à différentes qualités d'images. Les couches supérieures permettent d'obtenir une qualité plus fine de l'image, avec plus de précision.

En tronquant les couches supérieures, les applications multimédia peuvent mieux s'adapter à différents paramètres. Par exemple, le débit disponible dans le réseau en utilisant la hiérarchie en qualité ou encore la capacité d'affichage d'un terminal en utilisant la hiérarchie spatiale.

2.2.5.3 La QoS applicative pour la transmission vidéo

La transmission des paquets vidéo doit considérer les paramètres principaux qui caractérisent les réseaux IP, à savoir le débit, les pertes de paquets, le délai de bout-en-bout et la gigue. Durant ces dernières années, plusieurs techniques ont été utilisées par les applications multimédia afin de palier aux variations de ces facteurs et de minimiser leurs effets sur la qualité de la vidéo perçue par le récepteur. La Figure 2-8 présente une taxonomie de ces mécanismes qui seront détaillés dans la suite de cette section.

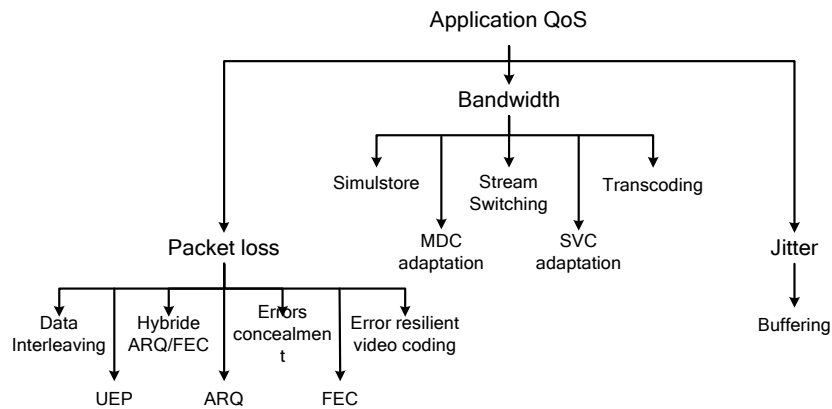


Figure 2-8 : Taxonomie de la QoS au niveau applicatif

2.2.5.3.1 L'adaptation au débit réseau (Bandwidth)

Ces mécanismes permettent principalement d'adapter le débit de la vidéo en fonction du débit disponible dans le réseau. Ce dernier est estimé grâce aux algorithmes de contrôle du débit présentés dans la section 2.2.4.1. Une brève description de chaque mécanisme est donnée ci-dessous :

- **Simulstore** [100] : La solution Simulstore propose de stocker sur un serveur plusieurs flux d'un même contenu avec différentes caractéristiques spatiales, temporelles et qualité (SNR). Le bon flux est choisi et transmis en fonction du débit de l'utilisateur. Cependant, cette solution est statique et ne peut pas s'adapter au changement dynamique du débit durant la transmission.
- **Stream switching** [101][102][103] : Le stream switching est une extension du simulstore pour permettre le basculement d'un flux vers un autre durant la transmission afin de répondre au changement dynamique du débit dans le réseau. Pour assurer la synchronisation des flux, le basculement est effectué sur les images I ou sur des images de basculement spéciale (SP-frame). Le principal inconvénient de cette solution est qu'elle nécessite un espace de stockage important.
- **Transcoding** [104][105][106] : Le transcoding permet de transformer une vidéo d'un format vers un autre en changeant la taille de l'image, le nombre d'images par seconde, la qualité de la vidéo ou le débit. Cette transformation peut être effectuée en gardant le même codage vidéo (de H.264 vers H.264) ou bien en basculant vers un autre format de codage (de MPEG-2 vers H.264). Cette solution ne nécessite pas un grand espace de stockage, par contre, elle ne supporte pas la montée en charge puisqu'elle consomme énormément de ressource de calcul pour exécuter le transcoding. De plus, elle introduit une latence supplémentaire qui peut être contraignante pour les services multimédia interactifs.
- **SVC adaptation (Scalable Video Coding)** [99] : Le codage hiérarchique, décrit dans la section 2.2.5.2, permet le codage de vidéo en plusieurs couches hiérarchiques, une couche de base et plusieurs couches d'amélioration. L'adaptation s'effectue simplement en supprimant une ou plusieurs couches d'amélioration. Plusieurs travaux se sont intéressés à ce type d'adaptation qui peut supporter la montée en charge [107][108][109]. L'inconvénient du codage hiérarchique est

la dépendance des couches supérieures des couches inférieures lors du décodage, ce qui nécessite la bonne réception des couches inférieures.

- **MDC adaptation (Multiple Description Coding)** : Avec le codage MDC, une vidéo est codée en plusieurs descriptions, ou flux, indépendants. Le MDC possède deux propriétés importantes : (1) Chaque flux peut être décodé indépendamment des autres en donnant une certaine qualité de vidéo, (2) les informations des flux sont complémentaires ce qui permet d'augmenter la qualité de la vidéo en augmentant le nombre de flux décodés simultanément. Plusieurs algorithmes de codage MDC ont été proposés par plusieurs travaux [110][111][112]. D'autres travaux [113][114][115][116] se sont intéressés à l'exploitation de ce type de codage dans des architectures de transmission vidéo.

2.2.5.3.2 La gestion des paquets perdus

Les pertes de paquets dégradent considérablement la qualité de la vidéo décodée à la réception. Cette dégradation est accentuée par l'effet de propagation d'erreurs [117] due à la dépendance du décodage des images I, P et B. En effet, la perte d'un paquet de l'image I provoque une erreur sur cette image et aussi des erreurs sur les images P et B qui appartiennent au même GOP (Groupe of Pictures). Ci-dessous, nous présentons les principaux mécanismes qui ont été développés afin de faire face aux pertes des paquets et de minimiser leur impact sur la qualité de la vidéo.

- **Data Interleaving** [118]: L'objectif de cette technique est de minimiser l'effet des pertes consécutives sur la qualité de la vidéo décodée côté récepteur. Pour cela, une réorganisation des paquets vidéo est effectuée avant leur transmission, au niveau de l'émetteur, afin d'éloigner les paquets séquentiellement proches. En cas de perte séquentielle de plusieurs paquets durant la transmission, cette perte sera répartie sur l'ordre original et minimisera la dégradation de la qualité de la vidéo. Dans [119], l'IETF propose un modèle d'interleaving pour le transport des flux MPEG-4. L'avantage de cette technique est qu'elle ne consomme pas de débit supplémentaire. Par contre, elle introduit un délai dû à la réorganisation des paquets vidéo du côté émetteur et du côté récepteur.
- **ARQ (Automatic Repeat reQuest)**[120][121] : La retransmission des paquets vidéo perdus est considérée comme un mécanisme simple. L'utilisation d'ARQ suppose la présence d'une voie de retour entre l'émetteur et le récepteur. ARQ peut fonctionner suivant deux modes : (1) ARQ côté émetteur en utilisant les acquittements positifs (les paquets perdus sont détectés par l'émetteur) et (2) ARQ côté récepteur en utilisant les acquittements négatifs (les paquets perdus sont détectés par le récepteur). Cependant, la retransmission introduit une latence supplémentaire (temps de retransmission additionné au temps de réordonnement des paquets). De plus, la mise en place d'un mécanisme ARQ pour les transmissions multicast est complexe et sa mise à l'échelle est limitée à cause du nombre d'acquittements que peut recevoir le serveur.
- **FEC (Forward Error Correction)** : Le mécanisme FEC ajoute des paquets de redondance aux flux de paquets original, au niveau de l'émetteur, afin de permettre au récepteur de reconstruire les paquets perdus [122][123]. Les paquets redondants sont générés par des codes correcteurs traditionnels (XOR, Reed Solomon, etc.) en considérant un paquet comme étant un

symbole. Dans le RFC [124], l'IETF propose une extension de l'en-tête RTP assez flexible pour supporter l'implémentation de n'importe quel mécanisme FEC. Une voie de retour n'est pas nécessaire dans un mécanisme FEC, ce qui le rend idéal pour les transmissions multicast et les transmissions temps réel. L'interleaving FEC [125][126] est une variante de FEC qui ajoute de la redondance sur des paquets qui ont subi un interleaving. Ceci permet d'additionner les avantages de l'interleaving et de la FEC pour faire face à la perte de plusieurs paquets consécutifs. Le principal inconvénient de la FEC est la consommation additionnelle du débit due à l'ajout de données redondantes en plus du délai supplémentaire introduit par le regroupement des blocs FEC au niveau du récepteur.

- **UEP (Unequal Error Protection)** : L'UEP offre une protection différente des paquets vidéo. Il se base principalement sur le codage hiérarchique où la couche de base doit être mieux protégée contre les pertes, comparée aux couches d'amélioration. La protection des paquets peut être assurée en utilisant un mécanisme FEC avec des taux de redondance différents suivant l'importance des paquets [127][128][129] ou/et un mécanisme ARQ en utilisant une retransmission sélective [130][131].
- **Hybride ARQ/FEC** [132][133][134] : Cette solution représente une combinaison de la FEC au niveau bit et de l'ARQ au niveau paquet. Il existe deux schémas pour cette combinaison : le schéma de type I et le schéma de type II. Pour le type I, le code de détection et le code de correction sont présents dans chaque transmission et retransmission. Tandis que pour le type II, les paquets sont transmis uniquement avec le code de détection d'erreur. En cas d'erreur, la retransmission inclut uniquement le code correcteur d'erreurs qui est combiné au paquet erroné pour reconstruire le paquet original.
- **Errors concealment** [135]: Cette technique opère au niveau image en essayant d'estimer les informations, ou pixels, perdues à partir d'informations correctement reçues. Pour cela, elle se base sur la forte corrélation spatiale et temporelle qui existent dans les images vidéos et qui est exploitée par les algorithmes de codage. L'estimation d'informations perdues se base sur deux approches : l'interpolation spatiale et l'interpolation temporelle. L'interpolation spatiale suppose que l'image est naturellement lisse (smooth). Donc, les pixels perdus peuvent être reconstruits à partir de pixels adjacents. Par contre, l'interpolation temporelle suppose que l'image est lisse (smooth) durant le temps. Ainsi, un groupe de pixels perdus sur une image sera remplacé par le groupe de pixels se trouvant au même endroit dans l'image précédente. Une variante plus sophistiquée de cette approche utilise une compensation de mouvement pour estimer l'emplacement du groupe de pixels sur l'image au lieu de le dupliquer au même emplacement. Cette variante se base sur les informations du vecteur de mouvement utilisé pour le décodage. Cependant, la technique du « errors concealment » est limitée à la perte d'une partie de l'image et elle ne peut pas résister à la perte de plusieurs images consécutives.
- **Error resilient video coding** [136]: Ces mécanismes sont généralement intégrés aux algorithmes de codage vidéo afin d'offrir une certaine résistance aux erreurs. Le principal problème qui peut se produire durant le décodage est la perte de synchronisation du flux décodé. Les mécanismes qui permettent une resynchronisation du décodage sont donnés ci-dessous :

- Resync Markers : se sont des marqueurs de synchronisation insérés dans le flux afin de permettre la reprise du décodage en cas d'erreur.
- Reversible Variable Length Codes (RVLCs) : Ils permettent un décodage inversé du flux afin d'exploiter les informations se trouvant entre une erreur et un marqueur de synchronisation.
- Data Partitioning : Il permet de placer les informations importantes pour le décodage d'un flux juste après un marqueur de synchronisation. Ceci offre une meilleure protection pour ces informations.
- Application Level Framing (ALF) : Il définit la décomposition d'un flux en paquets afin d'être transmis. Une décomposition adéquate permet de réduire l'effet de perte de paquets. Cette procédure est effectuée au niveau RTP. L'IETF a défini des formats de paquets RTP pour plusieurs codecs : MPEG-1/MPEG-2[137], MPEG-4 [138], H.264 [139].

2.2.5.3.3 *La gestion de la latence et de la gigue*

Au niveau applicatif, il n'existe pas de mécanisme qui permet de diminuer la latence des transmissions. Cette diminution est prise en charge par des mécanismes de QoS au niveau réseau et au niveau MAC exposés respectivement dans les sections 2.2.2.3 et 2.2.3. Par contre, pour faire face à la gigue qui représente la variation des délais de réception des paquets, un buffer est utilisé au niveau du récepteur afin d'absorber cette variation. L'utilisation d'un buffer permet le fonctionnement des autres mécanismes de QoS, présentés précédemment, comme l'ARQ, la FEC, etc. Ainsi, la taille du buffer est un compromis entre la latence introduite par ce dernier et la fiabilité de la transmission. Plusieurs travaux se sont intéressés à l'estimation de la gigue et à l'élaboration de modèle mathématique pour une taille de buffer adaptative [140][141][142].

2.2.5.4 **Discussion**

Jusqu'à présent, les applications et les protocoles de niveau applicatif se contentaient d'exploiter les services qui sont disponibles au niveau des couches inférieures. Étant donné que les conditions de transmission sont dynamiques (débit, taux de perte, latence, gigue), les applications multimédia subissent les variations de ces conditions avec des réactions limitées à leur niveau en utilisant des mécanismes de QoS présentés précédemment. De nouvelles interactions entre les couches applicatives et les couches réseaux permettront à ces mécanismes d'être plus dynamiques et plus adaptatifs pour une configuration adéquate suivant l'état instantané du réseau. De plus, les applications multimédia doivent être informées des mécanismes QoS déployés dans le réseau pour assurer une certaine QoS aux paquets transportant des informations importantes, par exemple une correspondance entre la couche de base dans le codage hiérarchique et les classes de service prioritaires au niveau IP et au niveau MAC.

Ainsi, les futures applications multimédia ne doivent plus être passives et doivent interagir avec les couches inférieures, d'une part, pour les informer de leurs besoins et, d'autre part, pour s'adapter aux conditions de transmission. Cette adaptation doit être dynamique et transparente par rapport à l'utilisateur.

2.3 Les architectures *Cross-layer* pour les réseaux sans fil

2.3.1 Le concept du *Cross-layer*

La multiplication des mécanismes de QoS et des techniques d'adaptation sur les différentes couches du modèle TCP/IP a engendré plusieurs problématiques causées principalement par l'isolation des couches. Ces problématiques peuvent être regroupées en trois grandes classes listées ci-dessous :

- **La redondance** : Elle est causée par la duplication d'un mécanisme sur plusieurs couches. Par exemple, la retransmission dans la couche transport en utilisant TCP et la retransmission dans la couche MAC 802.11.
- **L'annulation** : Les avantages introduits par certains mécanismes de QoS sur les couches supérieures ne sont pas respectés par des couches inférieures ce qui provoque leur annulation. Par exemple, la priorité introduite au niveau IP grâce aux classes de service n'est pas forcément assurée au niveau de la couche MAC 802.11.
- **La contradiction** : Dans certains cas extrêmes, l'effet de deux mécanismes présents sur deux couches distinctes est contradictoire. Par exemple, l'utilisation d'UDP au niveau transport pour éviter les retransmissions et l'utilisation de la retransmission au niveau de la couche MAC 802.11.

Afin d'apporter une solution à toutes ces problématiques et d'optimiser les performances des systèmes communicants, nous avons assisté ces dernières années à l'émergence d'un nouveau concept sous l'appellation de *Cross-layer*. Ce dernier autorise la violation de la structure protocolaire en couches dans le but d'améliorer les performances de transmission dans les réseaux sans fil et d'assurer une meilleure QoS pour les services multimédia.

Comme tous les nouveaux concepts, il est très difficile de trouver ou de proposer une définition exacte pour le *Cross-layer*. Même au niveau de la terminologie, nous trouvons dans la littérature plusieurs variantes : la conception *Cross-layer*, l'adaptation *Cross-layer*, l'optimisation *Cross-layer*, le retour d'information *Cross-layer*. Dans [143][144], les auteurs ont proposé une définition générique qui peut englober toutes les techniques et tous les mécanismes *Cross-layer* qui existent actuellement. Ainsi, le design *Cross-layer* est défini comme suit : « la conception d'un protocole en violation avec l'architecture en couches de référence est une conception *Cross-layer* à l'égard de cette architecture ». Le terme violation englobe :

- La définition de nouvelles interfaces entre les couches.
- La redéfinition des limites des couches.
- La conception d'un protocole sur une couche en se basant sur la conception d'un autre protocole sur une autre couche.
- La configuration commune des paramètres à travers les couches.

Le concept *Cross-layer* permet la définition de protocoles ou de mécanismes qui ne respectent pas l'isolation des couches du modèle OSI (respectivement TCP/IP). Ainsi, il autorise la

communication entre deux, ou plusieurs, couches adjacentes, ou non adjacentes, dans le but d'améliorer les performances globales du système. Ceci peut être réalisé par la définition de nouvelles interfaces au niveau des couches qui permettent de récupérer leurs paramètres de performance mais aussi de configurer certains paramètres d'une manière dynamique. Ces paramètres peuvent être utilisés par les protocoles et/ou les mécanismes d'adaptation pour améliorer la performance globale de la communication en se basant sur des politiques d'adaptation.

2.3.2 La communication dans les architectures *Cross-layer*

Le principe de base du concept *Cross-layer* est de permettre l'échange d'informations entre les couches adjacentes et non adjacentes afin d'améliorer les performances de transmission. Cet échange d'informations peut être mis en œuvre suivant différents schémas. Parmi toutes les architectures *Cross-layer* proposées dans la littérature, deux modèles de communication peuvent être distingués [143][145]: La communication directe entre les couches et une base de données partagée entre les couches. Nous présentons ci-dessous une description de ces deux modèles représentés dans la Figure 2-9.

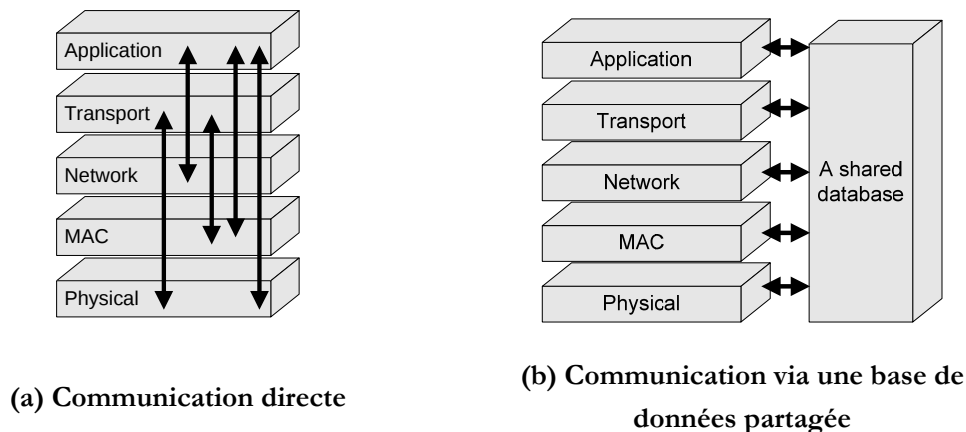


Figure 2-9 : Les modèles de communication *Cross-layer*

2.3.2.1 Communication directe entre les couches

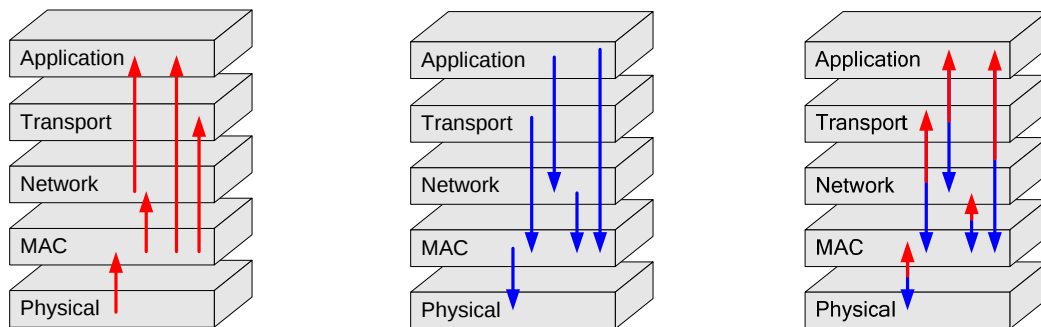
La communication directe entre les couches est le modèle le plus utilisé par les architectures *Cross-layer*. Il permet à une couche d'accéder directement aux paramètres et aux variables d'une autre couche sans passer par un intermédiaire. Cette communication peut être « in-band » en utilisant les en-têtes des protocoles qui sont déployés actuellement, par exemple, la couche IP accède aux champs de l'en-tête TCP « ECN » pour indiquer une congestion dans le réseau. Dans certains cas, des extensions d'en-tête sont nécessaires pour faire passer des informations supplémentaires. La communication directe peut être « out-of-band » en utilisant un nouveau protocole dédié. Le protocole de signalisation CLASS (Cross-Layer Signaling Shortcuts), présenté dans [146], est un parfait exemple de ce type de protocole. La communication « out-of-band » peut s'effectuer aussi en définissant de nouvelles interfaces (API : Application Programming Interface) au niveau d'une couche qui seront utilisées directement par d'autres couches pour récupérer et configurer des paramètres de fonctionnement. Un exemple de cette API est proposé dans [147].

2.3.2.2 Communication via une base de données partagée

Plusieurs architectures *Cross-layer* [148][149] proposent l'utilisation d'une base de données partagée afin de stocker et de récupérer des paramètres. Cette base est accessible par toutes les couches qui peuvent, ainsi, s'informer de l'état des autres couches ou récupérer des paramètres de configuration nécessaire à leur fonctionnement interne. La base de données est aussi accessible par un système d'optimisation responsable d'initialiser les paramètres avec les valeurs adéquates. La Figure 2-9 (b) illustre ce type de communication, la base de données est considérée comme une nouvelle couche en parallèle à toutes les autres. Cependant, pour mettre en œuvre cette base de données, il faut répondre à plusieurs contraintes conceptuelles, comme la localisation de la base de données (sur les nœuds communicants ou sur un nœud indépendant), le type de communication entre les couches et la base de données partagée, ainsi que le protocole de communication utilisé.

2.3.3 Les approches du *Cross-layer* dans les réseaux sans fil

Dans la littérature, plusieurs techniques *Cross-layer* ont été proposées pour améliorer les performances des transmissions sans fil. Au début, ces mécanismes étaient limités à l'interaction entre la couche physique et la couche liaison de données. De plus, les mécanismes proposés étaient indépendants et visaient l'amélioration d'une imperfection précise. Par la suite, nous avons assisté à l'apparition de plusieurs travaux proposant des interactions avec les couches supérieures et aussi à l'apparition d'architectures *Cross-layer* qui tentent de faire collaborer plusieurs couches, prenant en charge plusieurs paramètres, pour une optimisation globale.



(a) L'approche ascendante

(b) L'approche descendante

(c) L'approche mixte

Figure 2-10 : Les approches du *Cross-layer*

Afin de simplifier la présentation de ces travaux, nous les avons classés en trois grandes approches identifiées dans [143][144] et représentées dans la Figure 2-10. Nous présentons, ci-dessous, une description de chacun de ces approches.

- **L'approche ascendante (Bottom-up):** Les couches supérieures optimisent leurs mécanismes en fonction des paramètres (conditions) des couches inférieures.
- **L'approche descendante (Top-down):** Les couches supérieures décident des paramètres de configuration des couches inférieures. Ou bien, les couches inférieures considèrent certaines spécificités du niveau applicatif pour exécuter leurs traitements.

- **L'approche mixte (Integrated):** Cette approche exploite les deux approches précédentes dans une même architecture afin de trouver la meilleure configuration inter-couches pour un fonctionnement optimal du système.

2.3.3.1 Les approches ascendantes (Bottom-up)

Dans [85], les auteurs présentent les avancées majeures et les futures orientations de recherche dans le domaine des réseaux sans fil. La conception *Cross-layer* est présentée comme l'un des défis majeurs qui doit être relevé. L'article soulève le problème de contrôle de congestion du protocole TCP sur les réseaux sans fil (exposé dans la section 2.2.4.4) et propose l'utilisation du mécanisme de notification explicite de congestion (ECN : explicit Congestion Notification) afin de différencier les pertes causées par la congestion et les pertes causées par des interférences. Le mécanisme ECN propose d'utiliser le bit ECN présent dans l'en-tête du protocole TCP. Ce bit est initialisé à '0' par l'émetteur et peut être fixé à '1' par un routeur s'il est congestionné. Quand le paquet arrive au récepteur, ce dernier informe l'émetteur sur l'état du bit ECN. Ainsi, l'émetteur pourra faire la distinction entre un état de congestion et un état d'interférence. Cependant, les auteurs n'ont pas proposé un comportement à suivre par le protocole TCP dans le cas de pertes causées par des interférences.

Par la suite, l'article décrit comment utiliser l'état d'un canal dans un simple algorithme afin d'améliorer le débit du réseau. Pour cela, les auteurs supposent un canal de transmission avec deux états (ON, OFF). Le canal de transmission est partagé par trois stations et l'état du canal est différent pour chaque station. Avec un algorithme d'ordonnancement standard, le canal est partagé équitablement entre les utilisateurs, ce qui signifie que chaque station a $1/3$ du time-slot. Étant donné que les paquets ne peuvent pas être transmis quand l'état du canal est OFF, chaque station peut transmettre durant $1/6$ du time-slot. Cependant, si l'algorithme d'ordonnancement connaît l'état du canal pour chaque station et si l'envoi de paquets est interrompu uniquement quand tous les états sont OFF (ce qui se produit le $1/8$ du time-slot), le canal peut envoyer durant $7/8 = 1 - 1/8$ du time-slot. Ainsi, chaque station peut transmettre des paquets durant $7/24$ du time-slot qui représente le double du temps fourni avec l'algorithme standard.

Dans la même optique, les auteurs dans [150] proposent un algorithme similaire pour l'allocation de ressources dans les réseaux 3G exploitant un accès multiple à répartition par code (CDMA : Code Division Multiple Access). En effet, dans les réseaux CDMA, la variation du canal de transmission pour les stations mobiles (MS : Mobile Station) n'est pas identique. Cette diversité multi-utilisateurs est exploitée pour fournir des services en continu uniquement aux MSs dont la qualité du canal de transmission est élevée. Le MS, dont l'état instantané du canal est faible, reporte ses transmissions jusqu'à ce que son canal change d'état afin de ne pas pénaliser les autres MS.

Cependant, ce type d'algorithme montre des limites pour la transmission de flux temps réel puisque les paquets en attente d'amélioration du canal sont supprimés au bout d'un certain temps. Ceci montre l'importance d'introduire des contraintes temporelles dans ce type d'algorithme pour éviter que les MSs ne soient bloqués durant une longue période.

Un algorithme d'ordonnancement se basant sur l'état du canal (CSI : Channel State Information) a également été proposé pour les réseaux satellitaires dans [151]. L'algorithme est

implémenté au niveau liaison de données est exploité l'état du canal satellitaire pour décider de l'envoi d'un paquet. Le canal est aussi modélisé par deux états : un état bon et un état mauvais.

Les auteurs dans [152] explorent une architecture *Cross-layer* pour la transmission des flux vidéo sur des réseaux sans fil. L'architecture *Cross-layer* proposée maintient la structure en couche et identifie les principaux paramètres qui peuvent être échangés entre ces couches. Ainsi, une technique d'adaptation est proposée au niveau liaison de données qui détermine la taille optimale d'un paquet en fonction de la modulation et du codage qui sont à leur tour adaptés en fonction du SINR (Signal-to-Interference-plus-Noise Ratio). Le SINR donne une indication sur la qualité du canal de transmission. Dans le cas où le SINR est élevé, des modulations et des codages complexes sont utilisés et la taille des paquets est augmentée en conséquence. Quand le SINR décroît, des modulations et des codages plus simples sont utilisés afin d'augmenter la résistance du signal aux interférences et en conséquence la taille des paquets est diminuée. Les auteurs proposent un modèle analytique pour calculer la taille optimale du paquet qui maximise le débit en fonction de la taille de l'en-tête, le nombre de bits codés par symbole (voir section 2.2.1) et la probabilité d'erreur d'un symbole.

Dans [153], les auteurs proposent un système pour la transmission de la vidéo sur les réseaux sans fil. Le système se base sur deux mécanismes principaux : l'UEP basé sur la FEC et l'ARQ prioritaire. Ces deux mécanismes sont appliqués au niveau applicatif mais sur des paquets de taille égale aux paquets transmis sur l'interface du réseau sans fil, appelés RLP (Radio Link Protocol). Donc, un paquet applicatif est décomposé en plusieurs paquets RLP sur lesquels sont appliqués les deux mécanismes UEP et ARQ. L'UEP se charge de la protection inéquitable des paquets d'un flux vidéo suivant leur importance. En effet, le flux vidéo est organisé en quatre classes suivant l'importance des données pour le décodage de la vidéo. Ensuite, des paquets de redondance sont ajoutés à chaque classe, côté émetteur, avec des pourcentages différents afin de mieux protéger les classes importantes. La redondance est générée en utilisant l'algorithme RS (Reed-Solomon). Du côté récepteur, si des pertes de paquets sont détectées, le mécanisme UEP tente de décoder les paquets grâce à la redondance. Dans le cas où l'UEP échoue, l'ARQ prioritaire est exécuté pour décider de la retransmission ou non du paquet perdu suivant son importance. L'ARQ prioritaire décide de l'utilité de la retransmission en faisant des estimations sur le buffer de lecture et le RTT afin d'éviter la retransmission de paquets qui ne pourront pas être exploités à leur réception.

La valeur ajoutée de ce procédé est l'application de mécanismes de correction d'erreurs et de retransmission au niveau applicatif sur des paquets RLP pour éviter la perte d'un paquet applicatif à cause de la perte d'un paquet RLP. De plus, la retransmission se limite aux paquets RLP perdus afin de minimiser le débit de la retransmission. Ceci sous entend que la couche applicative doit prendre en considération la taille des paquets RLP utilisée au niveau de l'accès réseau en se basant sur des techniques *Cross-layer*.

D'autre part, plusieurs travaux se sont intéressés à la consommation d'énergie des terminaux mobiles utilisant des transmissions sans fil. Des techniques *Cross-layer* ont été proposées afin de trouver un compromis entre les performances des transmissions et leur consommation d'énergie. Dans [154][155][156], la gestion de l'énergie *Cross-layer* d'une manière ascendante est présentée. Elle consiste au développement de nouvelles adaptations qui considèrent le niveau de l'énergie disponible afin d'offrir une plus grande autonomie au terminal tout en assurant une certaine QoS

pour les communications. Les techniques proposées se basent principalement sur des algorithmes d'ordonnement de paquets et des techniques d'adaptation de la modulation du signal au niveau physique dans les réseaux sans fil. Ces techniques ont pour but de minimiser deux principaux paramètres : le délai de transmissions des paquets et l'énergie consommée durant leur transmission.

2.3.3.2 Les approches descendantes (Top-down)

L'ordonnement optimal pour minimiser la congestion et la distorsion (CoDiO : Congestion-Distortion Optimized) est l'une des principales techniques dans les approches descendantes étudiées dans plusieurs travaux [152][157][158]. La distorsion correspond à la différence de la qualité entre la vidéo encodée, côté émetteur, et la vidéo décodée, côté récepteur. Cette distorsion est causée par les pertes de paquets vidéo durant la transmission. Le CoDiO vise à minimiser cette distorsion suivant les contraintes du débit disponible dans le réseau. Il se base sur des algorithmes d'ordonnement qui décident des paquets vidéo à envoyer suivant leur importance et suivant le débit disponible. Ceci réduit la congestion et le délai de transmission dans le réseau.

Dans [150], les auteurs proposent une optimisation des performances du protocole TCP sur les réseaux sans fil 3G exploitant un accès multiple à répartition par code (CDMA : Code Division Multiple Access). Cette optimisation propose une variation de plusieurs paramètres au niveau de l'accès réseau afin de les faire correspondre au débit calculé au niveau TCP.

Le mapping entre les couches supérieures et les couches inférieures a été abordé dans [159]. L'article propose une architecture *Cross-layer* pour la transmission des flux H.264 sur des réseaux sans fil 802.11e. L'architecture se base principalement sur deux techniques : un partitionnement judicieux des données au niveau applicatif et un mapping efficace au niveau des classes de service de la couche liaison de données 802.11e. En effet, la norme H.264 introduit plusieurs techniques de résistance aux erreurs pour le codage des flux vidéo. Le partitionnement de données représente l'une de ces techniques qui permet d'organiser les données vidéo compressées en unités de différentes importances. Ainsi, trois partitions sont générées : A, B et C, où A est la partition la plus importante et C la moins importante. Pour être utile, les partitions de faible importance nécessitent la présence des partitions importantes. Ces différentes partitions sont encapsulées dans des unités NALUs (Network Abstraction Layer Units) qui peuvent être considérées comme des paquets puisque chaque NALU possède un en-tête et une charge utile. L'en-tête contient le champ NRI (Nal_Ref_Idc) de deux bits qui indique la priorité de la charge utile du NALU. Pour permettre un service différencié des NALUs suivant leur importance et leur priorité, l'architecture *Cross-layer* se base sur la nouvelle norme 802.11e et principalement les ACs définies par l'EDCA (voir la section 2.2.2.4.1). Ainsi, un algorithme de marquage est utilisé au niveau MAC afin de mapper les différents NALUs sur les ACs disponibles en se basant sur le champ NRI. Ce type d'architecture permet la continuité de la QoS à travers les couches réseaux en offrant le meilleur service aux données importantes afin de maximiser la qualité de la vidéo perçue côté récepteur.

Dans [160], les auteurs soulèvent le problème de la retransmission (ARQ) qui peut être présent au niveau liaison de données pour les retransmissions point-à-point et au niveau transport pour des retransmissions de bout-en-bout. En effet, la présence de ce mécanisme au niveau liaison de données dégrade sérieusement les performances du protocole TCP [161] et introduit un délai de

transmission désagréable pour les flux temps réel. Pour faire face à cette problématique, les auteurs proposent un système ARQ adaptatif qui permet de fournir un ARQ personnalisé en fonction des besoins des applications. Pour ce faire, les besoins en QoS des applications sont signalés à la couche liaison de données qui se charge de configurer les paramètres de l'ARQ en conséquence (le nombre de retransmission, le temps d'attente avant une retransmission, etc.).

D'autres techniques *Cross-layer* [154][162][163][164] utilisant le modèle descendant ont été proposées pour la préservation d'énergie. Dans ces techniques, les applications contrôlent les interfaces de communication et, par la même, l'énergie consommée par ces interfaces. Les applications se basent généralement sur des politiques d'adaptation pour décider de basculer vers des modes de consommation prédéfinies dans les couches inférieures [3]. Par exemple en mode minimale, en mode de veille, etc. Dans ce contexte, une politique peut décider d'éteindre l'interface réseau s'il n'y a pas de communication au bout d'un certain temps. Cependant, pour les applications multimédia temps réel, ces techniques sont inefficaces à cause de la continuité de la transmission exigée par ce type d'application [165].

2.3.3.3 Les approches mixtes (Integrated)

Les approches mixtes exploitent à la fois les interactions ascendantes et descendantes. Dans [144], les auteurs présentent une architecture *Cross-layer* pour analyser, sélectionner et adapter les différentes stratégies présentes sur les couches du modèle OSI. Ceci a pour but d'augmenter la qualité des flux multimédia, de préserver la consommation d'énergie des terminaux et d'optimiser l'utilisation spectrale des canaux de transmission. L'architecture proposée ne nécessite pas la redéfinition des protocoles existant mais plutôt une optimisation conjointe de leur fonctionnement à travers toutes les couches. La problématique du *Cross-layer* a été formalisée comme étant un problème d'optimisation avec l'objectif de sélectionner une stratégie commune à travers toutes les couches OSI. Dans cette modélisation, uniquement les couches applications (APP), liaison de données (MAC) et physique (PHY) sont considérées. Chaque couche, PHY/MAC/APP possède respectivement un nombre défini N_p , N_m , N_a de stratégies d'adaptation et de protection contre les erreurs. Pour la couche PHY, les stratégies PHY_i avec i appartenant à $[1, N_p]$ peuvent représenter les différentes combinaisons de modulation et de codage qui sont disponibles. Les stratégies MAC_i avec i appartenant à $[1, N_m]$, correspondent à des mécanismes ARQ, FEC, ordonnancement de trames, contrôle d'admission. Les stratégies APP_i avec i appartenant à $[1, N_a]$, correspondent à des mécanismes d'adaptation du codage vidéo, décomposition en paquets des images vidéo, lissage du trafic, classement du trafic suivant des priorités, ordonnancement de paquets et aussi les mécanismes FEC et ARQ appliqués au niveau paquets. La stratégie *Cross-layer* commune est définie par S :

$$S = \{ PHY_1, \dots, PHY_{N_p}, MAC_1, \dots, MAC_{N_m}, APP_1, \dots, APP_{N_a} \}$$

Suivant l'ensemble S , il existe $N = N_p \times N_m \times N_a$ stratégies possibles. L'optimisation *Cross-layer* a pour but de sélectionner la meilleure stratégie qui offre la meilleure qualité de service pour les flux multimédia sous les contraintes des transmissions sans fil ainsi que les contraintes du système. La sélection de la meilleure stratégie *Cross-layer* est très complexe. Ceci est dû à de nombreuses contraintes énumérées ci-dessous :

- Sélectionner la meilleure solution d'une manière analytique est très complexe, voir impossible, à cause du comportement non déterministe des contraintes et de la dépendance de certaines stratégies.
- Chaque stratégie a pour but d'optimiser chaque couche indépendamment des autres. De plus, les couches manipulent différentes unités de données et possèdent leurs propres paramètres de configuration et métriques de performances.
- La sélection de la stratégie doit être dynamique et continue puisque les réseaux sans fil sont caractérisés par leur caractère variable durant le temps.
- Au niveau de chaque couche, les stratégies doivent être groupées, ordonnées et initialisées avant le début de la sélection.
- Les contraintes techniques doivent être considérées comme le changement dynamique de certains paramètres.

L'article identifie, par la suite, les différentes interactions possibles entre les trois couches PHY, MAC et APP pour :

- La sélection d'une modulation optimale au niveau PHY qui offre une qualité de service élevée.
- Une consommation d'énergie optimale.

Dans [148], une nouvelle architecture *Cross-layer*, appelée CrossTalk, est proposée principalement pour les réseaux ad hoc. CrossTalk a pour but d'atteindre des objectifs globaux avec des comportements locaux. Dans cet article, le *Cross-layering* est considéré comme une amélioration de l'approche en couche en permettant un partage d'information entre les couches. L'objectif ultime est de préserver la structure en couches tout en permettant des améliorations de performance et de nouvelles adaptations. Le nouveau concept proposé par CrossTalk fournit une vue globale de l'état du réseau en utilisant plusieurs métriques présentes sur différentes couches. Cette vue globale permettra à chaque nœud réseau de comparer son état local avec l'état global du réseau afin d'appliquer les adaptations nécessaires. En effet, les architectures présentées dans la littérature décidaient des adaptations en fonction de l'état local uniquement. Le CrossTalk propose des adaptations locales en fonction des connaissances globales. L'architecture CrossTalk est constituée de deux entités de gestion de données. La première entité est responsable de l'organisation des informations locales. Ces informations peuvent être fournies par des couches protocolaires ou des composants systèmes. Ils informent sur l'état des batteries, le nombre de voisins, la puissance du signal, l'état du canal de transmission (SNR) et la localisation du nœud réseau. Chaque protocole local peut accéder et utiliser ces informations pour optimiser son fonctionnement. L'ensemble de toutes ces informations locales forme la vue locale. La deuxième entité se charge d'établir la vue globale des informations collectées par la vue locale. Cette vue globale est construite au niveau de chaque nœud par une procédure de dissémination d'informations locales en utilisant des paquets de contrôle ou les en-têtes des paquets qui transitent entre les nœuds afin de minimiser l'*overhead*.

Pour la construction de la vue globale, des algorithmes d'agrégation ont été utilisés, par exemple, par des moyennes pondérées dans le temps ou dans l'espace. La pondération temporelle donne plus

de poids aux informations récentes et la pondération spatiale donne plus de poids aux informations provenant de nœuds voisins immédiats. Les auteurs ne se sont pas intéressés aux types de données collectées dans la vue locale et globale. Cependant, ils ont fourni un exemple d'application de leur architecture pour le protocole de routage ad hoc AODV en utilisant la charge locale afin de réaliser un équilibrage de charge globale.

Dans [166], les auteurs proposent un nouveau mécanisme de protection *Cross-layer* qui fournit une QoS adaptative en exploitant conjointement le codage en couches des vidéos, les files d'attente prioritaires au niveau de la couche réseau et l'adaptation de la retransmission au niveau liaison de données des réseaux sans fil. Le mécanisme *Cross-layer* proposé a pour but de trouver un compromis entre le nombre de retransmission (ARQ) utilisé au niveau liaison de données et la taille des files d'attente au niveau de la couche réseau. En effet, pour offrir une protection optimale contre les erreurs de transmission, le nombre de retransmissions doit être le plus grand possible. Cependant, l'augmentation du nombre de retransmissions augmente le délai de transmission d'une trame, ce qui engendre automatiquement une augmentation de la taille de la file d'attente de la couche réseau. Cette augmentation peut engendrer une surcharge de la file d'attente et une suppression de paquets. Ainsi, pour un débit de flux et un état de canal donnés, le nombre de retransmissions approprié peut être déterminé afin d'offrir une protection optimale du flux et une préservation de sa QoS. Afin d'offrir une protection inéquitable (UEP) mieux adaptée au codage en couche des vidéos, le mécanisme *Cross-layer* propose d'utiliser des files d'attente prioritaires afin de mapper chaque couche sur une file au niveau réseau. Au niveau de la couche liaison de données, un algorithme adaptatif temps réel du nombre de retransmissions est utilisé pour chaque file d'attente afin de protéger chaque file suivant sa taille et l'état du canal de transmission.

Dans [149][167][168], une stratégie d'optimisation *Cross-layer* est proposée. Cette stratégie permet d'optimiser conjointement le fonctionnement des couches application, liaison de données et physique. L'optimisation *Cross-layer* dans cette nouvelle stratégie est pilotée par la couche applicative puisque l'objectif principal est de maximiser la satisfaction de l'utilisateur en relation directe avec l'application. L'architecture *Cross-layer* (CLA : Cross Layer Architecture) est définie d'une manière générique. Elle comprend N couches et un module d'optimisation *Cross-layer* (CLO : Cross Layer Optimizer). Le CLO optimise conjointement plusieurs couches en élaborant des prédictions sur leurs états et en sélectionnant les paramètres de configuration optimaux. Les étapes d'optimisation *Cross-layer* sont détaillées ci-dessous :

- 1- La couche d'abstraction calcule une abstraction des paramètres spécifiques à chaque couche. Ceci permet de réduire le nombre de paramètres utilisés par le CLO.
- 2- L'optimisation trouve les valeurs des paramètres de configuration des couches suivant une certaine fonction d'optimisation. Cette dernière maximise la qualité de la vidéo décodée au niveau récepteur.
- 3- La couche de reconfiguration redistribue les valeurs de configuration optimales à travers les couches de transmission. Il est responsable de traduire les paramètres abstraits de la couche d'abstraction, en paramètres spécifiques à chaque couche.

Ces étapes sont répétées périodiquement suivant les exigences de l'application et la variation du canal média utilisé. Le nombre de paramètres gérés par l'architecture *Cross-layer* étant important, les auteurs ont proposé une classification de ces paramètres que nous avons détaillés ci-dessous :

- Les paramètres directement configurables DT (Directly Tunable) représentent les paramètres qui peuvent être reconfigurés par le CLO. Par exemple, la réservation des time-slots dans le TDMA.
- Les paramètres indirectement configurables IT (Indirectly Tunable) représentent les paramètres de monitoring qui peuvent varier en changeant un paramètre directement configurable. Par exemple, BER (Bit Error Rate) qui dépend de la modulation et du codage utilisés.
- Les paramètres de description D (Descriptive). Ces paramètres statiques peuvent être lus par le CLO mais ne peuvent pas être reconfigurés. Par exemple, le nombre d'images par seconde, la taille d'une image vidéo.
- Les paramètres abstraits A (Abstracted) représentent les abstractions des paramètres DT, IT et D qui sont utilisés dans le CLO. Par exemple, la probabilité de transmission, modèle de perte de paquets à deux états.

Par la suite, les auteurs détaillent les éléments de cette architecture pour une optimisation *Cross-layer* pour le streaming vidéo sans fil. Dans cette optimisation, les couches application, liaison de données et physique sont optimisées conjointement pour une allocation de ressources optimale entre plusieurs utilisateurs. Enfin, l'*overhead* introduit par le *Cross-layer* en termes de puissance de calcul et de communication est soulevé afin de trouver le meilleur compromis entre optimisation et *overhead*.

Dans [147][169][170], les auteurs présentent l'architecture ECLAIR qui fournit un cadre générique pour la conception et l'implémentation d'adaptations *Cross-layer* pour les stations mobiles. ECLAIR exploite le fait que le comportement d'un protocole est déterminé par les valeurs de ces paramètres de configuration. Pour cela, l'architecture ECLAIR est décomposée en deux sous-systèmes :

- 1- Les couches de configuration (TLs : Tuning Layers) : Le but du TL est de fournir une interface aux structures de données qui déterminent le comportement d'un protocole. Les TLs ont été regroupées suivant leurs fonctions. Par exemple, TCPTL pour le protocole TCP et UDPTL pour le protocole UDP. Le TL a le droit de lire et de réinitialiser les paramètres des structures de données d'un protocole. Pour cela, le TL est décomposé en deux sous-couches :
 - a. Sous-couche de configuration générique : cette sous-couche détermine les interfaces du TL d'une manière générique et indépendante des implémentations.
 - b. Sous-couche d'accès dépendante de l'implémentation : cette sous-couche fournit une implémentation dépendante du système des interfaces TL définies dans la sous-couche de configuration générique.
- 2- Le sous-système d'optimisation OSS (Optimizing SubSystem) : Il regroupe plusieurs modules d'optimisation de protocoles PO (Protocol Optimizers). Le PO implémente un algorithme d'optimisation *Cross-layer* particulier qui décide des actions nécessaires à

entreprendre suivant l'état du protocole et les événements qui peuvent se produire sur différentes couches. Pour une action d'optimisation, le PO invoque une fonction de l'interface TL. Le PO s'enregistre pour un événement au niveau du TL en utilisant l'API « register ». Le TL notifie le PO quand un événement survient. Le PO peut aussi utiliser l'API du TL pour demander l'état des protocoles qui peuvent être modifiés.

Les auteurs détaillent leur architecture pour un contrôle par l'utilisateur du débit applicatif. Les applications visées exploitent le protocole TCP et l'optimisation se base principalement sur le contrôle de la fenêtre de réception du protocole TCP.

2.3.4 Discussion

À partir des travaux présentés ci-dessus, nous pouvons constater que le concept du *Cross-layer* est considéré comme une amélioration de la structure en couche qui existe actuellement. Cette amélioration n'est pas gratuite, elle possède un coût en terme de complexité de conception. En effet, les nouveaux protocoles et les nouvelles applications qui exploiteront ce nouveau paradigme seront plus difficiles à concevoir puisqu'ils doivent prendre en considération plusieurs paramètres présents sur différentes couches. Cette complexité est accentuée lorsque les mécanismes *Cross-layer* sont conçus d'une manière distribuée sans aucun contrôle centralisé. Parmi les approches que nous avons décrites, ci-dessus, l'approche mixte est plus attractive puisqu'elle englobe, en quelque sorte, les deux autres approches (ascendantes et descendantes) dans une seule architecture. Les travaux qui existent dans l'approche mixte proposent des architectures diverses pour améliorer conjointement les performances de toutes les couches. Cependant, ces travaux n'apportent pas de réponse à toutes les problématiques liées à la mise en œuvre d'une architecture *Cross-layer* fonctionnelle, par exemple, sur quelle couche se trouve le module d'optimisation ? Comment ce module communique avec les couches pour être le plus réactif possible ? Et quelle est l'apport de ces architectures dans l'amélioration de la QoS pour les services multimédia ? De plus, la majorité des travaux existants se basent sur des simulations pour démontrer l'efficacité des adaptations proposées.

L'architecture *Cross-layer* proposée dans cette thèse tente d'apporter des réponses à ces interrogations en se basant sur une approche *Cross-layer* mixte accompagnée d'un coordinateur central, qui concentre l'état du système et le contrôle des adaptations exécutées en un seul module. L'objectif premier du coordinateur central est de maintenir un fonctionnement cohérent du système afin d'améliorer la QoS des applications multimédia.

2.3.5 Les projets Européens traitant la problématique *Cross-layer*

Plusieurs projets financés par la commission Européenne ont traité, ou traitent, la problématique du *Cross-layer* en étudiant ce nouveau concept et en proposant de nouvelles interactions *Cross-layer* afin d'améliorer les performances des transmissions. Nous présentons, ci-dessous, une brève description de quelques projets.

2.3.5.1 Le projet 4MORE

Ce projet [171] s'inscrit dans le développement de systèmes 4G, la nouvelle génération des communications mobiles. La vision européenne de la 4G est un nouveau système basé sur IP et offrant tous les services, à n'importe quel moment en utilisant n'importe quel terminal. Pour cela, la 4G doit offrir des débits variant entre 2 et 100 Mbps couvrant plusieurs environnements mobiles (véhicules, piétons) et fixes (intérieures et extérieures) dans une plage de fréquences de 50 – 100 MHz. La technologie de transmission utilisée pour atteindre cet objectif est le MC-CDMA (Multi-Carrier-CDMA) qui est déjà adopté au Japon. L'objectif du projet 4MORE est d'étudier, développer, intégrer et valider un système sur puce, rentable et de faible consommation, pour les terminaux mobiles utilisant un système de transmission multi-antennes basé sur MC-CDMA. Ce système doit intégrer de nouvelles interactions entre la couche physique (MC-CDMA) et les algorithmes de la couche MAC.

2.3.5.2 Le projet PHOENIX

L'objectif de ce projet [172] est le développement d'un lien de communication de bout-en-bout basé sur un réseau sans fil optimisé. Pour cela, le projet propose une communication entre le monde applicatif (le codage source, le cryptage, etc.) et le monde des transmissions (le codage canal, la modulation) à travers le monde réseau (le protocole IPv6). Pour atteindre cet objectif, le projet se focalise sur trois axes principaux. Le premier axe développe un schéma innovant pour une optimisation conjointe du codage source et canal (JSCC : Joint Source Channel Coding). Ce premier axe inclut la définition de nouveaux codages et l'adaptation des codages qui existent actuellement. Le deuxième axe définit de nouvelles stratégies d'adaptation efficaces qui prennent en considération des paramètres réels (présence de cryptage, le nombre de saut sans fil, etc.). Enfin, le troisième axe définit une nouvelle architecture réseau pour les futurs systèmes de communication sans fil. Cette nouvelle architecture se base sur les optimisations étudiées dans les autres axes.

2.3.5.3 Le projet ENTHRONE II

Le projet ENTHRONE II [173] propose une solution globale pour gérer la chaîne de distribution des services audio/vidéo. Cette solution englobe la protection du contenu, la distribution à travers les NGNs et la réception au niveau de l'utilisateur final. Le but de ce projet n'est pas d'unifier ou d'imposer une stratégie mais plutôt d'harmoniser les différentes fonctionnalités afin de supporter une QoS de bout-en-bout à travers des réseaux hétérogènes distribuant des services IP multimédia pour divers utilisateurs. Pour atteindre cet objectif, le projet se base une architecture de gestion ouverte et décentralisée pour la distribution de bout-en-bout. Le modèle MPEG-21 est utilisé comme un support commun pour l'implémentation et la gestion des fonctionnalités de l'architecture. L'adaptation MPEG-21 *Cross-layer*, qui a pour objectif d'adapter les contenus multimédia pour maintenir la QoS, est considérée comme un élément principal de

l'architecture ENTHRONE II. Cette adaptation prend en considération plusieurs paramètres allant de la couche application (qualité objective du flux multimédia) jusqu'à la couche physique (débit physique dans les réseaux sans fil, niveau du signal) en passant par la couche réseau (latence, taux de perte).

2.4 Conclusion

Dans ce chapitre, nous avons présenté un état de l'art des différents mécanismes et techniques de QoS qui existe actuellement au niveau des différentes couches du modèle TCP/IP. En effet, l'introduction des réseaux sans fil et la multiplication des flux multimédia ont conduit à la définition de plusieurs mécanismes au niveau des différentes couches afin d'améliorer les performances de transmission. Nous avons discuté, au niveau de chaque couche, les inconvénients que pouvaient introduire certains mécanismes en considérant les performances globales du système de communication.

Ceci nous a amené à présenter le nouveau concept *Cross-layer* qui permet de faire face à ces inconvénients en autorisant un échange d'information entre les couches. Ce nouveau paradigme suscite un grand intérêt pour améliorer les performances des réseaux sans fil dont l'état du canal varie considérablement, comparé à un canal filaire. Le partage de l'état du canal avec les couches supérieures permettra à ces dernières de répondre efficacement à ses changements. Le *Cross-layer* permet aussi une collaboration entre les différents mécanismes de QoS qui existent sur les différentes couches TCP/IP afin d'assurer une continuité verticale de la QoS pour les flux multimédia.

Dans le chapitre suivant, nous allons voir comment ce nouveau concept peut être exploité dans une architecture réelle pour garantir la QoS des services IP multimédia transmis sur des réseaux sans fil 802.11.

Chapitre 3

Proposition d'une architecture *Cross-layer* pour les services IPTV sans fil

3.1 Introduction

Actuellement, la technologie IP est considérée comme la technologie de prédilection pour réaliser la convergence des réseaux et des services. L'interconnexion des réseaux traditionnels aux réseaux IP, filaire et sans fil, déployés actuellement représente un premier pas vers cette convergence en attendant la mise en œuvre et le déploiement des NGNs (Next Generation Networks). Cette interconnexion permet d'étendre l'accessibilité des services traditionnels sur des réseaux IP hétérogènes tout en permettant une certaine interopérabilité, principalement avec les réseaux téléphoniques. Cependant, plusieurs défis doivent être relevés pour permettre une interconnexion opérationnelle garantissant la QoS des services traditionnels et préservant le fonctionnement des réseaux IP.

Dans cette optique, nous proposons dans ce chapitre une architecture qui permet d'interconnecter les réseaux de diffusion DVB-T et les réseaux de données IP/802.11. L'objectif de cette architecture est d'étendre l'accessibilité des flux TV en offrant un service IPTV sans fil à des utilisateurs mobiles et hétérogènes connectés via un réseau 802.11. L'architecture que nous proposons couvre toute la chaîne de distribution des flux TV sur le réseau IP :

- La réception des flux TV du réseau de diffusion DVB-T et la transformation des services DVB-T en services IPTV au niveau transport et au niveau signalisation.
- La transmission des services IPTV dans le cœur du réseau IP en assurant la scalabilité.
- L'adaptation des flux IPTV basée sur la norme MPEG-21 pour permettre un accès universel au niveau des réseaux d'accès 802.11.
- Le maintien de la QoS des flux IPTV au niveau des réseaux d'accès 802.11 en exploitant le concept du *Cross-layer* détaillé dans le chapitre précédent.

Ainsi, nous commençons ce chapitre par décrire le contexte général dans lequel s'inscrit notre contribution. Dans ce contexte, nous présentons les deux concepts émergents, la convergence des réseaux et l'accès universel aux objets multimédia, en détaillant la vision de chaque concept, ainsi que les technologies qui en découlent. Dans le reste du chapitre, nous présentons notre contribution en décrivant l'architecture proposée pour interconnecter les réseaux DVB-T et les réseaux IP/802.11. Cette description est organisée en trois parties : présentation des réseaux DVB-T, l'interconnexion des réseaux DVB-T et le cœur du réseau IP et l'interconnexion du cœur du réseau et les réseaux d'accès 802.11. Nous présentons le principe de fonctionnement de chaque partie, ainsi que les entités clés qui permettent de réaliser les interconnexions.

3.2 Contexte général

3.2.1 La convergence des réseaux de diffusion, de télécommunications et de données

Jusqu'à la fin des années 90, le monde des réseaux était constitué de trois sous-réseaux distincts et indépendants, à savoir les réseaux de diffusion (Broadcast network), les réseaux de télécommunications (telecommunication network) et les réseaux de données (Internet network). Chaque réseau possède sa propre infrastructure et l'interconnexion entre ces réseaux était très limitée ou inexistante. Ainsi, le réseau de télécommunications est caractérisé par son interactivité et sa fiabilité en se basant sur un plan de contrôle rigide. Tandis que, le réseau de diffusion se caractérise par sa communication unidirectionnelle qui permet sa scalabilité en supportant un grand nombre d'utilisateurs. Enfin, le réseau de données, basé sur le protocole IP, représente un intermédiaire entre la fiabilité du réseau de télécommunications et la mise à l'échelle du réseau de diffusion tout en permettant l'interactivité.

Durant ces dernières années, Les trois réseaux ont connu des évolutions conséquentes. Le réseau de télécommunications est devenu numérique, sans fil et mobile grâce à l'émergence des réseaux mobiles de 2^{ème} et 3^{ème} générations qui ont augmenté le débit de transmission, et par la même, le nombre de services. Le même constat peut être fait pour le réseau de diffusion qui migre actuellement vers la télévision et la radio numérique grâce aux nouvelles normes de diffusion comme DVB (Digital Video Broadcast), ATSC (Advanced Television Systems Committee) pour la TV et DAB (Digital Audio Broadcasting), DRM (Digital Radio Mondiale) pour la radio. Ces nouvelles normes introduisent la mobilité et permettent la transmission d'autres services unidirectionnels. Concernant le réseau IP, l'augmentation du débit dans le cœur du réseau et dans les réseaux d'accès a permis le déploiement de nouveaux services et principalement des services multimédia intégrant la transmission de flux audio et vidéo (partage de fichiers, vidéoconférence, streaming, jeux interactifs, etc.). De plus, ces services sont devenus mobiles et fiables grâce, d'une part, à l'apparition de nouvelles technologies d'accès sans fil comme le IEEE 802.11 WLAN et le 802.16 pour le WMAN, et d'autre part, le développement de nouveaux protocoles qui permettent la transmission de flux multimédia et la gestion de la QoS au niveau IP.

Actuellement, les frontières entre les réseaux tendent à disparaître et les services se généralisent sur tous les réseaux. Par exemple, l'Internet et la TV sont disponibles sur les réseaux de

télécommunications, les communications téléphoniques peuvent être effectuées sur Internet, etc. Cette nouvelle mouvance a fait émerger la notion de convergence des réseaux qui a pour objectif de définir un cadre global pour le regroupement des réseaux sous une seule architecture. Cette dernière permettra la cohabitation de tous les services qui seront accessibles par tous les utilisateurs en utilisant n'importe quel réseau et équipement d'accès.

Il n'existe pas, actuellement, une définition précise de la convergence des réseaux, ni une architecture définitive pour réaliser cette convergence. Cependant, il y a un consensus qui se précise sur la technologie IP pour permettre une interconnexion et une interopérabilité entre les réseaux traditionnels. Ce consensus se concrétise par deux points principaux. Le premier est le concept du Tout-IP, ou ALL-IP, qui englobe le développement de nouveaux services basés sur le protocole IP et la migration des services traditionnels sur les réseaux IP (VoIP, IPTV, etc). Le deuxième point est la définition des réseaux de nouvelle génération NGNs qui exploitent le protocole IP dans leur couche transport.

3.2.1.1 Les réseaux de nouvelle génération (NGN : Next Generation Network)

Les NGNs représentent la prochaine génération des réseaux censée réaliser la convergence totale des services en une seule architecture. L'ITU-T a publié deux recommandations concernant les NGNs. La première recommandation, Y.2001 [174], définit les principales caractéristiques d'un NGN, tandis que la deuxième, Y.2011 [175], propose une architecture fonctionnelle. Le comité technique TISPAN de l'ETSI (European Telecommunications Standards Institute) [176] a aussi défini une architecture fonctionnelle pour les NGNs largement inspirée de celle proposée par l'ITU-T.

Un NGN, simplifié dans la Figure 3-1, est constitué principalement de deux couches, une couche transport qui fournit une connectivité IP aux différents composants d'un NGN tout en garantissant une QoS de bout-en-bout, et une couche de service qui fournit les fonctionnalités de base pour le fonctionnement des services avec ou sans session. Par exemple, des fonctionnalités relatives à l'enregistrement, la notification et la présence.

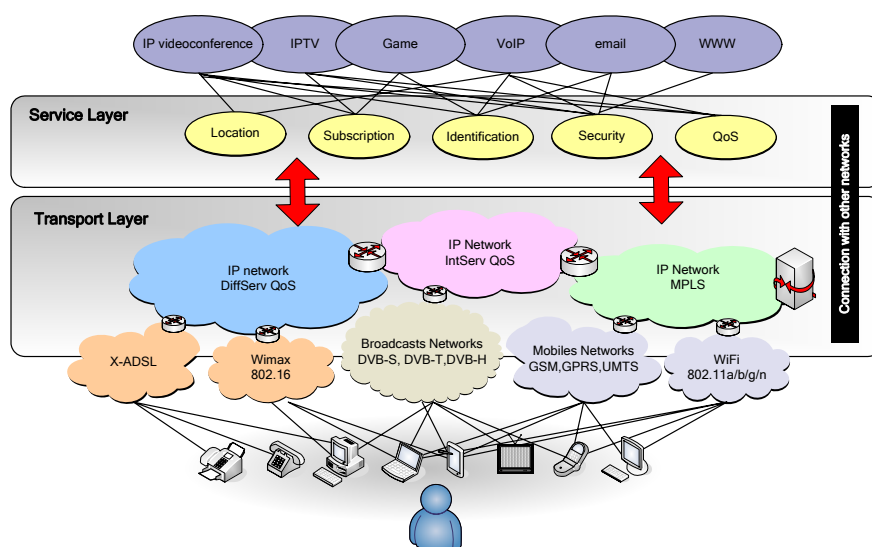


Figure 3-1 : L'architecture NGN

Ainsi, la définition d'un service est indépendante des technologies réseaux sous-jacentes. Un NGN est censé supporter n'importe quel réseau d'accès, filaire ou sans fil, et n'importe quel équipement terminal, fixe ou mobile. L'interaction entre les NGNs et les réseaux traditionnels est prise en charge par différentes passerelles d'interconnexion afin d'assurer une compatibilité avec les technologies déployées actuellement.

3.2.1.2 L'architecture IMS (Internet Multimedia Subsystem)

L'architecture IMS (Internet Multimedia Subsystem) [177] définie par le 3GPP démontre aussi le consensus sur la technologie IP pour réaliser la convergence. En effet, l'IMS se base sur le protocole IP pour permettre le déploiement d'applications et de services multimédia sur tout type de réseau fixe, mobile ou sans fil. Pour cela, l'architecture IMS, illustrée dans la Figure 3-2, est découpée en trois plans distincts qui séparent les fonctionnalités liées aux services, aux contrôles et aux réseaux. Ceci permet de passer d'un développement vertical des services qui sont définis indépendamment les uns des autres, sur tous les plans, vers un développement horizontal qui permet la réutilisation de fonctionnalités communes entre services. L'IMS se base sur les protocoles standards de l'Internet définis par l'IETF, principalement le protocole SIP pour le contrôle et l'administration des services. Le comité TISPAN a adopté l'architecture IMS dans son architecture NGN au niveau du plan de service pour assurer le contrôle des sessions multimédia. L'architecture IMS a été étendue par le comité TISPAN afin de supporter d'autres réseaux d'accès comme le xDSL (x – Digitale Line) et WLAN (Wireless Local Area Network).

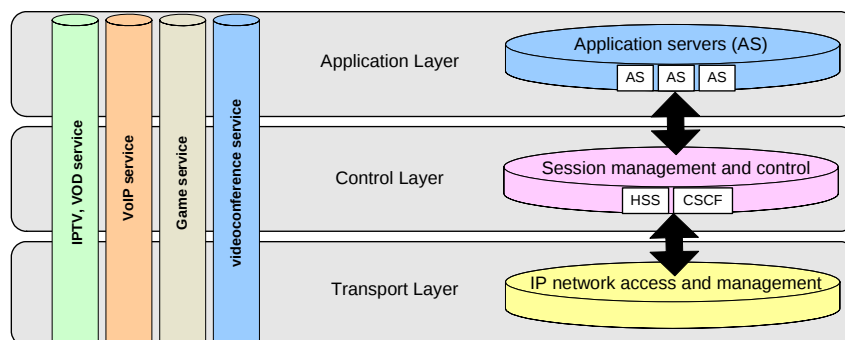


Figure 3-2 : L'architecture IMS

3.2.2 Le concept UMA (Universal Multimedia Access)

Le concept UMA est apparu comme une conséquence directe de la convergence des réseaux pour permettre l'accès aux contenus multimédia n'importe où, n'importe quand en utilisant n'importe quel réseau et n'importe quel terminal d'accès. En effet, Dans les NGNs, les réseaux d'accès peuvent avoir différentes caractéristiques et différentes capacités (le degré de fiabilité, le débit disponible, etc.). De même, les terminaux d'accès peuvent être limités en puissance de calculs, capacité de stockage, résolution d'affichage, durée de vie des batteries, etc. L'objectif du concept UMA est d'apporter une solution à toute cette hétérogénéité.

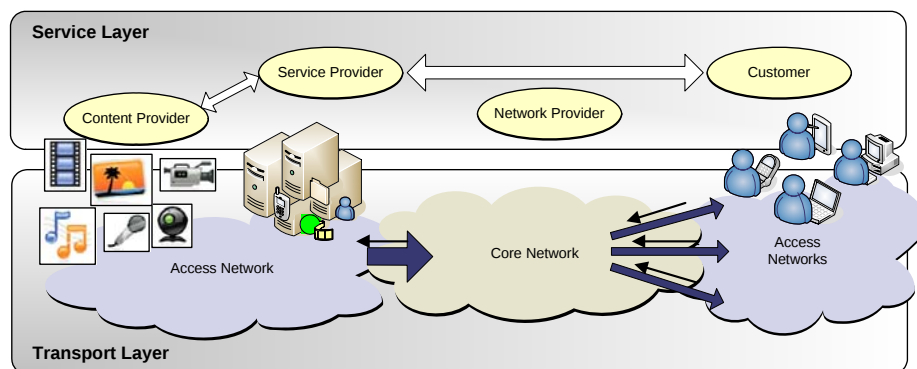


Figure 3-3 : L'architecture considérée pour le concept UMA

La Figure 3-3 illustre l'architecture préconisée dans le concept UMA, elle est inspirée principalement de l'architecture NGN. Le niveau service de cette architecture est constitué de quatre principaux acteurs :

- **Le fournisseur de contenu (CP : Content Provider)** : Il représente la source en fournissant du contenu multimédia suivant un format spécifique. Il fournit aussi toutes les descriptions relatives à ce contenu. Ces descriptions peuvent informer sur le contenu lui-même, son format, les droits d'auteur, la qualité du contenu, la possibilité d'adaptation, etc.
- **Le fournisseur de service (SP : Service Provider)** : Il est considéré comme le coordinateur principal de l'architecture UMA. Il se charge de la livraison du contenu multimédia du fournisseur de contenu au consommateur en passant par le fournisseur de réseau. Pour cela, il peut être amené à négocier la QoS du service avec le fournisseur de réseau suivant le contrat établi avec le consommateur. Il décide aussi d'éventuelles adaptations du contenu multimédia pour maintenir la QoS perçue par le consommateur.
- **Le fournisseur de réseau (NP : Network Provider)** : Il fournit une connectivité IP de bout-en-bout entre les différents acteurs tout en garantissant la QoS négociée avec le fournisseur de service pour le transport du contenu multimédia.
- **Le consommateur (CC : Content Consumer)** : C'est le dernier maillon de la chaîne qui reçoit et consomme le contenu multimédia. Il est lié au fournisseur de service par un contrat de service qui peut être négocié dynamiquement lors de la demande d'un service particulier.

Le niveau réseau est décomposé en deux grandes parties : le cœur du réseau et les réseaux d'accès. Le cœur de réseau est un réseau IP constitué de plusieurs domaines, fournissant de grandes capacités et plusieurs niveaux de QoS. Tandis que les réseaux d'accès sont limités en capacités et, dans la majorité des cas, ils n'offrent aucune QoS.

L'objectif de l'UMA est de déterminer la meilleure adaptation du contenu multimédia en considérant plusieurs paramètres (les caractéristiques du réseau, du terminal, du contenu multimédia) et différents acteurs (le fournisseur de contenus, le fournisseur de services, le fournisseur de réseau et le consommateur). Ceci permet de maximiser le degré de satisfaction du consommateur final et d'augmenter les revenus des fournisseurs.

3.2.2.1 Les techniques d'adaptation

L'adaptation du contenu multimédia est à la base du concept UMA. Cette adaptation englobe tout type de données qui peut être transmis dans un service : texte, image, vidéo et audio. Le principal avantage de l'adaptation est la réduction de l'espace de stockage dans les serveurs de contenus qui ne doivent plus stocker plusieurs formats du contenu original. De plus, l'adaptation est le seul moyen pour personnaliser les contenus temps réel.

Dans le cadre de cette thèse, nous nous intéressons aux techniques d'adaptation appliquées aux contenus vidéo. Les auteurs dans [178] proposent une classification de ces techniques suivant différents niveaux : sémantique, structurel et signal. Ci-dessous, nous donnons une brève description de chaque niveau d'adaptation :

- **L'adaptation au niveau sémantique :** Cette catégorie propose d'adapter la vidéo suivant l'importance des événements qu'elle véhicule. Par exemple, les scènes de but dans un événement sportif, le journal d'information dans un flux TV, un zoom sur la partie importante d'une image vidéo (joueurs, personnage, etc.), ou bien les scènes de mouvement dans une vidéo de surveillance. L'importance de la scène peut être définie par le fournisseur de contenu ou le consommateur. Ce type d'adaptation se base sur des Metadata qui sont inclus dans le flux vidéo afin de permettre une adaptation dynamique.
- **L'adaptation au niveau structurel :** Cette catégorie d'adaptation construit des résumés de vidéo en changeant sa structure. Par exemple, en affichant séquentiellement des images représentatives ou des images clés à partir d'une séquence vidéo. Ces images peuvent être déterminées suivant différents critères statistiques. Une autre technique consiste à construire une image mosaïque à partir de plusieurs images. Cette technique est très intéressante dans des scènes où un objet se déplace sur un fond statique. Dans l'image mosaïque, l'objet est dupliqué plusieurs fois sur le fond statique suivant la direction de son déplacement.
- **L'adaptation au niveau signal :** L'adaptation au niveau signal englobe le transcodage pour le codage standard et l'adaptation du nombre de couches pour le codage hiérarchique. Le transcodage permet de changer le format d'une vidéo vers un autre ou vers le même format en changeant plusieurs paramètres : la résolution temporelle (nombre d'images par seconde), la résolution spatiale (la taille de l'image), la qualité de l'image (SNR). L'adaptation du nombre de couches, quant à elle, permet uniquement de changer la résolution temporelle, spatiale ou SNR d'une vidéo sans changer son format de codage.

3.2.2.2 Les approches d'adaptation

Pour réaliser le concept UMA, il faut répondre à deux questions importantes d'un point de vue architectural : où décider de l'adaptation ? et où exécuter cette adaptation ? Dans la littérature, ces deux fonctions sont généralement localisées sur la même entité physique. C'est-à-dire que, l'entité qui décide de l'adaptation, se charge de l'exécuter.

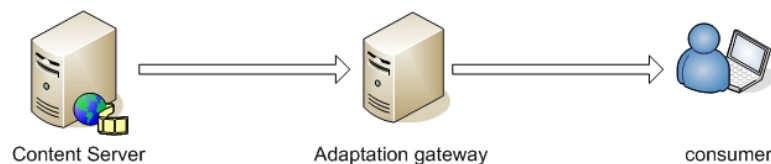


Figure 3-4 : La localisation de l'adaptation

La Figure 3-4 illustre les trois localisations qui se distinguent, actuellement, pour exécuter l'adaptation : au niveau du serveur de contenu (Content Server), au niveau du client (Consumer) ou au niveau d'une passerelle intermédiaire (Adaptation Gateway). Une brève description de chaque approche est donnée ci-dessous.

3.2.2.2.1 L'adaptation au niveau serveur

Dans cette approche, l'adaptation est exécutée au niveau du serveur de contenu. Ce dernier doit s'informer sur les capacités du réseau et du terminal d'accès afin de décider et d'exécuter la meilleure adaptation possible. Le principal avantage de cette approche est qu'elle fournit plus de contrôles sur les contenus multimédia en limitant l'altération du contenu suivant les exigences du propriétaire. Cependant, Cette approche n'est pas scalable parce que le serveur doit prendre en charge deux fonctionnalités : la fourniture du contenu et son adaptation. De plus, l'adaptation de contenu consomme énormément de ressources (CPU, mémoire, etc.) sur le serveur. Ceci limite le nombre de consommateurs qui peuvent être pris en charge. Enfin, le serveur doit mettre en œuvre des protocoles pour récupérer les informations relatives au réseau et au terminal du consommateur.

3.2.2.2.2 L'adaptation au niveau consommateur

Cette approche propose d'exécuter l'adaptation au niveau du consommateur suivant ses capacités. L'avantage de cette approche est que les capacités relatives du terminal sont directement disponibles au niveau du consommateur. Cependant, l'adaptation ne prend pas en considération les capacités du réseau. De plus, les terminaux qui ont une capacité limitée ne peuvent pas exécuter un transcodage très gourmand en ressource de calcul et en mémoire de stockage.

3.2.2.2.3 L'adaptation au niveau passerelle

L'adaptation effectuée au niveau d'une passerelle est un compromis entre les deux approches précédentes. La passerelle est localisée entre le serveur de contenu et le consommateur et elle est dédiée à l'adaptation. Ceci permet, d'une part, d'alléger le serveur de contenu de la charge de l'adaptation, et d'autre part, d'adapter le contenu suivant les capacités du réseau. Le déploiement de plusieurs passerelles d'adaptation permet d'assurer la mise à l'échelle, chaque passerelle se charge d'un certain nombre de consommateurs. Cette approche peut poser quelques inconvénients pour les échanges sécurisés dans lesquels le contenu est crypté. La passerelle doit pouvoir décrypter le contenu pour l'adapter et le crypter à nouveau pour le retransmettre au consommateur.

3.2.2.3 L'adaptation MPEG-21

L'importance grandissante de l'adaptation des objets multimédia pour réaliser le concept UMA a obligé le groupe MPEG à standardiser cette adaptation dans la partie 7 de la norme MPEG-21, appelée DIA (Digital Item Adaptation).

En effet, après le standard MPEG-7 [179] qui permet de décrire des contenus multimédia afin de faciliter la recherche et l'indexation de ces contenus, le groupe MPEG définit, actuellement, la norme MPEG-21 [180] qui a pour objectif de définir une architecture standardisée pour la livraison et la consommation de contenus multimédia. L'architecture doit permettre l'interopérabilité entre différents acteurs, du créateur de contenu jusqu'au consommateur final. MPEG-21 introduit deux nouveaux concepts importants : le DI (Digital Item) et l'utilisateur (user). Le DI représente l'unité de base qui est distribuée, échangée et consommée dans une architecture MPEG-21. Un DI peut être un fichier audio, vidéo, image, etc. Le concept de l'utilisateur est une abstraction de toute entité de l'architecture interagissant ou manipulant des DI. Afin de faciliter son adoption et son implémentation, le standard est composé de 18 parties indépendantes couvrant plusieurs aspects : la déclaration et l'identification d'un DI, la gestion des droits d'auteur, l'adaptation d'un DI, le streaming d'un DI, etc.

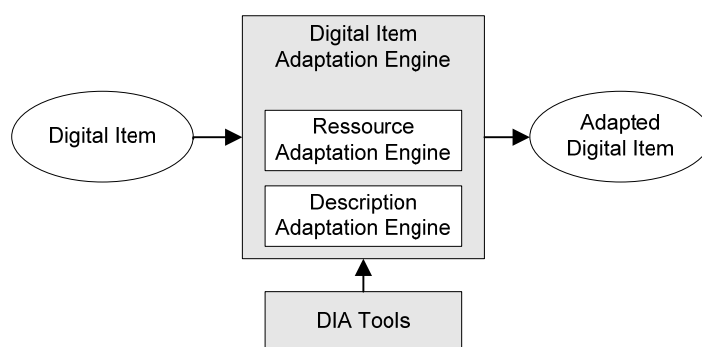


Figure 3-5 : L'architecture de MPEG-21 DIA

La partie 7 de MPEG-21 [181], appelée DIA (Digital Item Adaptation), se charge de l'adaptation de contenus multimédia. La Figure 3-5 illustre l'architecture conceptuelle d'adaptation proposée par le DIA. L'adaptation est exécutée par un DIAE (Digital Item Adaptation Engine) constitué d'un module «Description Adaptation Engine» responsable de l'adaptation de la description du DI en considérant plusieurs contraintes et d'un module «Ressource Adaptation Engine» responsable de l'adaptation du DI suivant sa nouvelle description. Cette adaptation est indépendante du format de codage du DI.

Il est important de préciser que le DIA du MPEG-21 ne définit pas de procédures ou de politiques d'adaptation, mais un ensemble d'outils pour gouverner cette adaptation. Ces outils correspondent à des langages qui permettent de générer des descriptions bien définies. La syntaxe de ces langages est définie par des schémas XML. Une description est un document XML qui correspond à un schéma particulier.

Les outils de DIA sont découpés en 7 groupes distincts : UED (Usage Environment Description), BSD (Bitstream Syntax Description), Terminal and Network Quality of service, Metadata Adaptability, Session Mobility, DIA Configuration. Nous présentons, ci-dessous, une description des principaux groupes.

3.2.2.3.1 L'outil « Usage Environment Description »

L'outil UED (Usage Environment Description) fournit des descriptions normalisées de propriétés relatives à l'utilisateur et à son environnement. L'UED couvre quatre aspects importants pour l'adaptation des DI :

- **Les caractéristiques de l'utilisateur (User Characteristics) :** L'utilisateur représente le consommateur final d'un DI. L'adaptation doit prendre en considération les caractéristiques et les préférences de l'utilisateur. L'UED reprend la description de l'utilisateur fournie par le standard MPEG-7 en l'enrichissant avec d'autres informations. Par exemple, le genre de contenu préféré, le degré de violence dans les scènes vidéos, la langue préférée, le sous-titrage, etc.
- **Les capacités du terminal (Terminal Capabilities) :** Cette description est indépendante du type du terminal qui peut être un téléphone mobile, un PDA, un PC, etc. La description couvre une large variété de caractéristiques et de capacités : le codage audio/vidéo supporté par le terminal, sa capacité d'affichage vidéo, sa capacité de sortie audio, ses caractéristiques de stockage, ses interfaces d'interactions avec l'utilisateur.
- **Les caractéristiques du réseau (Networks Characteristics) :** Les caractéristiques du réseau auquel est connecté l'utilisateur sont importantes pour l'adaptation des DI. La description de ces caractéristiques est découpée en deux parties. La première, appelée capacités du réseau « network capabilities », décrit les caractéristiques statiques d'un réseau. Par exemple, le débit maximal garanti, le taux d'erreurs moyen, les mécanismes de correction disponibles. La deuxième partie, appelée conditions du réseau « network conditions », décrit les paramètres de performances instantanés du réseau. Par exemple, le délai de transmission, la gigue, les taux de perte, le débit disponible, etc.
- **Les caractéristiques de l'environnement naturel :** cet aspect décrit l'environnement extérieur de l'utilisateur, par exemple, sa localisation, le degré de bruit audio, le degré de luminosité, etc.

3.2.2.3.2 L'outil « Bitstream Syntax Description –BSD »

Un flux de bits « Bitstream » représente une ressource multimédia avec n'importe quel format de codage. La description gBSD (generic BSD) offre une description générique (high-level) du bitstream indépendante de son format de codage. Le gBSD ne décrit pas le flux bit par bit mais son organisation générale, par exemple, une description au niveau image qui donne le type d'image (I, P, B) dans une vidéo, son emplacement dans la séquence et sa taille. La procédure d'adaptation consiste à générer un gBSD transformé à partir du gBSD original en considérant d'autres descriptions : steering description, BSD Transformation style sheet, etc. Toutes ces descriptions sont référencées dans une description racine, appelée BSDLink.

3.2.2.3.3 L'outil « Terminal and Network Quality of service »

La tâche principale de cette partie est la gestion de la QoS au cours de l'adaptation des DI aux contraintes imposées par le terminal et le réseau. Pour cela, elle se base principalement sur l'outil « AdaptationQoS » qui introduit deux notions importantes : les IOPins et les modules. Un IOPin

correspond à une variable unique qui décrit une caractéristique d'un DI, par exemple, la résolution verticale d'une image vidéo. Un module correspond à une fonction mathématique qui calcule les valeurs d'IOPins en fonction de valeurs d'autres IOPins. Trois types de module sont définis :

- **Look-up Table** : Ce module définit une table de n-dimensions qui fait correspondre toutes les valeurs possibles d'IOPins à d'autres valeurs d'IOPins. Par exemple, une table à trois dimensions qui donne l'échelle (scale) de la vidéo (1^{ère} dimension) pour chaque résolution verticale (2^{ème} dimension) et horizontale (3^{ème} dimension).
- **Stack Function** : Ce module définit une fonction qui calcule la valeur d'un IOPin en prenant comme paramètres les valeurs d'autres IOPins. Par exemple, le débit d'une vidéo à partir de sa durée et la taille du fichier vidéo.
- **Utility Function** : Ce module définit une relation formelle entre les contraintes, les opérations d'adaptation qui satisfont ces contraintes et l'utilité (le niveau de la QoS) qui résulte de ces adaptations. Chaque « utility Function » supporte un ou plusieurs opérateurs d'adaptation, par exemple le nombre de couches supérieures à supprimer dans un DI codé hiérarchiquement. Le résultat de l'adaptation est décrit par une partie utilité (utility) qui représente le degré objectif ou subjectif de la QoS, par exemple, la valeur du PSNR qui résulte d'une suppression d'une couche supérieure d'un DI.

3.2.2.3.4 *Universal Constraint Description tools*

L'UCD (Universal Constraint Description) est utilisé pour décrire les contraintes d'adaptation. Il représente le lien entre le « AdaptationQoS » et l'UED. Une description UCD définit, d'une part, des valeurs limites pour les IOPins, appelées contraintes limites, et d'autre part, des fonctions d'optimisation pour maximiser ou minimiser les valeurs de certains IOPins.

3.3 Motivations pour interconnecter les réseaux DVB-T et les réseaux IP/802.11

Durant ces dernières années, nous avons assisté au déploiement des réseaux DVB-T (Digital Vidéo Broadcast – Terrestrial) pour la diffusion TV numérique. La diffusion DVB-T permet une utilisation optimisée des ressources hertziennes, comparé à la diffusion analogique, en exploitant la compression numérique des flux TV. Ceci a permis d'introduire de nouvelles chaînes TV qui peuvent être reçues en utilisant une antenne analogique traditionnelle et un décodeur numérique bon marché en comparaison des décodeurs numériques satellitaires.

En parallèle du déploiement des réseaux DVB-T, nous avons assisté au déploiement des réseaux 802.11 au niveau des réseaux d'accès afin d'offrir une connexion IP sans fil. La technologie 802.11 est adoptée aussi bien par les professionnels que par les particuliers. Les professionnels déploient ce type d'accès dans différents espaces publics (universités, bibliothèques, gares, aéroports et centres commerciaux) afin d'offrir un accès payant ou gratuit en réduisant au maximum les coûts de déploiement. Les particuliers, quant à eux, déploient ce type de réseau dans leur maison afin de connecter différents équipements intégrant une interface de communication 802.11. Ce déploiement s'est accéléré avec l'apparition des nouvelles normes 802.11g/n qui ont

augmenté les débits de transmission à 54 Mbit/s et jusqu'à 540 Mbit/s. Cette augmentation a permis la transmission des services IP multimédia gourmands en bande passante réseau.

En considérant le contexte général qui tend vers la convergence des réseaux et l'accès universel aux services multimédia. L'interconnexion entre les réseaux DVB-T et les réseaux IP/802.11 se pose naturellement comme un défi à surmonter afin d'étendre l'accessibilité des services DVB-T pour des récepteurs mobiles et hétérogènes connectés via un réseau 802.11. Cette interconnexion engendre plusieurs problématiques vu la technologie différente des deux réseaux et la nature des données qu'ils transportent. Nous donnons, ci-dessous, les principales problématiques auxquelles il faudra faire face :

- **Le mode de transmission :** Le réseau DVB-T se base sur un seul mode de transmission qui est le broadcast. Tandis que les réseaux IP/802.11 peuvent utiliser trois mode de transmission : l'unicast, le multicast et le broadcast. Il est important de déterminer le meilleur mode de transmission des services DVB-T sur les réseaux IP/802.11 puisque chaque mode possède des avantages et des inconvénients. L'unicast ne permet pas la scalabilité parce que le flux TV est dupliqué pour chaque récepteur, par contre, il permet d'adapter le flux TV suivant les caractéristiques du récepteur. Le broadcast permet la scalabilité, cependant le flux est transmis dans le réseau même s'il n'y a aucun récepteur. De plus, le flux ne peut pas être adapté. Enfin, le multicast permet la scalabilité mais l'adaptation d'un flux pour chaque récepteur est impossible.
- **Le transport des flux TV :** Le transport des flux TV se base sur MPEG-2 TS (Transport Stream) qui utilise des paquets de petite taille (188 octets). La transmission de ces paquets dans un réseau IP/802.11 engendre un important *overhead* causé par les différents en-têtes des protocoles (UDP/IP/802.11 MAC/802.11 PHY).
- **La signalisation des services :** La signalisation des flux TV dans le réseau DVB-T se base sur des tables de signalisation MPEG-2 TS et DVB-T. Ces tables donnent des informations pour le démultiplexage des flux ainsi que des informations sur les flux eux-mêmes. Cette signalisation ne peut pas être maintenue dans les réseaux IP/802.11 qui possède des protocoles spécifiques pour la description des flux multimédia et leur signalisation.
- **L'hétérogénéité des récepteurs 802.11 :** Le réseau DVB-T diffuse des flux TV caractérisés par un format de codage audio/vidéo fixe est prédéfini. L'audio est codé suivant le format MPEG-2 layer II, le taux d'échantillonnage est de 44.1 KHz, le nombre de canaux est de 2 (stéréo) et le débit total du flux est de 192 Kbits/s. La vidéo est codée suivant le format MPEG-2 video, la résolution d'image est de 720x576 pixels, le nombre d'images par seconde est de 25 fps et le débit moyen du flux vidéo est de 2.5 Mbits/s. D'un autre côté, les terminaux de réception dans les réseaux IP/802.11 sont hétérogènes et possèdent des capacités différentes en termes de codecs, de puissance de calcul, d'affichage et de sortie audio. Si les flux TV sont transmis dans les réseaux IP/802.11 avec leur format original, certains terminaux ne pourront pas décoder ces flux. De plus, le débit fourni par les réseaux 802.11 reste limité et le débit des flux TV est relativement élevé.

- **Le maintien de la QoS au niveau des réseaux d'accès 802.11 :** Les flux TV sont des flux multimédia temps réel. La transmission de ce type de flux sur des réseaux sans fil 802.11 représente un défi important. Ceci s'explique principalement par le manque de fiabilité du canal sans fil qui est sujet à des interférences et des variations importantes durant le temps. De plus, le partage de ce canal entre plusieurs récepteurs engendre un important *overhead* qui réduit les performances de transmission.

Dans le reste de chapitre, nous allons détailler l'architecture que nous proposons pour interconnecter le réseau de diffusion DVB-T et les réseaux IP/802.11 en essayant d'apporter une solution pour chacune des problématiques citées ci-dessus.

3.4 L'architecture d'interconnexion du réseau DVB-T aux réseaux IP/802.11 WLAN

L'architecture d'interconnexion du réseau DVB-T et des réseaux IP/802.11 est présentée dans la Figure 3-6. L'objectif principal de cette architecture est de rendre accessible les services DVB-T à partir des réseaux WiFi 802.11 à l'instar de n'importe quel service IP. Le service IPTV doit s'intégrer dans une infrastructure IP existante d'une manière transparente en utilisant les standards des services IP.

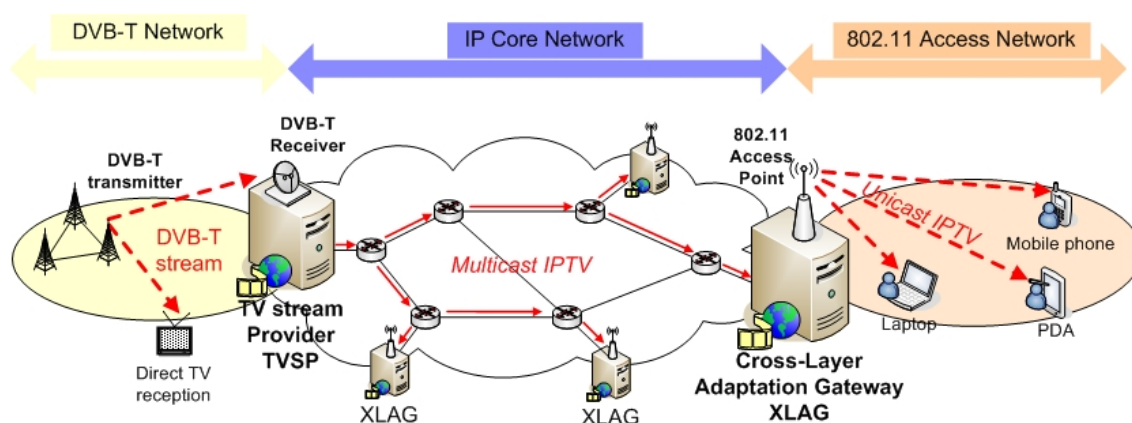


Figure 3-6 : L'architecture d'interconnexion des réseaux DVB-T et des réseaux IP/802.11

L'architecture est décomposée en trois segments : le réseau de broadcast DVB-T d'où sont récupérés les flux TV numériques, le cœur du réseau IP sur lequel transitent les flux IPTV et enfin les réseaux d'accès qui sont limités dans notre architecture aux réseaux 802.11. Le TVSP (TV Stream Provider) assure l'interconnexion entre le réseau DVB-T et le cœur du réseau IP. Il est responsable de la réception des flux DVB-T et de leur retransmission dans le cœur du réseau en flux IPTV. Ces derniers sont reçus par des passerelles XLAG (Cross Layer Adaptation Gateway) déployés entre le cœur du réseau et les réseaux d'accès. Ces passerelles représentent des entités clés de notre architecture. L'objectif principal d'un XLAG est de fournir des flux IPTV personnalisés à des stations 802.11 hétérogènes et mobiles. Deux aspects importants sont considérés dans la définition de notre architecture : la scalabilité et l'hétérogénéité.

Afin d'assurer la scalabilité, le mode de transmission multicast est utilisé pour transmettre les flux IPTV dans le cœur du réseau vers les XLAG. Ceci évite la duplication des flux transmis à

plusieurs XLAG et préserve, par conséquent, la bande passante disponible dans le cœur du réseau. Ainsi, chaque flux TV représente un groupe multicast et les XLAG doivent s'abonner aux groupes pour recevoir les flux TVs. Cependant, l'abonnement d'un XLAG à un groupe dépend principalement de l'existence de récepteurs du flux TV transmis par ce groupe dans le réseau d'accès géré par ce XLAG. Si aucun récepteur n'est intéressé par un flux TV, le XLAG se désabonne du groupe.

La transmission multicast n'est pas étendue au réseau d'accès. Les flux TV sont transmis en unicast dans ce dernier segment de l'architecture afin de faire face à l'hétérogénéité des récepteurs. Le XLAG offre un service LoD (Live on Demand) qui permet de transmettre en unicast un flux TV à la demande. Ce flux TV est personnalisé en considérant plusieurs caractéristiques du récepteur (la capacité du terminal, le débit du récepteur, les préférences de l'utilisateur, etc.). Le déploiement d'un XLAG dans chaque réseau d'accès fournit deux avantages. Le premier est la limitation du nombre de récepteurs pris en charge, puisque le XLAG doit exécuter, pour chaque flux TV transmis, des adaptations gourmandes en ressource de calcul. Le deuxième avantage est la communication directe entre le XLAG et le récepteur afin de récupérer facilement les caractéristiques du récepteur et d'exécuter rapidement les adaptations nécessaires.

Dans la suite de cette section, nous commençons par présenter les services IPTV qui suscitent un grand intérêt de la part de différents acteurs (fournisseurs de contenus, de services et de réseaux). Par la suite, nous détaillons tous les segments de l'architecture proposée dans la Figure 3-6 ainsi que toutes les entités la constituant.

3.4.1 Les services IPTV

Le service IPTV est le terme générique utilisé pour désigner les services de transmission audio/vidéo, *live* ou préenregistrés, sur des infrastructures IP. Actuellement, il n'existe pas une définition définitive d'IPTV, puisque la normalisation de ce dernier est encore un chantier ouvert dans plusieurs organismes de standardisation : ITU-T avec l'IPTV FG (Focus Group) et ETSI avec les récents drafts DVB-IPI (IP Infrastructure). Durant sa deuxième réunion, tenue en Corée du sud en 2006, l'IPTV FG a abouti à un consensus pour la définition d'IPTV [182] :

« IPTV est défini comme étant des services multimédia, télévision/vidéo/audio/texte/graphiques/données, délivrés sur un réseau IP contrôlé afin de fournir le niveau de QoS/QoE, sécurité, interactivité et fiabilité exigé »

Ainsi, IPTV couvre un large panel de services, comme la VoD, le LoD, le IPTV multicast, le IPTV pairs-à-pairs. Chaque service se base sur une architecture bien définie, un mode de transmission spécifique et des protocoles de contrôle appropriés. Nous présentons, ci-dessous, une brève description de chaque service :

- **VoD (Video on Demand)** : Le service VoD permet aux utilisateurs de recevoir à la demande des vidéos préenregistrées et stockées sur des serveurs. Le VoD suit une architecture client/serveur et le client doit explicitement demander la vidéo en utilisant un protocole de contrôle de session (RTSP, SIP, MMS, etc.) pour que le serveur entame la transmission. Ce service correspond à un vidéo club virtuel, où, le client paye pour voir un film à la demande à

n'importe quelle heure sans avoir à se déplacer. Le mode de transmission adopté par ce service est principalement l'unicast.

- **LoD (Live on Demand)** : D'une manière similaire au VoD, le LoD permet aux utilisateurs de recevoir à la demande des flux vidéo *live* et non pas des vidéos préenregistrées. Il se base sur une architecture client/serveur, dans laquelle le serveur représente le point de contrôle des flux *live*. Le client utilise les protocoles de contrôle de session pour demander la transmission en unicast d'un flux particulier. Le principal inconvénient de ce service est sa scalabilité puisque le flux *live* est dupliqué dans le réseau pour chaque récepteur. Cependant, cette duplication permet d'adapter et de personnaliser le flux *live* suivant l'environnement du récepteur.
- **IPTV multicast** : Ce service se base sur le multicast IP pour la transmission de flux *live* ou de sessions vidéo préprogrammées. Ceci permet un gain considérable en bande passante comparé au service LoD. L'IPTV multicast ne se base pas sur une architecture particulière, l'utilisateur doit s'abonner à un groupe multicast pour recevoir le flux vidéo.
- **IPTV pairs-à-pairs** : Ce service reprend le principe du partage de fichiers pairs-à-pairs pour l'adapter aux services IPTV. Il se base sur la notion de pair qui est à la fois client et serveur. Le pair partage les flux vidéo avec d'autres pairs. Il peut avoir plusieurs sources et peut retransmettre le flux reçu à plusieurs pairs.

Nous avons assisté ces dernières années à l'émergence des services IPTV grâce aux fournisseurs réseaux et télécoms qui proposent le service « triplay » intégrant dans un même abonnement Internet, la téléphonie VoIP, et la télévision IPTV. Ceci, dans le but d'augmenter le nombre de services offerts aux abonnés en exploitant la même infrastructure IP qui servait exclusivement à Internet.

La technologie IPTV promet de révolutionner notre manière de recevoir et de consommer les flux TV. La voie de retour disponible dans les réseaux IP va permettre à l'utilisateur d'interagir avec l'opérateur afin de choisir, indépendamment des autres, les programmes qu'il veut recevoir. Nous pouvons dire que le IPTV est orienté davantage utilisateur en comparaison avec la diffusion TV traditionnel qui est orientée exclusivement opérateur, puisque ce dernier décide seul des programmes que l'utilisateur peut recevoir à un moment précis. De plus, l'IPTV introduit de nouvelles fonctionnalités qui auraient été impossibles de mettre en œuvre auparavant. Les principales fonctionnalités sont listées ci-dessous :

- **Le Time-shift TV** : Cette fonctionnalité permet de contrôler la lecture des flux TV *live*. Ainsi, l'utilisateur peut effectuer des pauses sur des programmes *live*. Le reste du programme est enregistré en local grâce à un équipement appelé Set-Top Box (STB) qui se charge de recevoir les flux et de les décoder. L'utilisateur peut reprendre la lecture du flux ultérieurement.
- **Network based Private Video Recorder (nPVR)** : Cette fonctionnalité introduit un serveur dans le réseau qui capture et enregistre les flux TV *live*. Le client pourra par la suite accéder à n'importe quel programme suivant son choix. Contrairement au Time-shift TV, ce service n'oblige pas l'utilisateur à avoir un dispositif de stockage (disque dur) dans sa Set-Top Box.
- **La télévision payante à la carte (Pay-per-view)** : Cette fonctionnalité propose des programmes à la carte. L'utilisateur qui souhaite voir un programme, doit auparavant le

commander et le payer. Les prix peuvent varier considérablement d'un programme à un autre suivant son importance.

- **La TV interactive** : Cette fonctionnalité introduit un nouveau type de programme qui se base sur l'interactivité. Les téléspectateurs font partie du programme et jouent un rôle dans son déroulement. Par exemple, des jeux en groupe à grande échelle dans lesquels les téléspectateurs peuvent participer à partir de leur maison, des émissions de télé-achat où les téléspectateurs peuvent acheter des articles en temps réel ou des films avec des fins différentes et l'utilisateur peut choisir la fin qu'il veut suivant son humeur.

3.4.2 Les réseaux DVB-T (Digital Video Broadcast- Terrestrial)

Le projet DVB (Digital Video Broadcasting) [183] est un consortium industriel qui regroupe plus de 270 industriels et organismes de régulation répartis sur 35 pays. Les travaux de ce groupe ont débuté en 1993, ils ont pour objectif la définition de standards techniques libres de droit pour la diffusion (broadcast) de services TV et données numériques sur les réseaux satellitaires, les réseaux câblés et les réseaux terrestres. Les résultats de ces travaux ont abouti à la publication de plusieurs normes suivant le standard ETSI [184]. Ces normes définissent la couche physique et la couche transport pour chaque média de transmission : la DVB-S/DVB-S2 pour la diffusion par satellite, la DVB-C pour la diffusion par câble, la DVB-T pour la diffusion terrestre pour des terminaux fixes et la DVB-H pour la diffusion terrestre pour des terminaux mobiles. La couche physique définit les techniques de codage et de modulation du signal et la couche transport définit la signalisation qui sert à identifier les flux de données. Les standards DVB-S/C/T se basent sur la norme MPEG-2 pour le codage et le transport des flux audio/vidéo. Par contre, le standard DVB-H, plus récent, utilise la norme H.264 qui offre une meilleure qualité vidéo avec des débits relativement bas.

Dans le cadre de notre travail, nous nous sommes intéressé au standard DVB-T déployé actuellement dans plusieurs pays Européens pour la diffusion de la TV numérique qui, à terme, remplacera définitivement la TV analogique. DVB-T peut exploiter des canaux de différentes largeurs : 5MHz, 6MHz, 7MHz ou 8MHz. La modulation du signal se base sur COFDM (Coded Orthogonal Frequency Divisional Multiplexing) qui utilise un codage correcteur d'erreurs (1/2, 2/3, 3/4, 5/6, ou 7/8) associé à un entrelacement entre fréquences 2K (2048 porteuses), 4K (4096 porteuses), ou 8K (8192 porteuses). La modulation des bits en symboles exploite les schémas de modulation suivants : QPSK, 16 QAM et 64 QAM. Les symboles sont espacés par des intervalles de garde (Guard Interval) dont la longueur peut être de 1/32, 1/16, 1/8 ou 1/4 de la longueur d'un symbole. Ceci permet de mieux distinguer les symboles successifs. Le débit offert par le réseau DVB-T peut varier de 5 Mbit/s jusqu'à 32 Mbit/s suivant les techniques de codage, de modulation et d'intervalle de garde utilisés.

En juin 2006, le groupe DVB a désigné un groupe de réflexion, appelé TM-T2 (Technical Module on Next Generation DVB-T), pour développer la nouvelle génération des réseaux DVB-T, appelé DVB-T2. Suivant le « call for technologies » publié en avril 2007, le DVB-T2 utilisera des techniques de codage et de modulation plus sophistiquées et permettra une amélioration de la charge utile d'environ 30% par rapport à la DVB-T sous les mêmes conditions du canal de diffusion.

3.4.2.1 Le MPEG-2 Transport Stream

DVB-T se base sur MPEG-2 TS (Transport Stream) pour le transport de l'audio, de la vidéo et des données. MPEG2-TS permet de multiplexer plusieurs flux élémentaires dans un seul flux de transport composé de plusieurs paquets TS de taille fixe de 188 octets. La structure d'un paquet TS est détaillée dans l'annexe A.3. Il est constitué d'un en-tête de 4 octets et d'une charge utile de 184 octets. Les flux élémentaires présents dans un flux TS (multiplexe) sont identifiés par le champ d'en-tête PID (Packet IDentifier). MPEG-2 TS se base sur la notion de programme, chaque programme étant constitué de plusieurs flux élémentaires.

La signalisation dans MPEG-2 TS est assurée par des tables de signalisation, appelées PSI (Program Specific Information). Les tables PSI sont envoyées périodiquement, elles donnent une description des programmes présents dans le flux TS et une description des flux élémentaires qui composent ces programmes. Une table PSI est identifiée par un PID unique et elle est constituée de plusieurs sections de tailles différentes. Une brève description de ces tables est donnée ci-dessous :

- **PAT (Program Association Table) :** Elle donne le PID de chaque programme présent dans le multiplexe ainsi que le PID des tables PMT et NIT. Elle représente le point d'entrée pour la recherche des programmes et la valeur de son PID est standardisée 0x00.
- **CAT (Conditional Access Table) :** Elle informe sur le système d'accès restreint utilisé dans le multiplexe. Le PID de CAT est standardisé 0x01.
- **PMT (Program Map Table) :** Elle donne les PID des flux élémentaires de chaque programme. Le PID de cette table est donné par la table PAT.
- **NIT (Network Information Table) :** Elle fournit des informations sur le réseau physique utilisé pour la transmission du multiplexe.

3.4.2.2 Les tables de signalisation DVB

DVB a défini d'autres tables de signalisation en plus de celles définies par MPEG2-TS. Les principales tables sont décrites ci-dessous :

- **BAT (Bouquet Association Table) :** Elle fournit des informations sur des bouquets de programmes.
- **SDT (Service Description Table) :** Elle fournit une description d'un programme. Par exemple, le nom du programme, le fournisseur du programme, etc.
- **EIT (Event Information Table) :** Elle fournit des informations sur les événements ou des programmes. Par exemple, le nom de l'événement, l'heure et la date à laquelle débute l'événement, sa durée, etc.
- **RST (Running Status Table) :** Elle fournit des informations sur l'état des événements programmés.

3.4.3 L'interconnexion entre le réseau DVB-T et le cœur du réseau IP

3.4.3.1 Le TVSP (TV Stream Provider)

La tâche principale du TVSP est de récupérer les flux DVB-T pour les retransmettre dans le cœur du réseau IP. Pour cela, les services DVB-T sont transformés en service IPTV multicast. Cette transformation intervient uniquement sur la méthode de transmission et la méthode de signalisation des flux TV qui doivent être adaptées suivant les standards du monde IP en exploitant les protocoles de l'IETF. Par contre, le contenu et le format des flux TV ne subissent aucune adaptation au niveau du TVSP. Le format de codage est MPEG-2 et le débit moyen des flux est 192 Kbit/s pour l'audio et 2.5 Mbit/s pour la vidéo.

Afin de récupérer tous les services se trouvant sur tous les multiplexes, le TVSP utilise plusieurs récepteurs DVB-T. En effet, la bande de fréquences exploitée par le DVB-T est décomposée en plusieurs canaux. Chaque canal transmet un multiplexe composé de plusieurs services (4 ou 5 chaînes par multiplexe). Un récepteur DVB-T peut écouter un seul canal à la fois, ce qui signifie qu'un récepteur DVB-T ne peut récupérer qu'un seul multiplexe, d'où l'utilisation de plusieurs récepteurs.

Un service représente un flux TV composé d'un flux audio, un flux vidéo, et dans certain cas, un flux texte pour le sous-titrage. Chaque flux est constitué de paquets TS possédant le même PID. Chaque service possède une table PMT qui décrit ses composantes. Le service possède aussi une description supplémentaire, orientée utilisateur, dans la table SDT.

Dans le cadre de la transformation qu'effectue le TVSP, chaque service est transformé en une session multimédia transmise dans un groupe multicast. Afin de minimiser le traitement des flux au niveau du TVSP, le transport des flux basé sur le MPEG-2 TS est maintenu. C'est-à-dire que les flux gardent leur structure initiale en paquets TS et ils n'utilisent pas une nouvelle méthode d'encapsulation pour être transmis dans le réseau IP.

Chaque session utilise une adresse IP multicast et tous les flux composant la session sont transmis sur le même port en utilisant le protocole UDP. Ces deux derniers choix s'expliquent principalement par le maintien de l'en-tête TS qui permet, d'une part, de démultiplexer les flux audio, vidéo et texte en utilisant le champ PID, et d'autre part, de reconstituer d'ordre des paquets TS en utilisant le champ CC (Continuity Counter) (voir l'annexe A.3). Ainsi, l'utilisation du protocole RTP est inutile et l'utilisation d'UDP permet de réduire l'*overhead* (12 octets pour RTP contre 8 octets pour UDP). Cette *overhead* est réduit d'avantage en regroupant dans un seul paquet UDP plusieurs paquets TS. En prenant comme référence le MTU (Maximum Transfert Unit) de Ethernet de 1500 octets, un paquet UDP peut encapsuler 7 paquets TS de 188 octets afin d'avoir une taille de paquet UDP (en-tête inclus) égale à 1324 octets.

La description de la session multimédia multicast est assurée par le protocole SDP (session Description Protocol). Cette description reprend des informations de la table PMT du service qui est transformée en session multimédia. Ces informations correspondent principalement au PID et au format de codage des flux composant le service. La description reprend aussi des informations de la table SDT qui donne des renseignements supplémentaires sur les services, par exemple, le nom du service, le fournisseur du service et l'état du service. La description SDP est indispensable

pour le démultiplexage de la session multimédia au niveau des récepteurs. Dans notre architecture, les récepteurs correspondent aux passerelles d'adaptation XLAG.

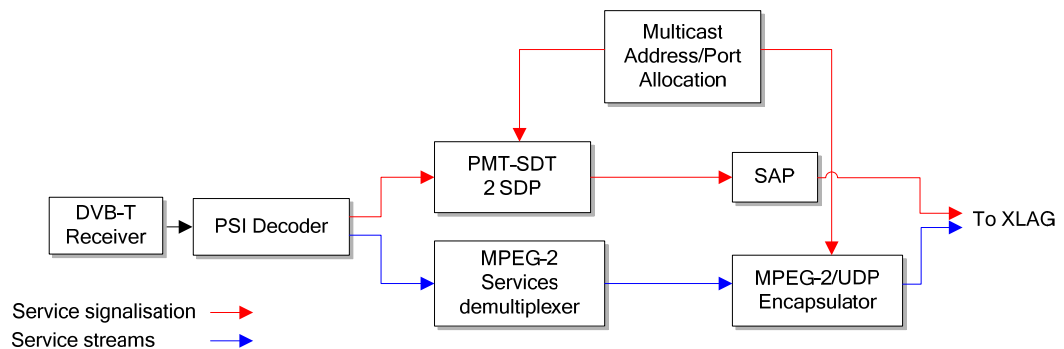


Figure 3-7 : L'architecture fonctionnelle du TVSP

La Figure 3-7 donne l'architecture fonctionnelle du TVSP. Elle est composée de 7 modules principaux :

- **DVB-T Receiver :** Ce module représente le pilote du récepteur DVB-T. Il se charge de la réception et du traitement du signal DVB-T au niveau physique afin de reconstituer les paquets TS d'un multiplexe. Ces paquets sont transmis au module « PSI Decoder ».
- **PSI Decoder :** Le décodeur PSI se charge du décodage du flux MPEG-2 TS afin de reconstruire les tables de signalisation PSI et de filtrer les paquets TS transportant les services. La première table décodée est la PAT (PID 0x00) qui contient les PID des autres tables : PMT et SDT. La PMT et la section SDT de chaque service sont transmises au module « PMT-SDT 2 SDP ». Les paquets TS contenant les services sont transmis au module « MPEG-2 services demultiplexer ».
- **Multicast Address/Port Allocation:** Ce module se charge de la gestion et de l'allocation des adresses multicast et des ports de destinations utilisés par les sessions multimédia. Une adresse et un port de destination sont alloués pour chaque service. Les politiques de gestion et d'allocation peuvent être fixées par le fournisseur de service ou le fournisseur de réseau.
- **PMT-SDT 2 SDP:** La tâche principale de ce module est de générer, pour chaque service, une description SDP à partir des informations contenues dans sa table PMT et sa section SDT. Ces dernières sont fournies par le module « PSI decoder ». L'adresse multicast et le port de destination de chaque session renseignée dans le SDP sont fournis par le module « Multicast Address/Port Allocation ». La description SDP est transmise au module « SAP ».
- **SAP (Session Announcement Protocol):** Ce module se charge de la transmission des descriptions SDP générées par le module « PMT-SDT 2 SDP » en utilisant le protocole SAP. Le protocole SAP envoie des annonces de services sur une adresse multicast réservée à cet effet.
- **MPEG-2 Services Demultiplexer :** Ce module démultiplexe les services présents dans le flux MPEG-2 TS fourni par le module « PSI decoder » en se basant sur les tables PMT. Les flux appartenant au même service sont envoyés au module « MPEG-2/UDP encapsulator ».

- **MPEG-2/UDP Encapsulator** : L'encapsulateur MPEG-2/UDP se charge de la transmission multicast d'un service sur le réseau IP. L'adresse multicast et le numéro de port employés par le service sont fournis par le module « Multicast Address/Port Allocation ». Le module utilise ces deux informations pour transmettre les paquets UDP. Dans chaque paquet UDP, le module encapsule 7 paquets TS sans faire une distinction entre les flux (audio/vidéo).

3.4.3.2 La transmission multicast dans le cœur du réseau

La transmission multicast dans les réseaux IP se base principalement sur deux types de protocoles : un protocole de gestion des groupes multicast au niveau des LANs fonctionnant entre les stations et le routeur d'accès et un protocole de routage multicast déployé entre les routeurs d'un WAN. Les protocoles de gestion de groupe définis par l'IETF et largement déployés actuellement sont le protocole IGMP v1/v2/v3 (Internet Group Multicast Protocol) [185] pour l'IPv4 et le protocole MLD v1/v2 (Multicast Listener Discovery) [186] pour IPv6. L'IETF a défini également un certain nombre de protocole de routage multicast : DVMRP (Distance Vector Multicast Routing Protocol) [187], MOSPF (Multicast Open Shortest Path First) [188], PIM (Protocol Independant Multicast) [189], et BGMP (Border Gateway Multicast Protocol) [190].

La Figure 3-8 montre la mise en œuvre des protocoles multicast dans notre architecture avec deux scénarios de fonctionnement très simple : la connexion et le départ d'un XLAG d'un groupe multicast. La transmission mutlicast s'effectue entre le TVSP et les passerelles d'adaptation XLAG. Nous supposons que ces deux dernières entités sont reliées au cœur du réseau IP par des routeurs d'accès. Le protocole IGMP est déployé entre le XLAG et son routeur d'accès et entre le TVSP et son routeur d'accès. Le protocole multicast déployé entre les routeurs est le protocole PIM-SM (PIM Sparse-Mode).

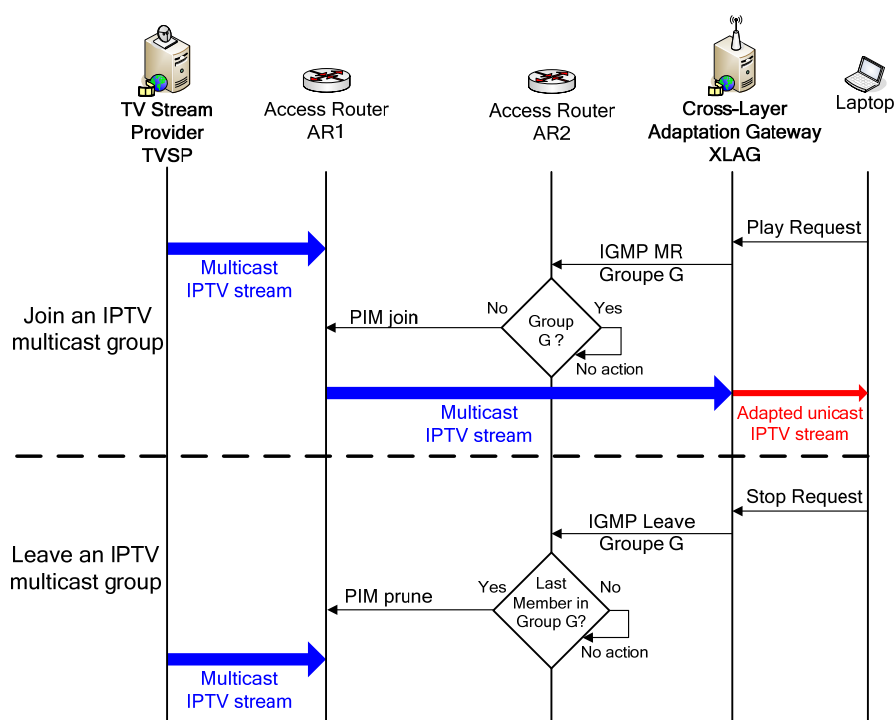


Figure 3-8 : La communication multicast entre le TVSP et le XLAG

La procédure d'abonnement d'un XLAG à un groupe multicast G se déroule comme suit : Le XLAG reçoit une demande explicite d'une de ses stations pour un flux TV particulier. Cette demande est simplifiée dans la Figure 3-8 par la requête « Play Request » (La communication entre la station et le XLAG va être détaillée ultérieurement dans la section 3.4.4.5). À la réception de cette demande, le XLAG déclenche la procédure d'abonnement au groupe G qui diffuse le flux TV demandé par la station. Cette procédure débute par la transmission d'une requête « IGMP MR » (Membership Report) en direction du routeur d'accès AR2. Ce dernier vérifie si ce groupe existe déjà au niveau du routeur. Dans le cas où le groupe existe, cela signifie que le flux IPTV du groupe multicast est déjà diffusé dans le LAN auquel est connecté le XLAG et le routeur AR2 n'a rien à faire. Dans le cas contraire, le groupe n'existe pas, le routeur d'accès utilise le protocole PIM-SM pour rejoindre l'arbre de diffusion du groupe G. Pour cela, le routeur transmet à son voisin une requête « PIM Join ». Dans une grande infrastructure, la requête est reliée jusqu'au routeur RP (Rendez-vous Point) (ou un routeur connaissant le groupe G). Dans notre scénario, nous supposons que le routeur d'accès du TVSP est le routeur RP. À la réception de la requête « PIM Join », le routeur d'accès AR1 retransmet le flux IPTV du groupe G vers le routeur AR2 qui le diffuse à son tour dans son LAN.

La procédure de désabonnement d'un XLAG du groupe multicast G suit le même schéma : Une demande d'arrêt « Stop Request » est envoyée par un récepteur au XLAG. Ce dernier arrête la réception du flux multicast et transmet automatiquement une requête « IGMP Leave » au routeur d'accès AR2. Ce dernier vérifie l'état du groupe G, c'est-à-dire, si le groupe G possède encore des membres dans son LAN. Dans le cas où le groupe G possède encore des membres, le routeur AR2 n'effectue aucune action. Dans le cas contraire, le routeur AR2 demande son élagage de l'arbre de diffusion du groupe G en envoyant au routeur AR1 une requête « PIM Prune ».

Les scénarios d'abonnement et de désabonnement détaillés ci-dessus illustrent l'avantage du multicast pour la transmission des flux TV dans le cœur du réseau. En effet, le débit du flux TV original est relativement élevé et variable. Le multicast évite la transmission d'un flux pour chaque XLAG. De plus, si aucun récepteur ne reçoit un flux TV au niveau d'un XLAG, ce dernier se désabonne du groupe afin d'éviter la transmission inutile d'un flux TV dans le réseau. Ceci permet de préserver les ressources qui peuvent être exploitées par d'autres services IP.

3.4.4 L'interconnexion entre le cœur du réseau IP et le réseau d'accès 802.11

3.4.4.1 La passerelle d'adaptation XLAG

La passerelle d'adaptation XLAG (Cross Layer adaptation Gateway) est l'entité principale de l'architecture IPTV proposée dans la section 3.4. Le XLAG gère un réseau WLAN 802.11 en mode infrastructure en intégrant directement un point d'accès. Il permet aux stations connectées à ce réseau d'accéder à un service LoD afin de recevoir les flux TV transmis par le TVSP.

Ainsi, la tâche principale d'un XLAG est de permettre à des terminaux hétérogènes et mobiles d'accéder au service LoD en préservant au maximum la QoS perçue par le consommateur final. Pour atteindre cet objectif, le XLAG, illustré par la Figure 3-9, intègre un nouveau système de transmission audio/vidéo adaptatif basé sur les interactions *Cross-layer*, appelé XLAVS (Cross Layer Adaptive Video Streaming). Ce nouveau système interagit avec toutes les couches réseaux

(application, transport, réseaux et lien d'accès) du XLAG ainsi que le terminal récepteur afin de déterminer la configuration optimale qui améliore les performances de transmission. La configuration optimale englobe l'adaptation audio/vidéo au niveau applicatif, le mapping de la QoS entre les couches, la configuration des paramètres de transmission et le paramétrage des mécanismes QoS sur toutes les couches.

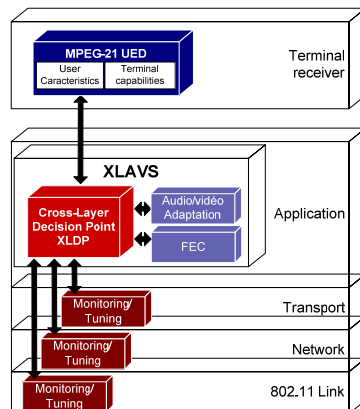


Figure 3-9 : L'architecture du XLAG

Le XLAVS se base sur l'approche intégrée, ou mixte, du concept *Cross-layer* qui permet une communication ascendante et descendante entre les couches réseaux (voir section 2.3.3). Ceci permet, d'une part, de traduire les besoins en QoS des flux transmis en métriques et mesures réseaux et lien d'accès, et d'autre part, de refléter au niveau applicatif les changements dynamiques de l'état du réseau sous-jacent.

L'originalité de notre système XLAVS est la centralisation de toutes les informations relatives à l'état du XLAG et de toutes les prises de décision concernant les adaptations et les configurations nécessaires en un seul module, appelé XLDP (Cross Layer Decision Point). La Figure 3-9 schématise cette centralisation qui a pour objectif la coordination de toutes les actions du système (configurations et adaptations) afin d'optimiser la transmission des flux TV dans le réseau géré par le XLAG. Le XLDP est situé au niveau applicatif, ce choix est motivé par plusieurs besoins :

- Permettre au nouveau système de transmission d'être déployé facilement et rapidement en étant autonome et indépendant des couches sous-jacentes. En effet, ces dernières sont implémentées au niveau des systèmes d'exploitation (transport et réseau) et au niveau des pilotes (lien d'accès) ce qui rend plus difficile leur modification pour l'intégration de nouvelles fonctionnalités.
- Une réaction et une adaptation rapide au niveau applicatif pour faire face aux variations des conditions de transmission.
- Maximiser la QoS perçue par le consommateur final. Le niveau applicatif offre le meilleur contexte pour maximiser cette QoS puisqu'il permet un contrôle direct des flux audio/vidéo.

Le XLAVS intègre deux autres modules importants pour assurer la QoS des flux TV : Le module d'adaptation qui permet de transcoder les flux audio/vidéo en temps réel et le module FEC (Forward Error Correction) qui ajoute de la redondance au niveau applicatif afin de permettre au récepteur de reconstruire les paquets perdus.

Les adaptations exécutées par le XLAVS se basent sur les outils proposés dans la partie 7 du standard MPEG-21, à savoir DIA (Digital Item Adaptation). Les principaux outils exploités sont l'UED (Usage Environment Description), l'adaptation QoS du « Terminal and Network Quality of service » et l'UCD (Usage Constraint Description). Ces adaptations ne se limitent pas aux flux audio/vidéo, elles incluent aussi l'adaptation des paramètres de couches sous-jacentes. La section suivante détaille l'architecture de notre nouveau système de transmission XLAVS.

3.4.4.2 L'architecture du XLAVS (Cross Layer Adaptive Video Streaming)

La Figure 3-10 illustre l'architecture détaillée du XLAVS qui permet une configuration dynamique de la transmission audio/vidéo. L'architecture XLAVS représente une combinaison fonctionnelle de l'adaptation MPEG-21 et des protocoles de l'IETF. Avant de décrire le principe de fonctionnement de cette architecture et les adaptations *Cross-layer* qu'elle exécute, nous présentons, ci-dessous, une description de chaque module la constituant.

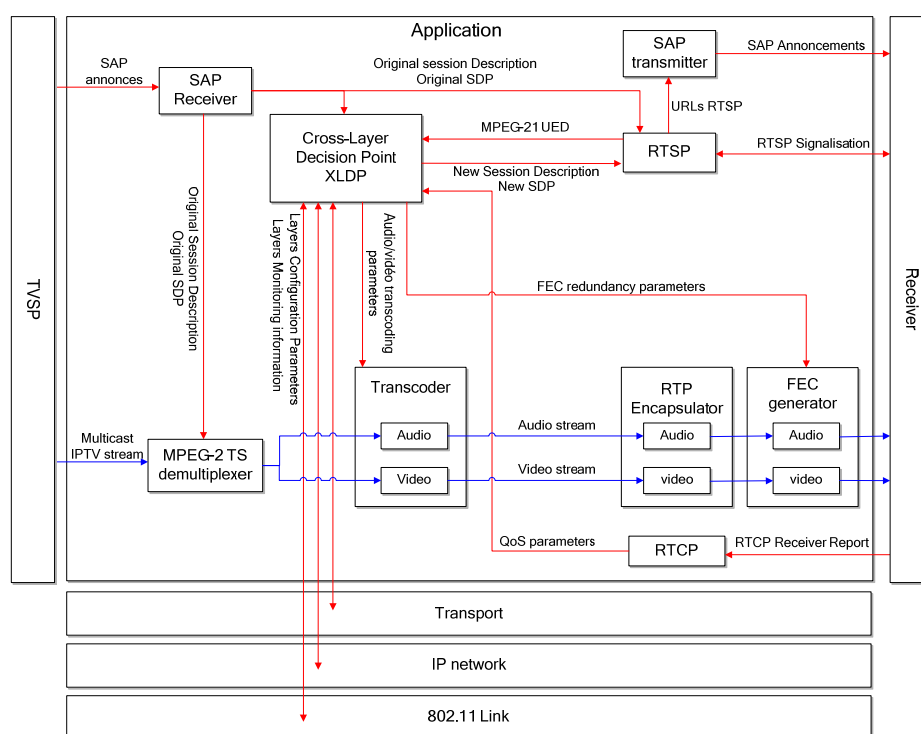


Figure 3-10 : L'architecture du XLAVS

- **SAP Receiver** : Il reçoit les annonces SAP transmises par le TVSP pour annoncer les flux IPTV multicast qu'il transmet. Ces annonces véhiculent des descriptions SDP des flux IPTV multicast. Les descriptions SDP sont transmises au module « MPEG-2 TS demultiplexer » et au module XLDP.
- **MPEG-2 TS demultiplexer** : Il démultiplexe un flux IPTV multicast, reçu du TVSP, suivant sa description SDP reçue du module « SAP Receiver ». Les flux sont transmis au module « transcoder ».
- **RTSP** : Il implémente le protocole de signalisation RTSP qui est utilisé pour le contrôle des sessions multimédia entre le XLAG et le récepteur. Dans notre architecture, ce module effectue

d'autres tâches en relation avec l'adaptation dynamique des flux audio/vidéo et l'annonce des services. Ces tâches sont listées ci-dessous :

- Le protocole RTSP véhicule une partie de la description MPEG-21 UED qui est transmise du récepteur vers le XLAG. Ce module se charge de l'extraction de cette partie de l'UED et de sa transmission vers le module XLDP.
- Le module RTSP reçoit, du module « XLDP », la nouvelle description SDP des flux audio/vidéo suivant la nouvelle adaptation. Cette nouvelle description est envoyée au récepteur durant la négociation RTSP afin qu'il puisse décoder les flux.
- Les URLs RTSP des flux IPTV disponibles sont communiquées au module SAP.
- **SAP transmitter** : Ce module transmet des annonces SAP au récepteur. Ces annonces véhiculent les URLs RTSP des flux IPTV. Ainsi, le récepteur récupère la liste des flux disponibles et leurs URLs RTSP pour y accéder.
- **Transcoder** : Ce module se charge du transcodage des flux audio/vidéo d'un format de codage vers un autre ou vers le même format en changeant plusieurs paramètres : le nombre d'images par seconde, le débit vidéo, la résolution d'image vidéo, le débit audio, le nombre de canaux audio. Les valeurs de ces paramètres sont communiquées par le module « XLDP » qui peut aussi activer ou désactiver le transcodage suivant la configuration choisie.
- **RTP encapsulator** : La fonction principale de ce module est le découpage et l'encapsulation des flux audio/vidéo en paquets RTP.
- **FEC generator** : Ce module génère les paquets de redondance FEC à partir des paquets RTP reçus du module « RTP encapsulator ». Le mécanisme FEC est indépendant pour chaque flux, c'est-à-dire que, chaque flux pourra utiliser son propre algorithme de codage et fixer son propre taux de redondance. Tous ces paramètres sont décidés et communiqués par le module « XLDP » qui peut aussi activer ou désactiver ce module.
- **RTCP** : Il récupère les rapports RTCP envoyés par le récepteur. Les informations relatives à la qualité de transmission, à savoir les taux de perte et la gigue, sont transmises au module « XLDP ».

3.4.4.3 L'architecture du XLDP (Cross Layer Decision Point)

Le XLDP représente le cœur de notre architecture. Il est responsable de toute la configuration du XLAG aussi bien au niveau applicatif qu'au niveau des couches sous-jacentes. L'architecture fonctionnelle du XLDP est illustrée dans la Figure 3-11.

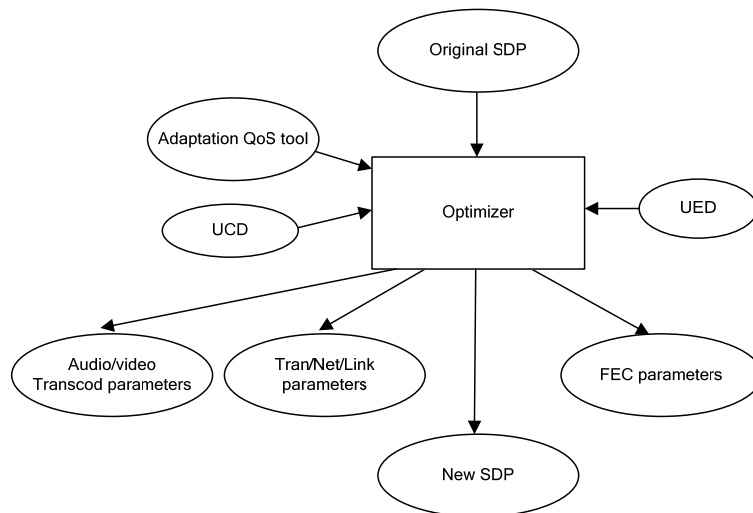


Figure 3-11 : L'architecture du XLDP

Il se base sur les outils MPEG-21 dont le concept est élargi à d'autres paramètres réseaux, puisque les outils MPEG-21 se limitent uniquement à l'adaptation des objets multimédia. Ces outils sont décrits, ci-dessous :

- **Adaptation QoS tool** : Il définit les adaptations possibles dans tout le système, c'est-à-dire, tous les paramètres du système dont les valeurs peuvent être changées par le module « Optimizer ». Cette description définit aussi les intervalles de variation de ces paramètres qui peuvent être continus ou discrets. Par exemple, la résolution d'image (QCIF, CIF et 4CIF), la classe de service au niveau IP (EF, AF), la taille des trames au niveau MAC (1500, 1024, 512 octets). Elle définit aussi les relations qui peuvent exister entre ces paramètres, par exemple, le changement de la résolution des images vidéo et le changement du débit vidéo.
- **UCD (Usage Constraint Description)** : Cette description définit les contraintes qui peuvent exister sur certains paramètres et qui doivent être respectées par le module « Optimizer ». Par exemple, le débit du flux (audio + vidéo) ne doit pas dépasser le débit disponible au niveau de la couche d'accès ou le taux de redondance FEC ne doit pas dépasser un certain seuil. Elle définit aussi les paramètres qui doivent être maximisés (le PSNR d'une vidéo) ou minimisés (les taux de perte).
- **UED (Usage Environment Description)** : Il définit l'environnement global dans lequel se déroule la transmission. Une partie de cet environnement est récupérée du récepteur, à savoir les caractéristiques de l'utilisateur et les capacités du terminal. L'autre partie, les caractéristiques statiques et dynamiques du réseau, est récupérée des différentes couches réseaux du XLAG.
- **Original SDP** : Il donne la description originale des formats de codage des flux audio/vidéo reçus du TVSP. Cette description est exploitée par le module « Optimizer » pour générer une nouvelle description SDP suivant la nouvelle adaptation des flux audio/vidéo.

- **Optimizer** : Ce module exploite les descriptions précédentes pour déterminer la configuration optimale de tous les paramètres décrits par « Adaptation QoS tool » suivant la description UED et en respectant les contraintes données par UCD. Ce module traite des problèmes d'optimisation multidimensionnelle qui sont actuellement un grand chantier de recherche. L'objectif de cette optimisation est d'améliorer les performances de transmission en se basant sur des politiques d'adaptation prédéfinies afin d'assurer une QoS maximale pour tous les récepteurs suivant les ressources disponibles.
- **Audio/video transcode parameters** : Il regroupe les paramètres de transcodage des flux audio/vidéo qui sont communiqués au module « transcoder » (Figure 3-10).
- **Trans/Net/Link parameters** : Il regroupe les paramètres de configuration des couches sous-jacentes. Les valeurs de ces paramètres sont fixées par le module « optimizer » et ils sont communiqués aux couches par des appels de fonction internes au système.
- **FEC parameters** : Il regroupe les paramètres de configuration du mécanisme FEC qui sont communiqués au module « FEC generator ».
- **New SDP** : Il correspond à la nouvelle description SDP générée par le module « Optimizer » suivant la nouvelle adaptation des flux audio/vidéo. Cette nouvelle description est communiquée au module RTSP afin qu'elle puisse être transmise au récepteur.

Le fonctionnement du XLDP est fortement couplé au fonctionnement du XLAVS. Dans la section suivante, nous allons détailler le fonctionnement du XLAVS en illustrant les adaptations *Cross-layer* possibles.

3.4.4.4 Le fonctionnement du XLAVS et les adaptations *Cross-layer*

Le XLAG représente, à la fois, un client pour le TVSP et un serveur pour les récepteurs finaux. Le fonctionnement du XLAVS est découpé en deux phases distinctes : la phase initialisation et la phase opérationnelle.

Durant la phase d'initialisation, le XLAVS découvre les flux IPTV qui sont transmis par le TVSP afin de générer les URLs RTSP qui correspondent à ces flux IPTV. Pour cela, le XLAVS exécute les tâches suivantes :

1. Il récupère grâce au module « SAP receiver » les annonces SAP transmises par le TVSP.
2. Il extrait la description SDP de chaque flux TV à partir des annonces SAP. Ces descriptions sont transmises au module RTSP afin qu'il génère une URL RTSP pour chaque flux TV.
3. L'ensemble des URLs RTSP sont transmises au module « SAP transmitter » qui transmet, à son tour, des annonces SAP au récepteur contenant ces URLs.

À la fin de cette phase, le XLAVS bascule dans la phase opérationnelle, durant laquelle, les récepteurs peuvent demander un flux TV en utilisant les URLs RTSP reçues dans les annonces SAP. Pour chaque demande, le XLAVS crée une session RTSP qui gère la transmission du flux TV entre le XLAG et le récepteur. Durant la vie d'une session RTSP, le XLAVS exécute des adaptations *Cross-layer* au cours de deux époques distinctes : Au début de la session, avant la

transmission des flux, et au cours de la session, durant la transmission des flux. Dans la suite de cette section, nous détaillons les adaptations *Cross-layer* exécutées durant ces deux époques.

3.4.4.5 L'adaptation au début de la session

Durant cette phase, le XLAVS adapte les flux audio/vidéo avant leur transmission suivant la description UED. Cette adaptation est spécifique à la session RTSP initiée par l'utilisateur afin de recevoir un flux TV particulier. La description UED est constituée de trois parties : les caractéristiques de l'utilisateur, les capacités du terminal et les caractéristiques du réseau. Ci-dessous, nous détaillons les paramètres de chaque partie qui sont considérés dans le XLAVS :

- **Les caractéristiques de l'utilisateur « UserCharacteristics »** (Figure 3-12) : Pour les préférences de l'utilisateur, un seul paramètre est considéré « ColorPreferences » qui permet à l'utilisateur de choisir de recevoir la vidéo en couleur ou bien en noir et blanc.

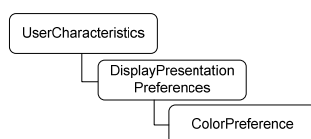


Figure 3-12 : Format des caractéristiques utilisateurs

- **Les capacités du terminal « TerminalCapability »** (Figure 3-13) : Elles sont découpées en trois grandes classes. La première classe « CodecCapability » donne les capacités de décodage du terminal. Elle comprend le format de codage « codec », le débit « Bitrate » pour le flux audio et vidéo et le nombre d'images par seconde « FrameRate » pour le flux vidéo. La deuxième classe « Display » décrit les capacités d'affichage du terminal en termes de résolution verticale et horizontale. La dernière classe décrit les capacités de sortie audio du terminal qui comprend le nombre de canaux audio (mono, stéréo) « numchannels » et le taux d'échantillonnage « SamplingFrequency ».

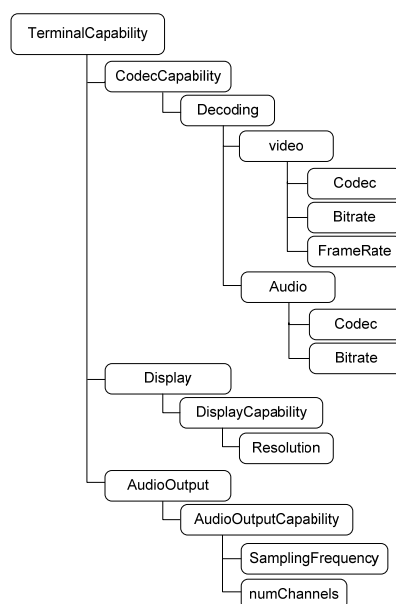


Figure 3-13 : Format des capacités du terminal

- **Les caractéristiques du réseau « NetworksCharacteristics »** (Figure 3-14) : Cette partie de l'UED est récupérée en local au niveau des couches réseaux du XLAG. Elle comprend les capacités réseaux « NetworksCapability » et des conditions réseaux « NetworksCondition ». Les capacités réseaux regroupent les paramètres invariants du réseau : Le débit maximal de la connexion « MaxCapacity », par exemple 11 Mbits/s pour la norme 802.11b et 54 Mbits/s pour la norme 802.11g, et le taux d'erreur moyen de la connexion. Tandis que les conditions réseaux regroupent les paramètres variables du réseau : le débit disponible « AvailableBandwidth », le taux de perte « LossRate », la gigue « jitter », la puissance du signal « Signal Strength ».

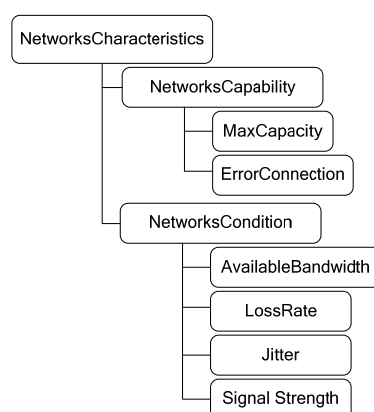


Figure 3-14: Format des caractéristiques réseaux

L'adaptation exécutée au début de la session permet de transcoder les flux audio/vidéo suivant les paramètres de l'UED excepté la partie « NetworksCondition ». La majorité de ces paramètres ne peuvent être changés durant la session, comme le format de codage audio/vidéo, la résolution des images vidéo, le taux d'échantillonnage et le nombre de canaux pour l'audio. Ceci s'explique par l'établissement d'un contexte du côté de l'émetteur et du récepteur (codage/décodage, affichage, sortie audio) qui, une fois établi, ne peut être changé sans interrompre la transmission des flux audio/vidéo. De plus, le changement de ce contexte correspond au démarrage d'une nouvelle session avec une nouvelle description UED.

Les deux premières parties de la description UED sont transmises du récepteur vers le XLAG durant la signalisation RTSP illustrée dans la Figure 3-15. Le protocole RTSP est un protocole de signalisation basé sur des requêtes/réponses en mode texte. Pour entamer la signalisation, le client récupère à partir des annonces SAP, l'URL RTSP du flux qu'il veut recevoir. Au début de la signalisation, le client transmet une requête « OPTION » pour demander au XLAG le type de requêtes qu'il supporte. Le XLAG communique ces requêtes dans la réponse. Ensuite, le client demande la description de la session SDP qui comprend des informations concernant les flux audio/vidéo (format, type d'encapsulation) et des informations concernant le transport (les adresses et les ports d'émission et de réception, les protocoles de transport utilisés). Avec le nouveau système XLAVS, la description SDP n'est pas statique, elle est générée par le module XLDP suivant l'adaptation décidée par ce dernier pour chaque nouvelle session. Pour cela, la description UED est intégrée à la requête « DESCRIBE » envoyée par le client. Le module RTSP communique l'UED reçu au module XLDP (voir Figure 3-10). Ce dernier génère le nouveau SDP qui correspond aux nouvelles adaptations. Le reste de la signalisation se déroule d'une manière standard. La nouvelle description SDP est transmise dans la réponse de la requête « DESCRIBE ». Le client et le XLAG

s'échangent les requêtes/réponses « SETUP » pour confirmer la configuration de chaque flux présent dans la session. Le XLAG transmet les flux audio/vidéo adaptés après réception de la requête « PLAY ». Le client peut stopper la transmission avec la requête « STOP » et détruire la session avec la requête « TEARDOWN ».

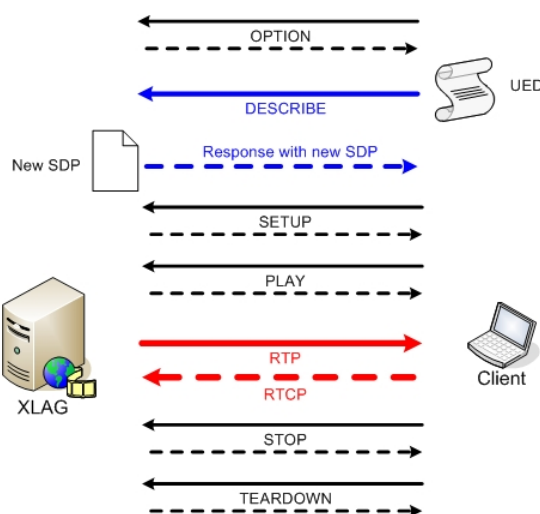


Figure 3-15 : La signalisation RTSP entre le XLAG et le récepteur

3.4.4.6 L'adaptation au cours de la session

L'adaptation *Cross-layer* au cours de la session, c'est-à-dire, durant la transmission des flux audio/vidéo est aussi importante que l'adaptation au début de la session. En effet, les réseaux sans fil 802.11 sont caractérisés par leur dynamique et les conditions réseaux peuvent varier durant la transmission des flux. Le XLAVS doit prendre en considération ces variations en adaptant les flux audio/vidéo et/ou les mécanismes de QoS en conséquence. L'objectif de ces adaptations est de préserver la QoS des flux TV perçue par l'utilisateur final durant la dégradation des conditions de transmission.

Actuellement, ces adaptations *Cross-layer* font l'objet de plusieurs travaux. Dans le cadre de notre architecture, nous avons exploré plusieurs techniques d'adaptation qui permettent d'assurer la QoS des flux TV transmis dans le réseau d'accès 802.11. Ces adaptations s'intéressent en particulier à la transmission des flux vidéo qui engendrent plusieurs problématiques dues principalement à leurs débits relativement élevés, comparés aux débits des flux audio.

Les adaptations identifiées, implémentées et expérimentées sont organisées en deux grandes classes : les adaptations *Cross-layer* ascendantes exécutées au niveau applicatif et les adaptations *Cross-layer* descendantes exécutées au niveau de la couche MAC 802.11. Nos différentes contributions concernant ces adaptations font l'objet d'une présentation détaillée dans le reste de cette thèse.

3.4.5 Mise en oeuvre de l'architecture

Un prototype de l'architecture a été réalisé en utilisant des bibliothèques et des outils libres. Le logiciel de base qui est utilisé dans les trois entités (le TVSP, le XLAVS, et le récepteur final) est l'application VLC (VideoLan Client) [191]. Ce dernier représente une solution logicielle complète

pour la lecture et la diffusion de flux multimédia. En réalité, VLC se base sur plusieurs outils libres pour offrir un large panel de fonctionnalités. Pour l'implémentation de notre architecture, nous avons modifié l'application VLC ainsi qu'un certain nombre de bibliothèques (FFmpeg, LiveMedia, libdvbpsi). Nous donnons, ci-dessous, quelques détails sur le prototype illustré par la Figure 3-16.



TVSP



XLAG



Clients

Figure 3-16: Les entités du prototype

- **L'entité TVSP** : Cette entité est mise en œuvre sur une station de travail « DELL Workstation Precision 670 » équipée d'un processeur « Intel Xeon » 3.2 GHz. La station fonctionne avec une plateforme Linux Fedora 4 [192] utilisant le noyau version 2.6.17. Pour recevoir les flux DVB-T, trois récepteurs DVB-T « WinTV Nova-T USB2 » [193] sont connectés à la station. L'architecture TVSP présentée dans la section 3.4.3.1 a été réalisée à partir d'une modification de l'application VLC et de la bibliothèque « libdvbpsi » [194]. Cette dernière permet le décodage et le filtrage des paquets TS.
- **L'entité XLAG** : Elle est, également, mise en œuvre aussi sur une station de travail « DELL Workstation Precision 670 » équipée d'un processeur « Intel Xeon » 3.2 GHz et fonctionnant avec une plateforme Linux Fedora 4 utilisant le noyau version 2.6.17. Le point d'accès intégré au XLAG correspond à une carte réseau 802.11a/b/g 3COM PCI [195] basée sur le chipset *Atheros* [196] et fonctionnant à l'aide d'un pilote libre, appelé *MadWiFi* [197]. Ce dernier offre plusieurs fonctionnalités qui permettent d'utiliser une carte WiFi comme un point d'accès ou une station. De plus, il offre plusieurs interfaces qui permettent de récupérer des métriques et de configurer des paramètres au niveau MAC et au niveau physique. L'architecture du XLAVS se base aussi sur une modification de VLC et deux principales bibliothèques, FFmpeg [198] pour le transcodage et LiveMedia [199] pour les protocoles multimédia (RTSP, RTP/RTCP, SAP). Les modifications ont permis d'intégrer toutes les fonctionnalités décrites précédemment.
- **Les stations de réception** : Trois types de terminaux sont utilisés comme stations 802.11 :
 - Une station de travail mobile Dell precision M70 (Intel pentium M 2.26 GHz), fonctionnant avec une plateforme Linux Fedora 4 utilisant le noyau version 2.6.17. La station est équipée d'une carte 802.11a/b/g 3COM (PC Card).
 - Un PDA (Personal digital Assistant) Dell axim X51 (Intel PXA270 520 MHz), fonctionnant avec un système WindowsCE version 5.0. Il intègre une interface 802.11b.

- Un téléphone PDA HP (Intel PXA270 416 MHz), fonctionnant avec un système WindowsCE version 5.0. Il intègre aussi une interface 802.11b.

Les trois terminaux utilisent une version modifiée de VLC. Les modifications permettent à l'application cliente de transmettre une partie de l'UED au XLAVS durant la négociation RTSP (voir section 3.4.4.5) et de décoder les paquets FEC pour régénérer les paquets perdus. Un compilateur croisé (cross compiler) est utilisé sur la plateforme Linux afin de générer un exécutable VLC pour la plateforme WindowsCE.

3.5 Conclusion

Dans ce chapitre, nous avons proposé une architecture complète pour l'interconnexion des réseaux de diffusion DVB-T et les réseaux IP/802.11. Cette architecture s'inscrit dans le cadre de la convergence des réseaux et des services afin d'étendre l'accessibilité des services DVB-T sur les réseaux IP en général et les réseaux 802.11 en particulier. Nous avons vu que l'interconnexion de ces deux réseaux engendrait plusieurs problématiques : la scalabilité, l'hétérogénéité et la QoS. Dans notre architecture, nous proposons des solutions adaptées pour toutes ces problématiques. Pour cela, l'architecture est décomposée en trois segments : le réseau DVB-T, le cœur du réseau IP et les réseaux d'accès 802.11.

L'interconnexion entre le réseau DVB-T et le cœur du réseau IP est assurée par l'entité TVSP (TV Stream Provider) qui a pour tâche principale la conversion des services DVB-T en service IPTV. Cette conversion englobe le transport des flux audio/vidéo et leur signalisation en partant des standards DVB-T vers les standards IP. Pour assurer la scalabilité, les flux IPTV sont transmis en multicast dans le cœur du réseau afin d'éviter la duplication des flux TV et préserver, par la même, les ressources disponibles. Ces flux sont réceptionnés par des passerelles d'adaptation déployées au niveau des réseaux d'accès 802.11, appelées XLAG (Cross Layer Adaptation Gateway). Ces passerelles assurent la connexion entre le cœur du réseau IP et les réseaux d'accès 802.11.

La fonctionnalité principale du XLAG est de fournir un service LoD pour les stations hétérogènes et mobiles connectées au réseau sans fil 802.11. Pour assurer cette fonctionnalité, le XLAG intègre un nouveau système de transmission audio/vidéo adaptatif, appelé XLAVS (Cross Layer Adaptive Streaming). Ce dernier se base sur deux nouveaux concepts : L'adaptation MPEG-21 DIA pour adapter les flux TV suivant l'hétérogénéité des récepteurs et les interactions *Cross-layer* pour maintenir la QoS des flux TV transmis sur les réseaux sans fil. L'originalité du XLAVS est la mise en œuvre du module XLDP (Cross Layer Decision Point) qui centralise toutes les prises de décision concernant les adaptations et les mécanismes de QoS exécutés au niveau du XLAG. Ceci permet d'optimiser la transmission des flux TV en maximisant la QoS perçue par les récepteurs.

Les adaptations *Cross-layer* exécutées par le XLAVS sont découpées en deux phases : les adaptations au début de la session et les adaptations au cours de la session. Dans ce chapitre, nous avons détaillé le premier type d'adaptation basé sur la description UED. Dans les chapitres suivants, nous allons nous intéresser aux adaptations *Cross-layer* exécutées durant la transmission des flux IPTV.

Chapitre 4

Les adaptations *Cross-layer* ascendantes exécutées au niveau applicatif

4.1 Introduction

Nous avons présenté dans le chapitre précédent le nouveau système de streaming XLAVS qui est embarqué sur des passerelles d'adaptation XLAG. Le XLAVS exécute des adaptations *Cross-layer* au début et au cours d'une session IPTV. Dans ce chapitre, nous nous intéressons aux adaptations exécutées au cours d'une session, c'est-à-dire, durant la transmission des flux IPTV.

Les deux contributions détaillées ici reposent sur l'approche *Cross-layer* ascendante (voir section 2.3.3) qui permet aux couches supérieures de s'adapter dynamiquement aux variations des couches inférieures. Effectivement, les conditions de transmission dans les réseaux 802.11 peuvent varier suivant plusieurs paramètres. Il est important que le XLAVS ait connaissance de ces paramètres afin d'adapter les mécanismes de QoS au niveau applicatif.

Dans la première contribution, nous proposons une adaptation dynamique du débit vidéo en fonction de la variation du débit physique dans les réseaux 802.11. Pour cela, nous commençons par présenter les algorithmes de contrôle du débit physique 802.11 qui sont situés au niveau de la couche MAC 802.11. Nous continuons par présenter les différentes techniques qui permettent d'adapter le débit vidéo dynamiquement durant la transmission. Ensuite, nous détaillons la mise en œuvre de cette contribution au niveau du XLAG (Cross Layer Adaptation Gateway) en proposant un modèle mathématique pour estimer le débit physique réellement disponible. Enfin, nous présentons les résultats des expérimentations effectuées afin d'évaluer cette première contribution.

La deuxième contribution permet l'adaptation du taux de redondance FEC (Forward Error Correction) et du débit vidéo en fonction de la puissance du signal récupérée au niveau physique et des taux de perte applicatifs. Au début, nous détaillons le mécanisme FEC en expliquant l'importance de la FEC au niveau paquets, comparée à la FEC au niveau bits. Ensuite, nous détaillons notre contribution qui se base sur une adaptation conjointe de la FEC et du débit vidéo

pour rendre le flux vidéo plus résistant aux pertes tout en maintenant un débit fixe pour le flux. Nous terminons par une présentation des tests conduits pour évaluer cette contribution ainsi que les résultats obtenus.

4.2 L'adaptation du débit vidéo en fonction du débit physique 802.11

4.2.1 Motivation

Actuellement, les réseaux 802.11 a/b/g peuvent fonctionner avec différents débits physiques variant entre 1 Mbits/s et 54 Mbits/s. Cette originalité, appelée « multi-rate capability », a été présentée dans la section 2.2.1.3. Elle est due à la définition de plusieurs modulations et taux de codage au niveau physique. La combinaison de ces modulations et taux de codage offre un débit précis. Toutefois, la zone de couverture du réseau est inversement proportionnelle au débit utilisé, c'est-à-dire que, lorsque le débit est élevé la zone de couverture de ce débit est réduite et inversement pour les débits faibles. Donc, un débit élevé peut être utilisé entre deux stations relativement proches, mais lorsque les stations s'éloignent, il est nécessaire d'utiliser des débits plus faibles pour faire face à l'évanouissement lent du signal et maintenir la communication.

Les standards IEEE 802.11 se sont limités à la définition des débits physiques sans spécifier les algorithmes qui permettent de basculer d'un débit vers un autre. Durant ces dernières années, nous avons assisté à l'apparition de plusieurs algorithmes qui prennent en charge cette fonctionnalité. La variation du débit physique n'est pas sans conséquences sur les applications multimédia qui exigent un débit relativement stable pour transmettre leurs flux.

Dans le cadre de notre architecture, le XLAG intègre un point d'accès qui offre une connexion WiFi à des stations mobiles. Ces stations peuvent se trouver à différentes distances du XLAG. Ce dernier utilise un débit physique spécifique pour chaque station afin de lui transmettre son flux IPTV. Ce débit physique varie suivant un algorithme de contrôle du débit mis en œuvre par le point d'accès intégré au XLAG. À partir de là, nous avons identifié deux problématiques importantes :

- Lorsque la distance entre le XLAG et la station est importante, le débit physique utilisé par le XLAG pour transmettre le flux IPTV peut être réduit jusqu'à 1 Mbits/s. Si la station reçoit initialement un flux IPTV (audio/vidéo) de 2.7 Mbits/s, le débit physique ne peut pas satisfaire le débit du flux et un pourcentage considérable de ce flux ne peut être transmis.
- Dans le cas où le XLAG transmet plusieurs flux IPTV à des stations en utilisant un débit physique différent pour chaque station, le débit physique réellement disponible pour une station diffère considérablement du débit physique utilisé par le XLAG pour transmettre à cette station le flux IPTV.

Ces deux problématiques introduisent de nouvelles exigences pour maintenir la QoS des flux IPTV. Premièrement, la nécessité d'informer la couche applicative sur la variation du débit physique des stations. Deuxièmement, le besoin de déterminer dynamiquement le débit physique effectif de chaque station afin d'adapter efficacement les flux qui leur sont destinés. Enfin, la nécessité de mettre en œuvre un mécanisme qui permet d'adapter le débit vidéo en temps réel.

Dans cette contribution, nous proposons une solution à chaque besoin, cependant avant de détailler ces solutions, nous présentons dans la section suivante, les algorithmes de contrôle du débit physique et leur mode de fonctionnement.

4.2.2 Les algorithmes de contrôle du débit physique 802.11

L'objectif des algorithmes de contrôle du débit physique, appelés RCA (Rate Control Algorithm), est de déterminer le débit physique adéquat pour transmettre un signal entre deux stations communicantes. En d'autres termes, ils déterminent la meilleure modulation et taux de codage qui permettent la transmission d'un signal avec un minimum d'altération de ce dernier.

Dans les réseaux 802.11, chaque trame peut être transmise avec un débit physique renseigné dans le champ « SIGNAL » (8 bits) de l'en-tête PLCP (voir l'annexe A.1). Ceci permet au RCA de l'émetteur de changer le débit physique pour la transmission de chaque trame sans aucune signalisation supplémentaire.

Puisque le standard IEEE 802.11 n'a pas spécifié un RCA précis, les premiers RCAs développés ont été élaborés par les constructeurs de carte WiFi afin de permettre aux cartes 802.11 d'adapter dynamiquement le débit physique sans intervention de l'utilisateur. Cependant, l'utilisateur peut fixer ce débit suivant sa volonté. Certains constructeurs ont implémenté le RCA directement au niveau du chipset des cartes et d'autres l'ont intégré comme un module aux pilotes des cartes. D'autres RCAs ont été proposés dans le cadre de différents travaux de recherche. La majorité de ces travaux se sont limités aux simulations pour valider le fonctionnement des RCAs proposés.

La différence entre les RCAs réside dans les informations utilisées pour juger de l'état du canal et aussi dans la politique adoptée pour faire varier le débit. Dans la Figure 4-1, nous représentons une taxonomie des RCAs qui ont été définis jusqu'à présent [200]. Nous présentons, ci-dessous, une description de chaque classe.

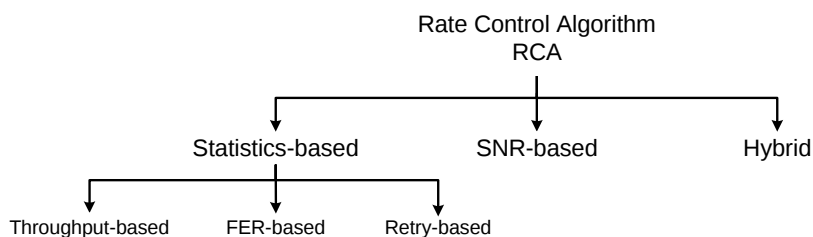


Figure 4-1 : Taxonomie des RCAs

4.2.2.1 Les RCAs basés sur des statistiques (Statistics-based)

Pour déterminer la qualité du canal de transmission, des statistiques sur le déroulement des transmissions passées peuvent être calculées au niveau MAC de l'émetteur. Ceci est facilement réalisable puisque chaque trame transmise au niveau MAC est acquittée par le récepteur. Donc l'émetteur, peut savoir si une transmission a réussi ou a échoué.

Il existe trois types de RCAs basés sur des statistiques. Les types se distinguent suivant le paramètre de performance considéré : Le débit d'émission effectif (throughput-based), le taux de perte des trames (FER-based : Frame Error Rate), le nombre de retransmission (retry-based). Ces trois paramètres sont en relation directe avec le débit effectif reçu par une station. La variation du débit d'émission permet d'optimiser ces paramètres, ce qui améliore indirectement le débit de réception. Le fonctionnement de chaque type d'algorithme est détaillé ci-dessous :

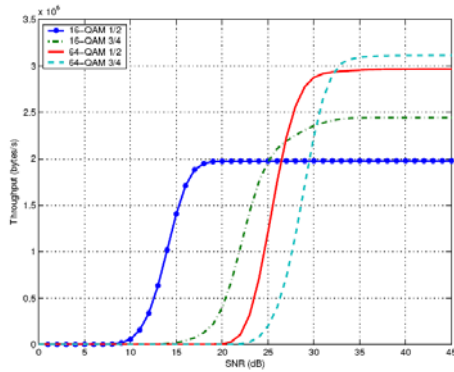
- **Le RCA basé sur le débit d'émission effectif :** Dans cette approche une proportion des données (10% : 1 trame sur 10) est transmise avec deux débits physiques différents, le débit supérieur adjacent et le débit inférieur adjacent, durant un intervalle de temps précis. Par exemple les deux débits adjacents à 18 Mbits/s sont 12 Mbits/s et 24 Mbits/s (voir le Tableau 2-1). L'intervalle de temps doit être relativement long pour avoir des statistiques fiables. Ensuite, les débits effectifs réalisés par chaque débit physique (nombre de trames/intervalle de temps) sont comparés. Le RCA bascule vers le débit physique qui réalise la meilleure performance. Cet algorithme est utilisé par la carte Atheros 802.11a basé sur le chipset AR5000.
- **Le RCA basé sur le FER :** Cette approche considère le taux de perte des trames au niveau MAC. Son mode de fonctionnement repose sur une adaptation basique qui consiste à basculer vers un débit bas lorsque le taux de perte dépasse un certain seuil et à basculer vers un débit élevé lorsque le taux de perte reste nul pendant un intervalle de temps prédéfini. Cet intervalle de temps ainsi que le seuil des taux de perte possèdent un impact sur les performances de ce type de RCA. La configuration optimale de ces deux paramètres doit être déterminée en fonction de l'état du canal et du type de flux transmis. Toutefois, dans la majorité des cas, ils sont initialisés avec des valeurs fixes.
- **Le RCA basé sur la retransmission :** Ce type d'algorithme se base sur le nombre de retransmissions des trames MAC. Lorsque ce nombre dépasse un seuil, la trame est transmise avec un débit faible. L'avantage de cette approche est la prise en charge des variations rapides de l'état du canal. Par contre, l'inconvénient majeur est l'oscillation continue du débit physique. Le RCA ARF (Auto Rate Fallback) [201] publié et implémenté dans la carte waveLan II se base sur ce type d'algorithme. Pour stabiliser le changement de débit dans ARF, une amélioration de ce dernier, appelée AARF (Adaptive ARF) est proposée par les auteurs dans [202].

4.2.2.2 Les RCAs basés sur le SNR (SNR-based)

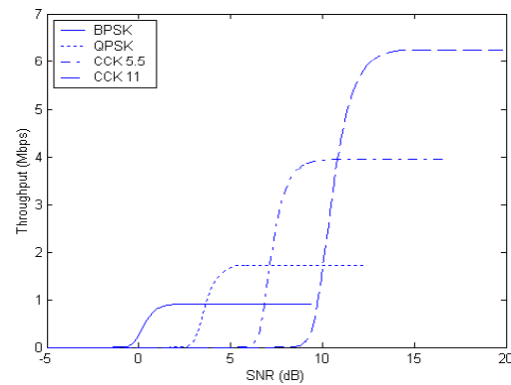
Nous avons expliqué dans la section 2.2.1.3 que chaque modulation et taux de codage nécessite un certain niveau de SNR pour que les symboles du signal soient décodés correctement au niveau du récepteur. Les RCAs basés sur SNR exploitent cette relation pour déterminer le meilleur débit physique en fonction du SNR perçu par le récepteur. Les deux graphes de la Figure 4-2 ont été proposés respectivement par les travaux [202] et [203]. Ils donnent respectivement le débit physique théorique en fonction du SNR pour différentes modulations du standard IEEE 802.11a et IEEE 802.11b.

Cependant, pour utiliser ce type de RCA sur une plateforme réel, il faut répondre à plusieurs défis techniques. Le premier défi consiste à trouver une correspondance réaliste entre le SNR et le débit physique parce que les résultats présentés dans la Figure 4-2 sont des résultats théoriques qui

se basent sur une modélisation simpliste du canal de transmission. Le second défi est l'estimation du SNR en temps réel au niveau du récepteur puisque la seule information disponible au niveau du récepteur est le SSI (Signal Strength Indicator) qui informe sur la puissance du signal avec laquelle les trames sont reçues. Le dernier défi est de rendre disponible le SNR perçu par le récepteur au niveau de l'émetteur puisque le RCA fonctionne à ce niveau.



IEEE 802.11a



IEEE 802.11b

Figure 4-2 : Le débit physique en fonction du SNR pour le standard IEEE 802.11a et IEEE 802.11b

Pour surmonter ces défis, les auteurs dans [203] proposent d'utiliser le SSI en remplacement du SNR. De plus, au lieu d'utiliser le SSI des trames reçues au niveau du récepteur, l'algorithme utilise le SSI des acquittements reçus au niveau de l'émetteur. Les auteurs justifient cette solution par le fait que le canal de transmission est symétrique, donc, le SNR du canal est identique sur n'importe quel point. Le mapping entre le SSI et les débits est dynamique et adapté continuellement durant la transmission pour faire face aux interférences.

Le RBAR (Received –based Auto Rate) proposé dans [204] appartient à cette classe de RCA. Le choix du débit physique dans RBAR se fait côté récepteur en se basant sur différents paramètres. Les auteurs n'ont pas proposé une nouvelle technique pour le choix du débit physique mais ils se sont focalisés sur la définition d'un protocole basé sur les trames RTS/CTS pour transmettre à l'émetteur le débit choisi par le récepteur. Le RBAR a été évalué par des simulations en utilisant le SNR comme métrique pour le changement de débit physique.

4.2.2.3 Les RCAs hybrides (Hybrid)

Les deux classes précédentes possèdent des avantages et des inconvénients. Les RCAs basés sur des statistiques permettent une adaptation du débit à long terme mais leur réaction est limitée aux changements rapides de l'état du canal. Par contre, les RCAs basés sur le SNR réagissent rapidement aux changements de l'état du canal, mais ils manquent de fiabilité à cause de l'estimation approximative du SNR.

Pour trouver un compromis entre les deux précédentes classes en maximisant leurs avantages et en minimisant leurs inconvénients, les auteurs dans [200] proposent un RCA hybride qui exploite à la fois les statistiques de transmission et le SNR. Le RCA utilise à la base un algorithme statistique qui calcule le débit d'émission effectif et ajuste le débit physique en conséquence. Cependant, le

débit physique peut être changé par un deuxième algorithme qui se base sur le SNR afin de réagir au changement rapide de l'état du canal.

4.2.3 L'adaptation dynamique du débit vidéo au cours de la transmission

L'une des fonctionnalités principales du nouveau système de transmission XLAVS (Cross Layer Adaptive Video Streaming) est l'adaptation des flux vidéo grâce à son module « transcoder ». Cette adaptation est exécutée au début de la session et les paramètres de transcodage sont maintenus durant toute la vie d'une session (voir section 3.4.4.5).

Cependant, les conditions de transmission dans un réseau 802.11, principalement le débit physique, peuvent varier durant la transmission des flux. Par conséquent, il est important d'assurer une adaptation dynamique du débit des flux au cours de la session et non seulement au début. Puisque le débit vidéo est nettement plus élevé que le débit audio, nous nous sommes focalisés sur l'adaptation dynamique du débit vidéo et particulièrement les flux vidéo MPEG.

4.2.3.1 Le codage vidéo MPEG

Le codage vidéo MPEG se base sur deux caractéristiques principales du signal vidéo numérique, à savoir la redondance temporelle et la redondance spatiale. La première redondance est due principalement à la similitude qui existe entre les images vidéo successives. Pour réduire cette redondance, le codage temporel est appliqué pour coder la différence qui existe entre les images. Ainsi, des images de prédiction sont générées à partir d'une image de référence. Pour augmenter les performances de la prédiction, le codage temporel utilise des estimations de mouvement entre les images. Cette estimation, donnée sous formes de vecteurs de mouvement, permet au décodeur de compenser le mouvement des blocs de pixels entre l'image prédite et l'image de référence. Le codage temporel génère trois types d'image représentés dans la Figure 4-3:

- L'image I (Intra-coded): Elle représente l'image de référence et elle ne subit aucun codage temporel.
- L'image P (Predictive): Cette image est codée en utilisant une image précédente de type I ou P.
- L'image B (Bidirectional) : Cette image est codée en utilisant une image précédente (I ou P) et une image successive (P).

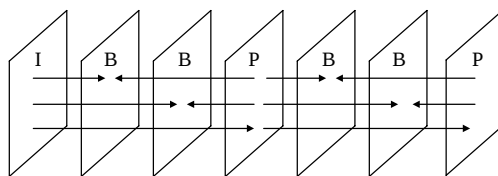


Figure 4-3 : Les types d'images dans le codage MPEG

Suite au codage temporel, chaque image subit un codage spatial qui permet de réduire la similitude entre les blocs de pixel. La Figure 4-4 illustre les principales procédures d'un codage spatial.

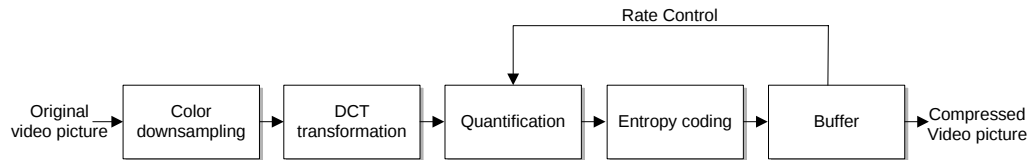


Figure 4-4 : Le codage spatial dans la norme MPEG

Au début, l'image subit un sous-échantillonnage des couleurs qui permet de réduire le nombre de bits par pixel tout en donnant plus d'importance à la luminance (Y) par rapport à la chrominance (Cb et Cr), puisque la vision humaine est plus sensible à la composante luminance (Y). Ensuite, une transformée en cosinus discrète DCT (Discrete Cosine Transform) est appliquée à l'image pour transformer les données d'un domaine spatio-temporel (pixels) vers un domaine fréquentiel (fréquences). Cette transformée permet une meilleure représentation de l'information afin qu'elle puisse être compressée. La matrice DCT générée subit un dernier traitement avant la compression réel. Ce traitement est la quantification des données qui divise la matrice de DCT par une matrice de quantification afin de rendre nulle les fréquences élevées imperceptibles pour la vision humaine. La matrice de quantification est calculée à partir d'une équation en utilisant un facteur de quantification, appelé aussi facteur de qualité. L'équation Eq. 4-1 illustre un exemple de ce calcul.

$$Q(i, j) = 1 + (1 + i + j) \times Fq \quad \text{Eq. 4-1}$$

Enfin, un codage en entropie est appliqué pour compresser la matrice finale en tenant compte des propriétés statistiques de ses données. Les images compressées sont stockées dans un tampon qui a pour objectif de réguler le débit vidéo. Il existe deux types de codage pour le débit vidéo :

- Le codage à débit variable VBR (Variable Bit Rate) : Avec ce codage, aucun contrôle de débit n'est effectué. Pour cela, le facteur de quantification est fixe et le débit vidéo est variable.
- Le codage à débit constant CBR (Constant Bit Rate) : Avec ce codage, le débit vidéo généré par le codec est constant. Pour cela, un algorithme de contrôle de débit est utilisé entre le tampon de stockage et la procédure de quantification. Cet algorithme utilise une mémoire tampon pour estimer le débit vidéo compressé et il fait varier le facteur de quantification pour maintenir un débit stable.

4.2.3.2 Techniques d'adaptation dynamique du débit vidéo

Dans le nouveau système de transmission XLAVS, deux techniques peuvent être utilisées pour faire varier le débit vidéo dynamiquement : la variation de la résolution temporelle et la variation de la résolution SNR.

4.2.3.2.1 Variation de la résolution temporelle

La résolution temporelle d'une séquence vidéo représente le nombre d'images par seconde. La variation de cette résolution correspond à la suppression d'images en suivant la hiérarchie temporelle, c'est-à-dire, le type d'image généré par le codage temporel. Comme le décodage d'une image B dépend du décodage d'une image P qui, à son tour, dépend d'une image I ou d'une image

P précédente, la suppression des images commence inévitablement par les images B suivies des images P.

En variant le nombre d'images par seconde suivant la hiérarchie temporelle, nous varions le débit de transmission d'une vidéo sans affecter la qualité des images transmises (pas d'artefacts dans les images durant la lecture de la vidéo). La Figure 4-5 montre la variation de ce débit pour deux flux vidéo, vidéo 1 et vidéo 2, possédant des débits initiaux différents.

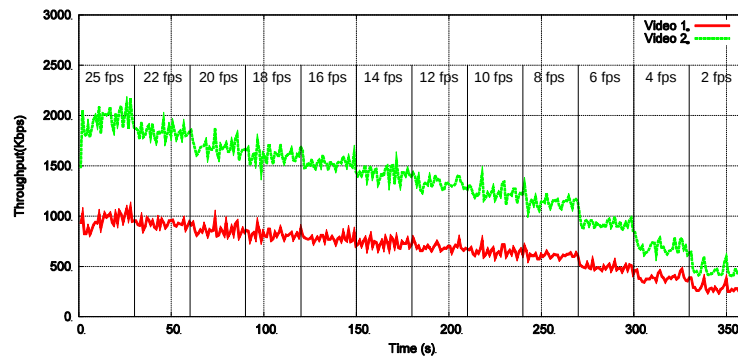


Figure 4-5 : La variation du débit vidéo suivant la variation de la résolution temporelle

Les caractéristiques techniques des deux flux sont résumées dans le Tableau 4-1.

	Vidéo 1	Vidéo 2
Format de Codage	MPEG4	
La taille d'un GOP	12	
Structure d'un GOP	IBBPBBPBBPBB	
Résolution	CIF : 352x288 pixels	
Nombre d'images par seconde	25 fps	
Débit moyen/max/min	970/1064/904 Kbits/s	1940/2032/1800 Kbits/s
Durée de la séquence	360 secondes	

Tableau 4-1 : Les caractéristiques techniques des flux vidéo

La Figure 4-6 donne le débit moyen réalisé par les deux flux vidéo pour chaque résolution temporelle. Nous constatons que le débit des deux vidéos diminue avec la diminution du nombre d'images par seconde, mais cette diminution est en relation directe avec le débit vidéo initial. De plus, la pente de la diminution est plus importante à partir de 8 images par seconde qui correspond au début de la suppression des images P dont les tailles sont plus grandes comparées aux tailles des images B.

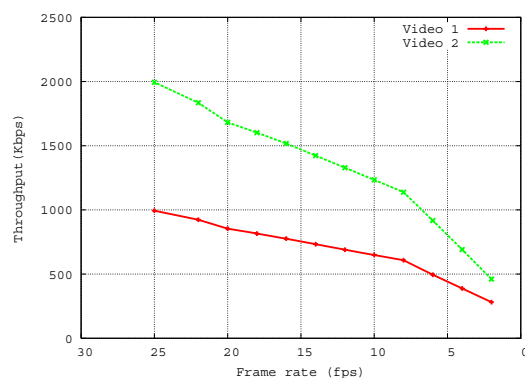


Figure 4-6 : Le débit vidéo moyen en fonction du nombre d'images par seconde

Le premier inconvénient de la variation de la résolution temporelle est la dégradation de la qualité vidéo. En effet, le manque d'image affecte principalement le mouvement des objets dans une scène. Le mouvement n'est plus continu mais il évolue en saccade. Cet effet est plus perceptible avec des scènes d'action, comparées à des scènes fixes. Le deuxième inconvénient est la difficulté de contrôler efficacement le débit puisqu'il dépend principalement du débit vidéo initial. De plus, pour réduire un débit vidéo initialement élevé, il faut réduire d'avantage le nombre d'images par seconde. Par exemple, pour avoir un débit de 600 Kbits/s, la vidéo 1 doit transmettre 8 images par secondes, par contre la vidéo 2 doit transmettre 4 images par secondes. Ainsi, cette méthode s'avère peu attractive vu la dégradation de la qualité et la complication liée à la prédiction du débit.

4.2.3.2.2 Variation de la résolution SNR

La deuxième technique qui permet de faire varier dynamiquement le débit vidéo se base sur la résolution SNR. Cette dernière correspond au degré de finesse des images vidéo définies par les fréquences élevées d'une matrice DCT. Ces fréquences sont réduites durant le codage spatial par la procédure de quantification en utilisant le facteur de quantification. Lorsque ce dernier est élevé, le nombre de fréquences nulles dans la matrice DCT augmente. Ceci réduit la finesse de l'image, mais réduit aussi le débit de la vidéo. Par contre, lorsque le facteur de quantification est faible, la matrice DCT contiendra plus de fréquence qui font augmenter la finesse des images, et par la même, le débit vidéo.

L'algorithme de contrôle de débit utilisé par CBR se base sur le facteur de quantification pour maintenir un débit vidéo stable. Donc, pour contrôler efficacement le débit vidéo, il suffit d'agir sur cet algorithme en le réinitialisant à chaque fois avec le débit voulu. Nous rappelons que dans notre architecture, le flux vidéo est transcodé en temps réel, ce qui rend possible cette variation dynamique du débit vidéo.

La Figure 4-7 illustre parfaitement cette variation pour les deux flux vidéo précédents (voir Tableau 4-1). Les deux flux sont transcodés dans le même format de codage MPEG-4 avec la même résolution d'image. Le débit des deux flux est réduit suivant le débit souhaité dans chaque intervalle de temps. Ceci démontre le contrôle parfait du débit réalisé en variant la résolution SNR. Ce contrôle est complètement indépendant du débit vidéo initial. Le temps de réaction de l'algorithme est de l'ordre de 2 secondes et la précision de l'adaptation est de l'ordre de 50 Kbits/s.

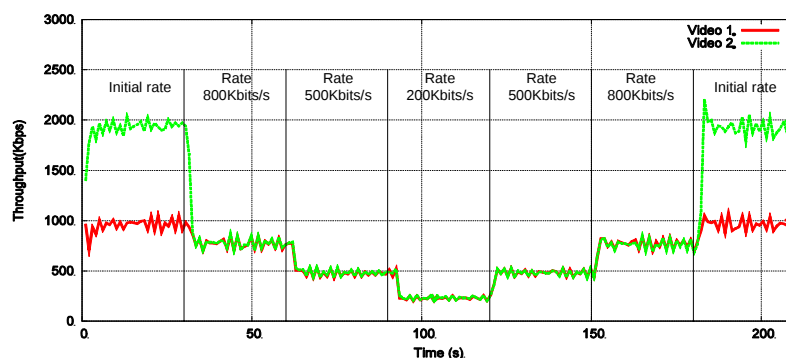


Figure 4-7 : La variation du débit vidéo en fonction de la résolution SNR

Il est clair qu'en réduisant le débit, nous réduisons la qualité de la vidéo, mais la séquence vidéo reste fluide. La Figure 4-8 montre des captures de la même image vidéo avec différents débits : 2 Mbits/s, 1 Mbits/s, 800 Kbits/s, 500 Kbits/s, 200 Kbits/s. Nous remarquons difficilement la différence entre les images et la qualité globale reste acceptable. De plus, les mouvements des objets dans les scènes ne sont pas saccadés.

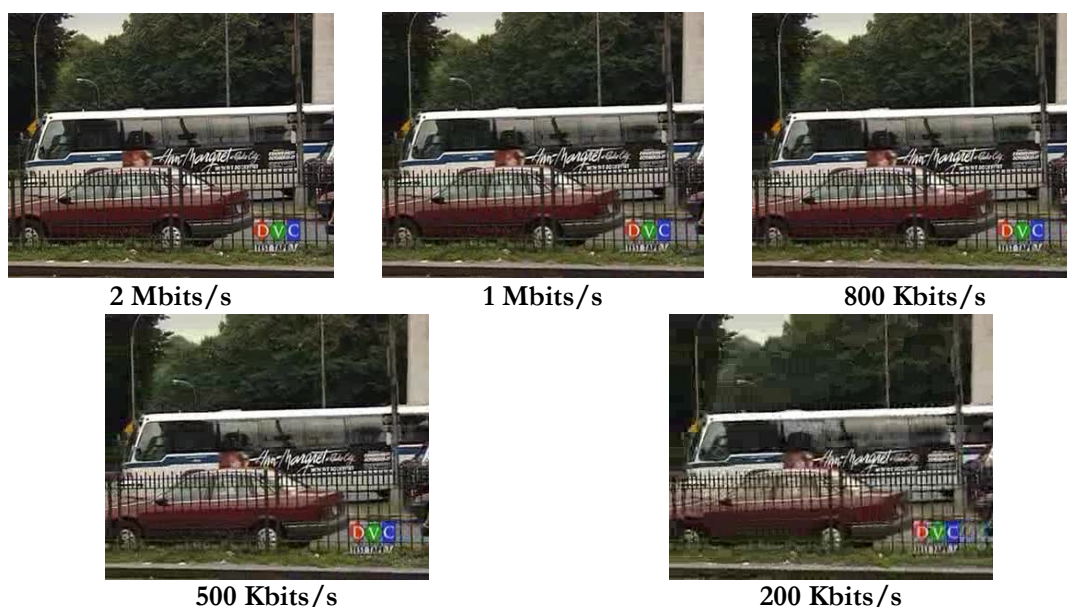


Figure 4-8 : Captures d'images vidéo codées avec différents débits

4.2.4 L'adaptation du débit vidéo en fonction du débit Physique 802.11

L'objectif de cette adaptation est d'assurer une corrélation entre le codage source et le codage canal. Le codage source correspond au débit vidéo généré par le XLAVS et le codage canal correspond au débit physique utilisé par le XLAG. La correspondance entre ces deux débits est primordiale pour optimiser la QoS des flux vidéo tout en assurant une exploitation optimale du débit physique disponible.

Pour mettre en œuvre cette correspondance, le XLAVS doit se tenir informé du débit disponible au niveau physique. Ce débit est choisi par un algorithme de contrôle de débit (RCA) implémenté au niveau de la couche MAC 802.11. Une instance du RCA est associée pour chaque

station connectée au XLAG et le fonctionnement des instances est indépendant entre elles. Par conséquent, le XLAG peut utiliser un débit physique différent pour chaque station.

4.2.4.1 L'estimation du débit applicatif à partir du débit physique

L'un des défis majeurs que doit effectuer le XLAVS à travers son module « XLDP » pour concrétiser cette adaptation est l'estimation du débit applicatif qui peut être réalisé avec un débit physique. En effet, ces deux débits ne sont pas identiques malgré l'existence d'une certaine corrélation.

Afin d'expliquer la différence entre le débit applicatif et le débit physique, nous supposons un réseau 802.11 en mode infrastructure et nous étudions deux situations possibles : (1) une seule station connectée au point d'accès et (2) plusieurs stations connectées de sorte que le point d'accès utilise un débit physique différent pour chaque station. Nous détaillons ces deux cas par la suite :

- **Situation 1 - Une station** : Lorsque le débit physique est fixé à un certain seuil au niveau du point d'accès (AP), le débit applicatif, entre le point d'accès et la station qui peut être réalisé avec ce débit physique, est toujours inférieur. Effectivement, le débit physique est un débit théorique calculé à partir des modulations et des codages utilisés. Dans la réalité, ce débit n'est jamais atteint à cause d'un délai supplémentaire *overhead* introduit pour la transmission de chaque trame de données. Ce délai est causé principalement par : (1) les en-têtes des différentes couches réseaux (RTP, UDP, IP, LLC, 802.11 MAC et 802.11 PHY), (2) les intervalles de temps utilisés par la couche 802.11 MAC pour partager l'accès au canal et (3) l'attente de l'acquittement et la retransmission au niveau MAC. Le Tableau 4-2 donne le débit applicatif moyen réalisé avec les différents débits physiques. Ces mesures ont été obtenues en transmettant un flux UDP d'un point d'accès vers une station. La taille des trames est fixée à 1512 octets.

Débit physique	Débit applicatif	Ratio débit APP/débit PHY
54 Mbits/s	36 Mbits/s	66 %
48 Mbits/s	32.8 Mbits/s	68 %
36Mbits/s	26.3 Mbits/s	73 %
24 Mbits/s	18.7 Mbits/s	78 %
18 Mbits/s	14.5 Mbits/s	80 %
12 Mbits/s	10 Mbits/s	83 %
11 Mbits/s	8.15 Mbits/s	74 %
9 Mbits/s	7.7 Mbits/s	85 %
6 Mbits/s	5.2 Mbits/s	86 %
5.5 Mbits/s	4.4 Mbits/s	80 %
2 Mbits/s	1.7 Mbits/s	85%
1 Mbits/s	900 Kbits/s	87 %

Tableau 4-2 : Les débits applicatifs réalisés avec différents débits physiques IEEE 802.11

- **Situation 2 - plusieurs stations** : Dans cette situation, l'AP peut utiliser un débit physique pour chaque station. Le débit physique est choisi par l'instance de RCA associée à chacune des stations. Si l'AP transmet simultanément des flux pour toutes les stations, le débit applicatif d'un flux transmis ne dépend pas seulement du débit physique utilisé par l'AP pour transmettre ce flux (Situation 1), mais il dépend aussi des débits physiques utilisés pour transmettre les autres flux vers les autres stations. Cette dépendance s'explique par deux raisons principales. La première raison est le partage équitable du canal entre toutes les stations, c'est-à-dire, la même quantité de données est envoyée pour chaque station durant un intervalle de temps t . La deuxième raison est liée au débit physique lui-même puisque la transmission avec un débit faible occupe le canal pour une durée plus longue. Par exemple, le temps nécessaire pour transmettre une trame de 1512 octets avec un débit de 54 Mbits/s est de 0.0267 ms, mais avec un débit de 1Mbits/s, le temps augmente à 1.44 ms. Ainsi, le débit physique le moins élevé tire vers le bas les débits applicatifs de tous les flux. Pour illustrer ce phénomène, le Tableau 4-3 donne le débit applicatif moyen de trois flux transmis d'un AP vers les trois stations en utilisant des débits physiques différents. Le tableau présente toutes les combinaisons possibles entre trois débits physiques : 1 Mbits/s, 11 Mbits/s et 54 Mbits/s. Nous constatons que les débits applicatifs des flux sont relativement égaux malgré la différence de débit physique. Ceci confirme le partage équitable du canal. De plus, il n'existe plus de corrélation entre le débit applicatif et le débit physique pour une station. Par exemple, dans la troisième combinaison, la station 3 possède un débit physique de 54 Mbits/s, mais son débit applicatif moyen est de 417 Kbits/s.

Débits Physiques			Débits Applicatifs		
Station 1 Mbits/s	Station 2 Mbits/s	Station 3 Mbits/s	Station 1	Station 2	Station 3
1	1	1	289 Kbits/s	303 Kbits/s	310 Kbits/s
1	1	11	416 Kbits/s	411 Kbits/s	446 Kbits/s
1	1	54	423 Kbits/s	425 Kbits/s	417 Kbits/s
1	54	54	858 Kbits/s	828 Kbits/s	834 Kbits/s
1	11	54	808 Kbits/s	804 Kbits/s	805 Kbits/s
1	11	11	732 Kbits/s	730 Kbits/s	737 Kbits/s
11	11	11	2.70 Mbits/s	2.74 Mbits/s	2.69 Mbits/s
11	54	11	3.62 Mbits/s	3.65 Mbits/s	3.75 Mbits/s
11	54	54	5.54 Mbits/s	5.82 Mbits/s	5.79 Mbits/s
54	54	54	12.10 Mbits/s	12.04 Mbits/s	12.11 Mbits/s

Tableau 4-3 : Les débits applicatifs réalisés avec différentes combinaisons de débits physiques

À partir de cette étude, nous avons proposé un modèle mathématique pour calculer le débit physique effectif pour chaque station au niveau d'un AP en utilisant les débits physiques théoriques de toutes les stations connectées. Ce modèle se base sur le fait que le canal est partagé équitablement entre les stations. Par conséquent, le nombre de trames transmises de l'AP vers chaque station est identique durant une seconde. Nous supposons aussi que la taille des trames est identique pour toutes les stations. Tout cela se traduit par l'équation Eq. 4-2.

$$\frac{Nbf \cdot fs}{r_1} + \frac{Nbf \cdot fs}{r_2} + \frac{Nbf \cdot fs}{r_3} + \dots + \frac{Nbf \cdot fs}{r_i} + \dots + \frac{Nbf \cdot fs}{r_{Nb_{sta}}} = 1s \quad \text{Eq. 4-2}$$

Nbf : Le nombre de trames transmises

fs : La taille des trames

r_i : Le débit physique théorique utilisé pour transmettre les trames à la station i

Le débit physique effectif réalisé par chaque station Ri_{phy} est calculé par l'équation Eq. 4-3

$$Ri_{phy} = Nbf \cdot fs \quad \text{Eq. 4-3}$$

Le Nbf dans l'équation Eq. 4-4 est calculé à partir de l'équation Eq. 4-2

$$Nbf = \frac{1}{fs \sum_{i=1}^{Nb_{sta}} \frac{1}{r_i}} \quad \text{Eq. 4-4}$$

En remplaçant l'expression du Nbf dans l'équation Eq. 4-3, le calcul de Ri_{phy} peut être simplifié en utilisant l'équation eq. 4-5

$$Ri_{phy} = \frac{1}{\sum_{i=1}^{Nb_{sta}} \frac{1}{r_i}} \quad \text{Eq. 4-5}$$

Il est important de préciser que le débit physique effectif est supérieur au débit applicatif puisqu'il ne considère pas les délais supplémentaires introduits par les en-têtes des couches réseaux et le mécanisme d'accès 802.11 MAC. Cependant, il permet une meilleure estimation du débit applicatif avec un taux d'écart acceptable comparé au débit physique théorique. Le Tableau 4-4 présente le débit effectif calculé à l'aide de l'équation Eq. 4-5 et la moyenne des débits applicatifs des stations 1, 2 et 3 pour toutes les combinaisons du Tableau 4-3. En comparant les deux dernières colonnes du Tableau 4-4, nous constatons clairement que le débit physique effectif offre une meilleure estimation du débit applicatif.

Débits Physiques			Débits Physiques effectifs calculés par l'équation Eq. 4-5	Débits applicatifs moyens pour les stations 1, 2 et 3
Station 1 Mbits/s	Station 2 Mbits/s	Station 3 Mbits/s		
1	1	1	341 Kbits/s	300 Kbits/s
1	1	11	489 Kbits/s	424 Kbits/s
1	1	54	507 Kbits/s	421 Kbits/s
1	54	54	988 Kbits/s	840 Kbits/s
1	11	54	924 Kbits/s	805 Kbits/s
1	11	11	867 Kbits/s	733 Kbits/s
11	11	11	3.7 Mbits/s	2.71 Mbits/s
11	54	11	5.05 Mbits/s	3.67 Mbits/s
11	54	54	7.93 Mbits/s	5.71 Mbits/s
54	54	54	18 Mbits/s	12 Mbits/s

Tableau 4-4 : Les débits physiques effectifs calculés pour différentes combinaisons de débits physiques

4.2.4.2 Mise en œuvre de l'adaptation du débit au niveau du XLAG

Au niveau du XLAG, l'adaptation du débit vidéo est exécutée par le module « transcoder » en variant la résolution SNR comme nous l'avons expliqué dans la section 4.2.3.2.2. Toutefois, le débit vidéo est décidé par le module « XLDP » qui doit récupérer le débit physique de chaque station connectée au XLAG.

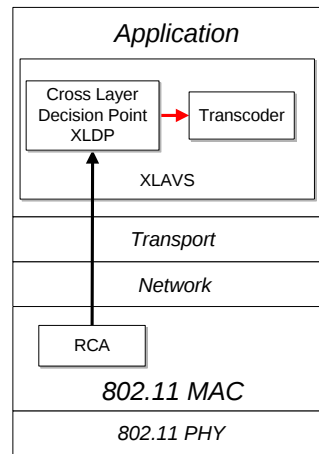


Figure 4-9 : L'interaction entre les modules au niveau du XLAG pour l'adaptation du débit vidéo en fonction du débit physique

La Figure 4-9 illustre les interactions entre les différents modules qui mettent en œuvre l'adaptation dynamique du débit vidéo. Périodiquement, le module « XLDP » interroge le module RCA qui gère les débits physiques au niveau de la couche MAC 802.11. Cette interrogation permet de récupérer deux informations importantes : (1) le nombre de stations connectées et (2) le débit physique de chaque station. Elle est effectuée par un appel de fonction « *liste_station (sta_nb, phy_rate[])* » interne au XLAG.

Par la suite, le module « XLDP » calcule le débit physique effectif $R_{i_{phy}}$ en utilisant l'équation Eq. 4-5 proposée dans la section précédente. À partir de ce débit, le module « XLDP » calcule le débit vidéo de chaque session IPTV en permettant au flux vidéo d'occuper une proportion du débit physique effectif. Cette proportion doit prendre en considération la différence entre le débit physique effectif et le débit applicatif. De plus, elle doit empêcher le flux vidéo d'occuper toute la bande passante, par exemple, le débit vidéo occupe 50% du débit physique effectif. La valeur de la proportion fait partie de la politique d'adaptation qui peut être configurée par un administrateur externe.

Enfin, le débit vidéo est communiqué au module « transcoder » qui réinitialise l'algorithme de contrôle du débit pour le flux vidéo avec la nouvelle valeur reçue. La périodicité de cette adaptation doit considérer le temps de réaction de l'adaptation vidéo, c'est-à-dire, le temps entre l'initialisation du module « transcoder » et la diminution effective du débit vidéo. Ce temps est de l'ordre de 2 secondes (voir section 4.2.3.2.2). Il est donc inutile d'exécuter l'adaptation avec une périodicité inférieure à 2 secondes. Ceci n'est pas contraignant puisque la diminution et l'augmentation du débit physique au niveau des RCAs n'est pas drastique, mais s'effectue étape par étape (voir section 4.2.2).

4.2.5 Tests et évaluations de performance

L'adaptation dynamique du débit vidéo en fonction du débit physique a été implémentée au niveau de la passerelle d'adaptation XLAG. Le module « transcoder » qui se base sur la librairie FFMPEG [198] a été modifié pour permettre une réinitialisation du débit vidéo en temps réel durant la transmission du flux. Au niveau MAC, le XLAG utilise le pilote libre *MadWiFi* [197] pour les cartes 802.11 basées sur un chipset *Atheros* [196]. Le pilote *MadWiFi* implémente trois RCA différents : *Onoe* [197], *Amrr* [202] et *Sample* [205]. Ces RCAs appartiennent à la classe basée sur des statistiques ; *Onoe* et *Amrr* se basent sur les retransmissions et *Sample* se base sur le débit d'émission effectif (voir section 4.2.2.1). Nous présentons ci-dessous une description de chaque RCA :

- **Onoe** : Ce RCA adapte le débit physique chaque seconde. Pour cela, il utilise un crédit qui est incrémenté ou décrémenté suivant la qualité des transmissions. Ainsi, le RCA débute avec un débit initial de 24 Mbits/s et un crédit égal à zéro. Ensuite le RCA exécute les étapes suivantes chaque seconde :
 - Si aucun paquet n'a pu être transmis, le RCA bascule vers le plus petit débit physique.
 - Si la transmission de 10 trames a nécessité une moyenne de retransmission supérieure à 1, le RCA bascule vers le débit inférieur adjacent.
 - Si plus de 10 % des trames transmises ont nécessité des retransmissions, le crédit est décrémenté jusqu'à atteindre zéro sinon il est incrémenté.
 - Si le crédit atteint 10, le RCA bascule vers le débit supérieur adjacent.

Avec ce modèle d'adaptation, la réaction de cet RCA est très limitée, particulièrement, pour l'augmentation du débit.

- **Amrr** : Ce RCA est une implémentation de l'algorithme AARF (voir section 4.2.2.1). Dans *Amrr*, l'adaptation du débit physique n'est pas périodique, mais elle dépend principalement du nombre d'échantillons recueillis afin d'avoir des statistiques représentatives. Donc, ce RCA adapte le débit physique uniquement lorsque le nombre d'échantillons dépasse un certain seuil. *Amrr* utilise aussi un crédit. L'adaptation suit les étapes suivantes :
 - Si moins de 10 % des trames transmises ont nécessité une retransmission, incrémenter le crédit.
 - Si plus de 33 % des trames transmises ont nécessité une retransmission, le RCA bascule vers le débit inférieur adjacent.
 - Si le crédit atteint 10, le RCA bascule vers le débit supérieur adjacent.

Le principe de fonctionnement d'*Amrr* ne diffère pas considérablement de celui d'*Onoe*, mais *Amrr* réagit plus rapidement puisque sa périodicité dépend du nombre d'échantillons.

- **Sample** : Ce RCA se base sur une estimation du temps de transmission. Sa périodicité est d'une seconde. Il démarre la transmission avec le débit physique le plus élevé. Si la transmission d'une trame échoue 4 fois, le RCA bascule vers le débit inférieur adjacent jusqu'à ce que la trame puisse être transmise correctement. Pour augmenter le débit, *Sample* transmet une trame sur dix

avec un débit physique choisi aléatoirement. Par la suite, il compare le temps de transmission moyen réalisé par chaque débit physique. Le RCA bascule vers le débit qui réalise le plus bas temps de transmission. Ce mode de fonctionnement provoque une variation drastique du débit physique qui n'évolue plus étape par étape.

Le pilote *MadWiFi* offre aussi une interface qui permet de récupérer le nombre de stations et le débit physique de chaque station. Cet appel de fonction a été intégré au niveau du module « XLDP ».

Les tests ont été organisés en deux parties. Dans la première partie, nous évaluons l'adaptation dynamique du débit vidéo en utilisant les trois RCAs décrits ci-dessus. Dans la deuxième partie, nous évaluons la capacité de l'adaptation dynamique à préserver la QoS des flux vidéo en utilisant plusieurs stations.

La Figure 4-10 illustre la plateforme d'expérimentation utilisée pour évaluer l'adaptation dynamique du débit vidéo en fonction du débit physique. La plateforme est constituée de la passerelle d'adaptation XLAG et de trois stations connectées au XLAG. Ce dernier embarque le point d'accès et le système de transmission XLAVS.

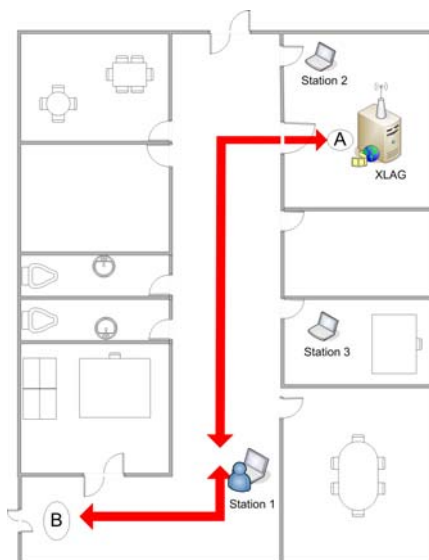


Figure 4-10 : La plateforme d'expérimentation

4.2.5.1 L'adaptation en utilisant différents RCAs

Dans cette partie des tests, nous utilisons uniquement le XLAG et la station 1 (voir la Figure 4-10). Le test consiste à transmettre un flux vidéo du XLAG à la station 1 et la durée totale d'un test est de 180s. Durant cette période, l'utilisateur de la station 1 se déplace du point A, à côté du XLAG, vers le point B avec une vitesse moyenne d'un pas par seconde. La durée de ce déplacement est de 60s. L'utilisateur reste au niveau du point B durant 60s et il revient vers le point A au bout de 60s. L'objectif de ce déplacement est de voir à la fois la réaction des RCAs au niveau MAC et la réaction de notre adaptation pour varier le débit vidéo. Pour cela, nous avons défini deux scénarios :

- **Scénario 1** : Ce scénario reproduit le comportement d'un système traditionnel. Le RCA adapte le débit physique mais le débit vidéo n'est pas adapté en conséquence.

- **Scénario 2 :** Le débit vidéo est adapté suivant la variation du débit physique effectif. Puisque nous utilisons une seule station, le débit physique effectif de cette station est identique à son débit physique théorique. La politique d'adaptation stipule que le flux vidéo peut occuper 50 % du débit physique effectif. Par exemple, si le débit physique effectif est de 2 Mbits/s, le débit vidéo peut utiliser 1Mbits/s.

Les caractéristiques techniques du flux vidéo utilisé dans les tests sont données dans le Tableau 4-1, avec un débit moyen de 2.425 Mbits/s, le débit maximum est de 2.584 Mbits/s et le débit minimum est de 2.256 Mbits/s. Pour chaque scénario, nous avons effectué dix tests pour chaque RCA. Dans ce qui suit, nous présentons un résultat représentatif qui permet de comparer les RCAs dans les deux scénarios.

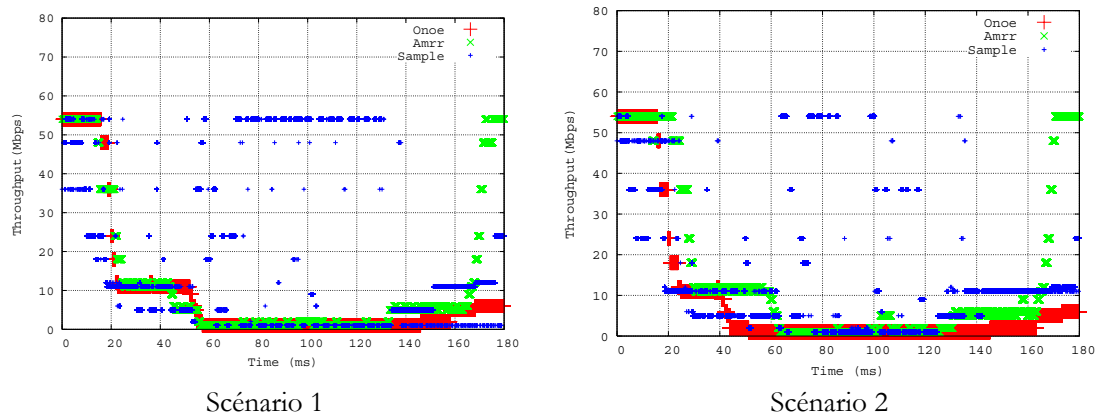
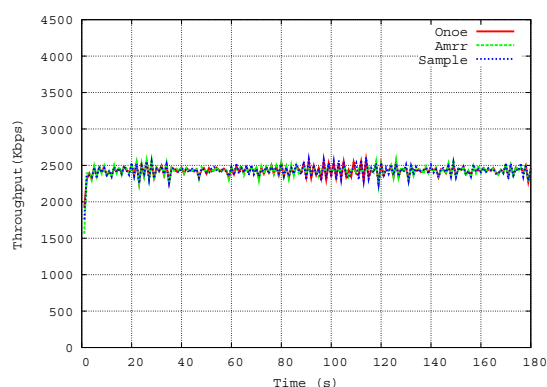


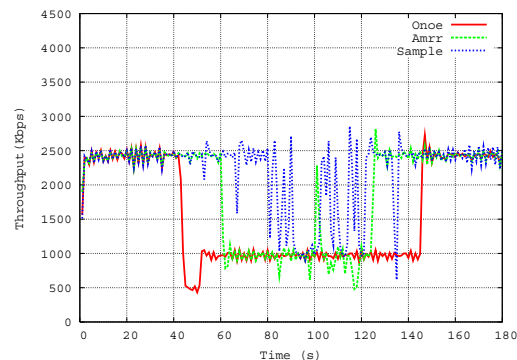
Figure 4-11 : la variation des débits physiques dans les scénarios 1 et 2

La Figure 4-11 montre la variation des débits physiques durant les tests pour les différents RCAs, *Onoe*, *Amrr* et *Sample* dans les deux scénarios. Nous constatons que la variation est différente pour les trois RCAs, mais elle est relativement identique pour un RCA dans les deux scénarios. Nous remarquons que la diminution du débit physique est la même pour tous les RCAs dans la phase de l'éloignement de la station 1 du point A et son immobilisme au niveau du point B, intervalle [0s, 120s]. Cependant, la variation du débit est instable avec *Sample* qui s'explique par son mode de fonctionnement détaillé dans la section précédente. Durant l'intervalle [120s, 180s], la station 1 entame son retour au point A. L'augmentation du débit physique est rapide pour *Amrr*, mais elle est très lente pour *Onoe* et *Sample*.

Pour voir l'impact de ces variations sur le débit physique. Les Figures 4-12, 4-13 et 4-14 donnent respectivement le débit d'émission, les taux de perte et le débit de réception du flux vidéo pour chaque RCA dans les deux scénarios.

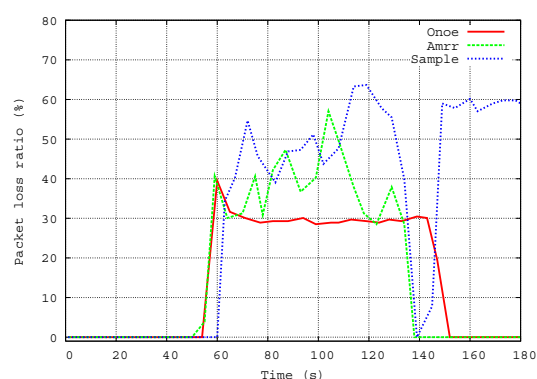


Scénario 1

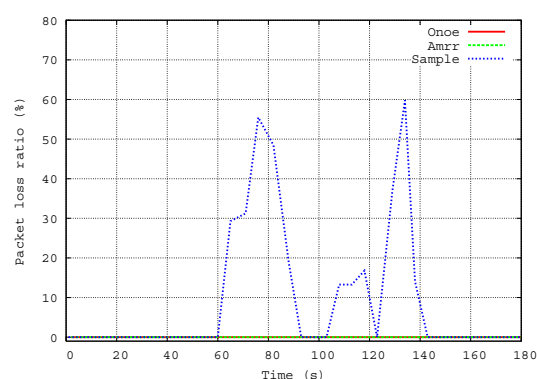


Scénario 2

Figure 4-12 : Le débit d'émission des flux vidéo dans les scénarios 1 et 2

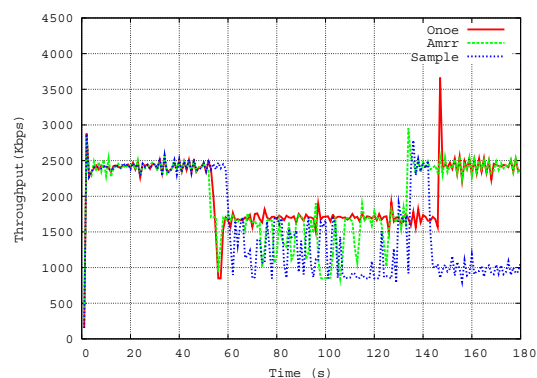


Scénario 1

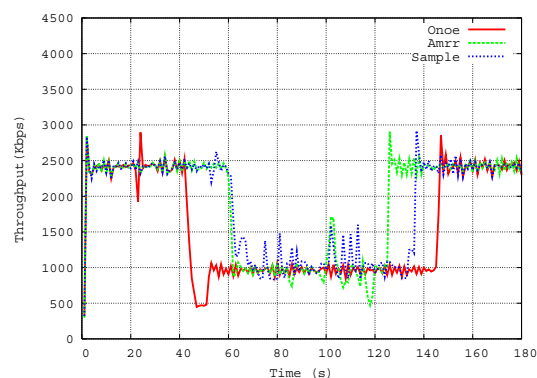


Scénario 2

Figure 4-13 : Les taux de perte des flux vidéo pour les scénarios 1 et 2



Scénario 1



Scénario 2

Figure 4-14 : Le débit de réception des flux vidéo dans les scénarios 1 et 2

Dans le scénario 1, le débit d'émission du flux vidéo reste stable à 2.4 Mbits/s. Dans l'intervalle [0s,55s], le débit physique est encore élevé et il peut couvrir le débit vidéo. Cependant, dans l'intervalle [55s, 140s], le débit physique varie entre 1Mbits/s et 2Mbits/s. Ceci ne permet pas de couvrir la totalité du débit vidéo dont un grand pourcentage est supprimé au niveau de la file d'attente MAC du XLAG. Ceci explique les taux de perte élevés durant cette période. Enfin durant l'intervalle [140, 180s] le débit physique augmente différemment suivant le RCA, sauf pour le RCA *Sample* dont le débit reste instable. En conséquence, les taux de perte diminuent en premier pour *Amrr*, ensuite, pour *Onoe*. Par contre, les pertes restent élevées pour *Sample* à cause de son instabilité.

Dans le scénario 2, le débit d'émission vidéo varie suivant la variation du débit physique. Nous constatons cette variation uniquement durant l'intervalle de temps [40s, 150s]. Ceci est dû à notre politique d'adaptation qui permet au flux vidéo d'occuper la moitié du débit physique. Donc, lorsque le débit physique varie entre 1Mbps/s et 2Mbps/s, le débit vidéo varie entre 500Kbits/s et 1Mbps/s. Mais quand le débit physique est supérieur ou égale à 5.5 Mbps/s le débit vidéo est réinitialisé à son débit original 2.4 Mbps/s. Cependant, la variation du débit vidéo diffère suivant le RCA. La variation optimale est réalisée par *Amrr* puisque le débit vidéo diminue uniquement dans l'intervalle [50s, 125s] contre [40s, 150s] pour *Onoe*. Par contre, l'instabilité de *Sample* se reflète sur la variation du débit vidéo et le flux vidéo enregistre plusieurs pertes. Ces pertes sont causées par l'utilisation d'un débit physique élevé lorsque la station 1 est au niveau du point B. Dans cette situation, les modulations des débits élevés ne résistent pas aux interférences parce que la distance entre le XLAG et la station est importante.

La Figure 4-15 donne les taux moyens de perte pour les dix tests effectués pour chaque RCA dans les deux scénarios. La différence entre les deux scénarios est claire pour tous les RCAs. Nous constatons que des taux de perte sont enregistrés pour chaque RCA dans le scénario 2. Pour *Onoe* et *Amrr*, ceci s'explique par des pertes intempestives causées par des interférences, mais les taux de perte sont modérés, ils ne dépassent pas 1%. Par contre, pour *Sample*, le taux de perte est plus élevé à cause de son instabilité.

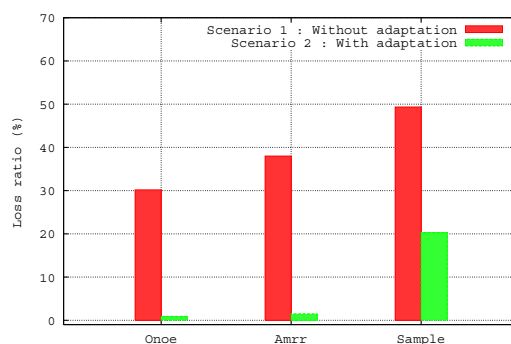


Figure 4-15 : Les taux moyens de perte des flux vidéo pour dix tests dans les scénarios 1 et 2

Pour montrer l'impact de l'adaptation dynamique du débit vidéo sur la qualité vidéo, nous avons calculé deux métriques qui permettent de mesurer la qualité objective de la vidéo, à savoir le PSNR (Peak Signal-to-Noise Ratio) [206] et le SSIM (Structural SIMilarity) [207]. Le calcul de ces deux métriques s'effectue en comparant la séquence vidéo reçue, par la station 1 dans chaque test, avec la séquence vidéo originale. Le PSNR mesure la distorsion entre les images de la vidéo originale et la vidéo reçue et son unité de mesure est le décibel (dB). Le SSIM mesure, quant à lui, le degré de similarité entre les images. Ce degré est donné en pourcentage (%). Les deux métriques sont calculées pour chaque composante couleur de l'image (Cr, Cb, et Y) mais seule la composante Y est tracée dans les graphiques.

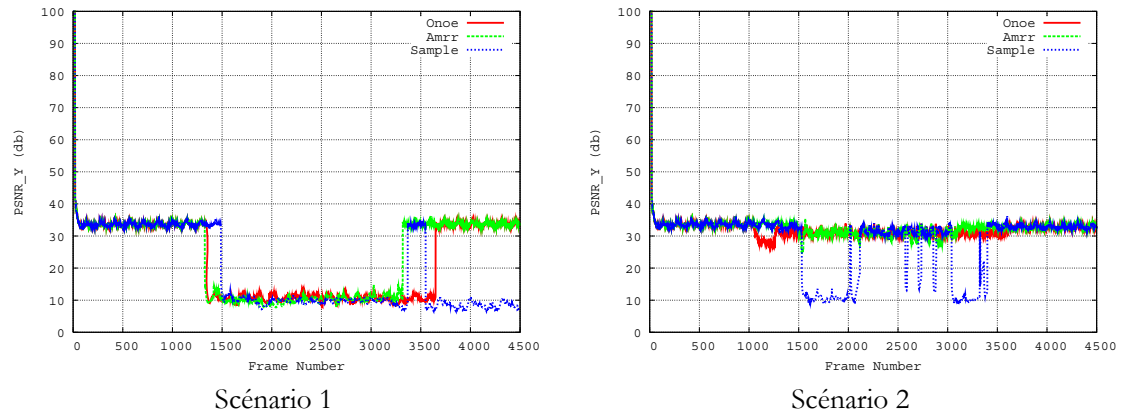


Figure 4-16 : PSNR des flux vidéo dans les scénarios 1 et 2

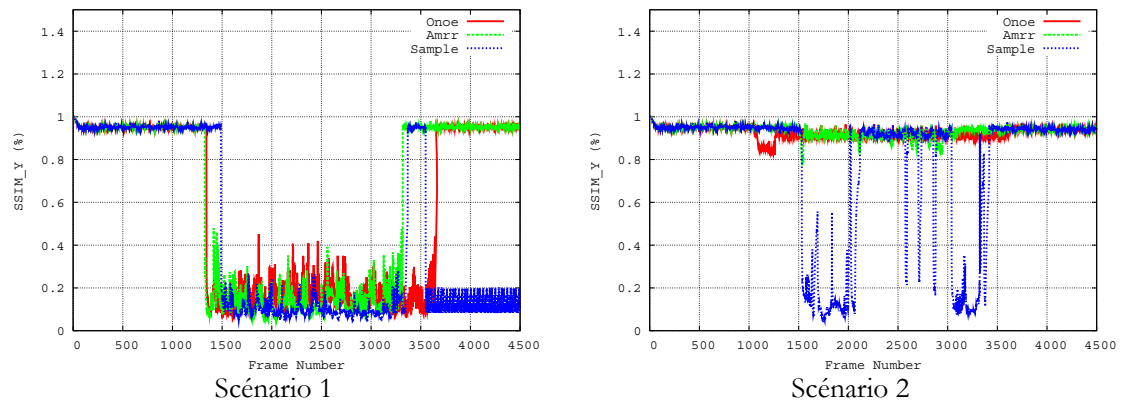


Figure 4-17 : SSIM des flux vidéo dans les scénarios 1 et 2

La Figure 4-16 et la Figure 4-17 donnent respectivement le PSNR et le SSIM de la composante luminance (Y) pour chaque RCA dans les deux scénarios. Nous constatons clairement la différence entre le scénario 1 et le scénario 2. Les pertes de paquets dans le scénario 1 réduisent considérablement la qualité de la vidéo reçue jusqu'à 10 dB pour le PSNR et 10 % pour le SSIM. Par contre, la réduction de la qualité dans le scénario 2, causée par la diminution du débit vidéo, est minime, 30 dB pour le PSNR et 90 % pour le SSIM.

Ainsi, la diminution du débit vidéo possède moins d'impact sur la qualité vidéo comparée aux pertes de paquets. Ceci démontre l'utilité fonctionnelle de notre adaptation dans la préservation de la QoS des flux vidéo durant la transmission.

4.2.5.2 L'adaptation en utilisant plusieurs stations

Cette partie des tests évalue la capacité du débit physique effectif proposé dans la section 4.2.4.1, à évaluer le débit applicatif en présence de plusieurs stations utilisant des débits physiques différents. Pour cela, nous utilisons les trois stations de la plateforme d'expérimentation illustrée dans la Figure 4-10. Les stations 2 et 3 sont immobiles de manière que leur débit physique au niveau du XLAG reste fixe à 54 Mbits/s pour la station 2 et à 11 Mbits/s pour la station 3. Par contre, l'utilisateur de la station 1 est mobile et il se déplace suivant le modèle décrit dans la section précédente. Le XLAG utilise le RCA *Amrr* et le test consiste à transmettre simultanément un flux vidéo pour chaque station. L'objectif de cette configuration est de voir l'impact de la diminution du débit physique de la station 1 sur le débit applicatif des flux destinés aux autres stations.

Nous avons exécuté les deux scénarios décrits dans la section précédente. Cependant, dans cette nouvelle configuration, le débit physique effectif de chaque station varie suivant la mobilité de la station 1. Ci-dessous, nous comparons les résultats obtenus dans les deux scénarios.

Les Figures 4-18, 4-19 et 4-20 donnent respectivement le débit d'émission vidéo, les taux de perte et le débit de réception vidéo pour chaque station dans les deux scénarios.

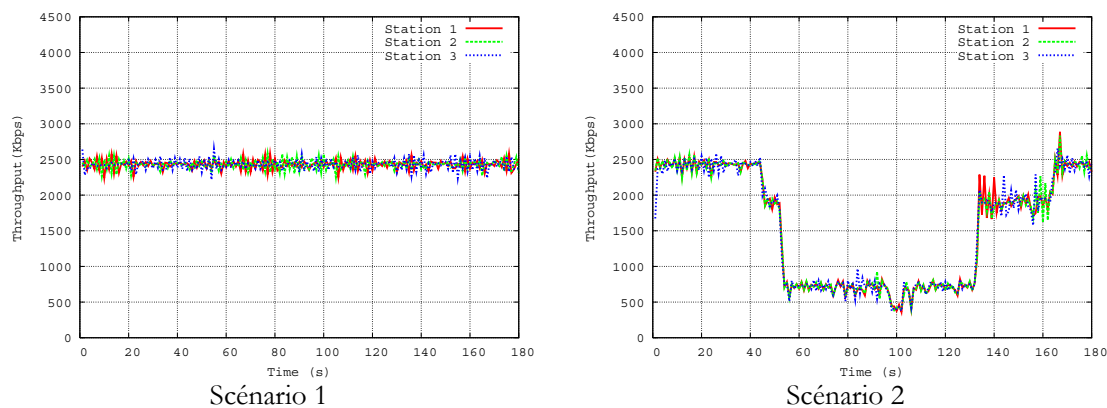


Figure 4-18 : Le débit d'émission des flux vidéo dans les scénarios 1 et 2

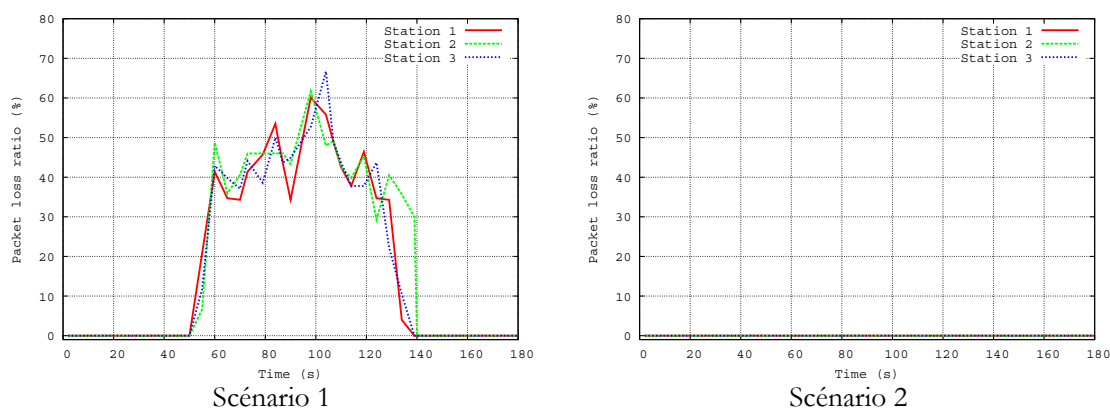


Figure 4-19 : Les taux de perte des flux vidéo pour les scénarios 1 et 2

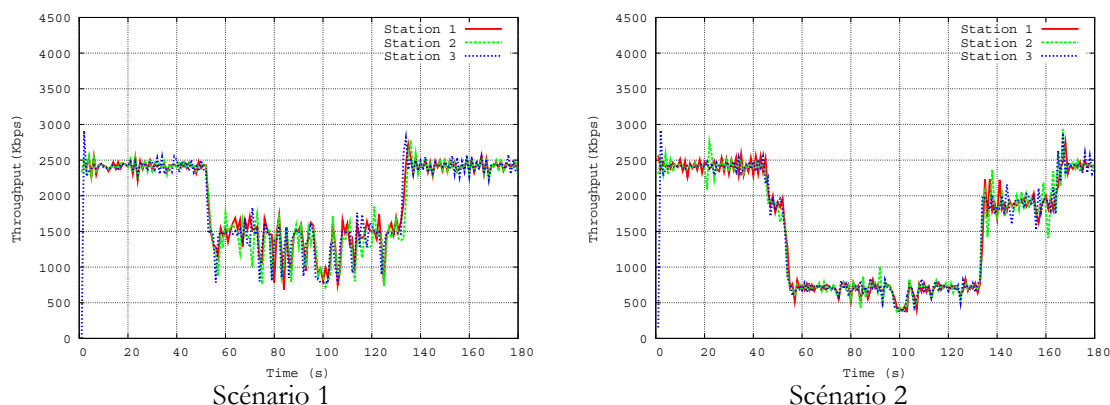


Figure 4-20 : Le débit de réception des flux vidéo dans les scénarios 1 et 2

Dans le scénario 1, lorsque la station 1 est à côté du XLAG, pendant l'intervalle $[0s, 50s]$, les flux ne subissent aucune perte parce que le débit physique de la station 1 est élevé. Cependant, lorsque la station 1 s'éloigne vers le point B et pendant le temps qu'elle reste au niveau de ce point (c-à-d durant l'intervalle $[40s, 140s]$), les trois flux vidéo enregistrent des taux de perte élevés allant

jusqu'à 65 %. Ceci démontre que la réduction du débit physique d'une station affecte le débit applicatif de tous les flux destinés aux autres stations. Enfin, les taux de perte s'annulent lorsque la station 1 entame son retour.

Dans le scénario 2, la station 1 effectue le même déplacement, mais les trois flux n'enregistrent aucune perte. Ceci est dû à la variation du débit d'émission des trois flux vidéo en fonction du débit physique effectif calculé à partir des débits physiques des stations suivant le modèle mathématique introduit. En effet, le débit physique effectif de toutes les stations varie en fonction de la variation du débit physique de la station 1. L'adaptation du débit vidéo n'est pas limitée au flux destiné à la station 1 mais elle est appliquée à tous les flux en fonction du débit physique effectif calculé. Ceci permet de préserver la QoS de tous les flux vidéo dont le débit ne dépasse pas le débit physique réellement disponible. Enfin, lorsque la station 1 se rapproche du point A, le débit d'émission de tous les flux augmente graduellement jusqu'à atteindre le débit original.

Les mesures du PSNR et du SSIM de la composante luminance (Y) pour chaque flux vidéo dans les deux scénarios sont données respectivement dans les Figures 4-21 et 4-22. Ces mesures confirment que les pertes de paquets affectent d'avantage la qualité de la vidéo, comparée à la réduction du débit vidéo basée sur la résolution SNR.

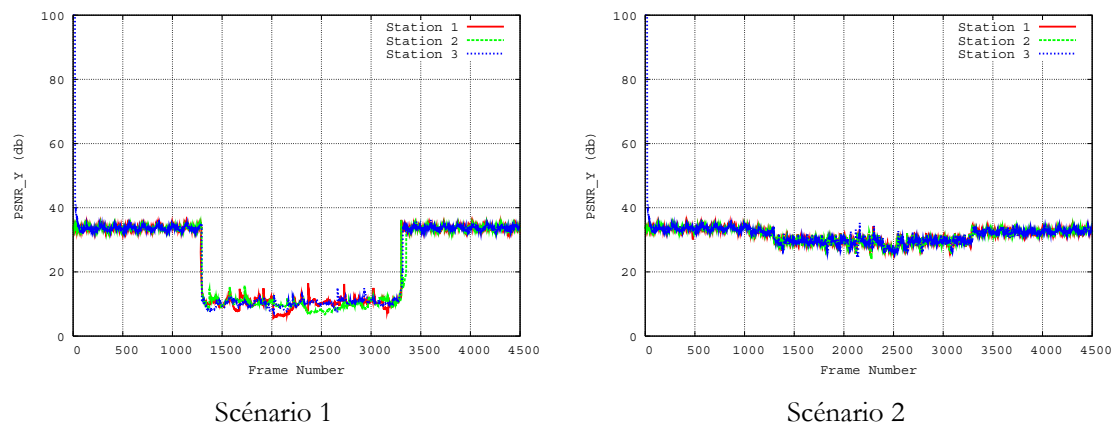


Figure 4-21 : PSNR des flux vidéo dans les scénarios 1 et 2

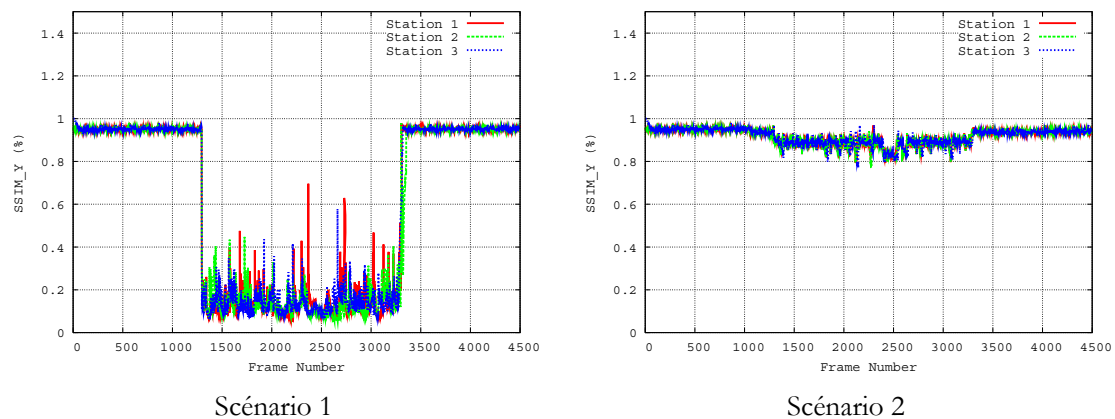


Figure 4-22 : SSIM des flux vidéo dans les scénarios 1 et 2

4.3 Adaptation conjointe de la FEC et du débit vidéo en fonction de la puissance du signal et des taux de perte

4.3.1 Motivation

Dans la section 2.2.5.3.2, nous avons présenté plusieurs mécanismes de QoS de niveau applicatif qui permettent aux applications multimédia de gérer les pertes de paquets. Parmi ces mécanismes, la FEC est une technique très utilisée pour les transmissions temps réel.

Le principe de base de la FEC est la génération de paquets vidéo redondants à partir du flux vidéo original en utilisant un code correcteur particulier. Les paquets originaux et les paquets redondants sont transmis simultanément afin de permettre au récepteur de générer les paquets perdus en temps réel durant la réception du flux.

Ainsi, la FEC représente un mécanisme parfait pour la transmission de flux vidéo sur les réseaux sans fil caractérisés par leur manque de fiabilité et leur dépendance de l'état du canal. Toutefois, deux inconvénients majeurs sont engendrés par l'utilisation de la FEC. Le premier est l'augmentation du débit du flux vidéo causée par l'ajout des paquets redondants. Cette augmentation est variable et dépend principalement du taux de redondance utilisé au niveau du code correcteur. Le deuxième inconvénient est l'augmentation du délai de transmission de bout-en-bout due à l'introduction de délais supplémentaires. Ces derniers sont engendrés par le codage des paquets redondants au niveau de l'émetteur et le décodage des paquets perdus au niveau du récepteur.

Dans le nouveau système de transmission XLAVS, le deuxième inconvénient est moins contraignant puisque le service IPTV est un service à temps différé plus flexible que les services interactifs conversationnels. Par conséquent, l'augmentation du délai causé par la FEC peut être supportée par le service IPTV sans remettre en cause son bon fonctionnement.

Par contre, l'augmentation du débit causé par les paquets redondants représente un inconvénient réel puisque la bande passante au niveau des réseaux 802.11 reste limitée et le partage de cette bande passante entre plusieurs stations la limite d'avantage. Par conséquent, il est impératif d'éviter l'augmentation du débit d'un flux lorsque le mécanisme FEC est utilisé.

De plus, le mécanisme FEC est utile lorsque le canal sans fil se retrouve dans un mauvais état et provoque plusieurs pertes. Par contre, l'utilité de la FEC est réduite lorsque l'état du canal s'améliore.

À partir de ce constat, nous proposons dans cette contribution une utilisation efficace et profitable du mécanisme FEC en se basant sur plusieurs indicateurs. Nous proposons ainsi une technique qui permet de garder un débit de flux stable malgré l'utilisation de la FEC.

Avant de détailler cette contribution, nous présentons dans la section suivante le mécanisme FEC ainsi qu'un ensemble de codes correcteurs largement utilisés.

4.3.2 Le mécanisme FEC au niveau paquet

Actuellement, il existe plusieurs codes correcteurs qui peuvent être exploités par un mécanisme FEC. Chaque code possède des caractéristiques spécifiques comme sa complexité algorithmique et son taux de redondance maximal. Les codes correcteurs sont définis théoriquement sur des symboles qui peuvent avoir différentes tailles. La génération des symboles redondants est appelée codage et la génération des symboles perdus à partir des symboles redondants est appelée décodage.

Pour la transmission de données, ces codes peuvent être appliqués sur deux niveaux distincts : au niveau bit et au niveau paquet. Dans le premier niveau, le symbole est un bit et le code correcteur génère des bits redondants. Par contre, dans le deuxième niveau, un paquet est considéré comme un symbole et des paquets redondants sont codés à partir des paquets originaux.

Lorsque le mécanisme FEC est utilisé au niveau applicatif. La FEC au niveau paquet est plus avantageuse pour différentes raisons listées ci-dessous :

- Les pertes dans le réseau s'effectuent par paquet et la FEC au niveau bit ne peut pas décoder un paquet perdu entièrement.
- La FEC au niveau paquet offre une capacité de décodage supérieure puisque un paquet redondant, codé à partir d'un groupe de paquets, peut décoder n'importe quel paquet perdu dans ce groupe.
- Lorsque la FEC au niveau bit est utilisée, les paquets corrompus durant la transmission, principalement dans les réseaux sans fil, n'atteignent jamais la couche applicative puisque ces paquets sont supprimés automatiquement au niveau MAC durant la vérification du CRC.

Pour la transmission de flux multimédia, le mécanisme FEC est implémenté au niveau du protocole RTP afin d'exploiter la numérotation séquentiel des paquets présente dans l'en-tête RTP. Pour cela, l'IETF a défini dans le RFC 2733 [124], le format des paquets RTP qui utilise un mécanisme FEC. Le RFC est défini d'une manière générique pour assurer les fonctionnalités citées ci-dessous :

- Indépendance vis-à-vis du flux protégé (audio, vidéo ou texte).
- Indépendance du code correcteur utilisé.
- Adaptation dynamique des paramètres FEC sans signalisation supplémentaire.
- Support de plusieurs schémas pour la transmission des paquets FEC.

Dans la suite de cette section, nous présentons les codes correcteurs qui ont rencontré un grand succès durant ces dernières années et qui peuvent être utilisés au niveau paquet en se basant sur le RFC 2733.

4.3.2.1 Le code XOR (eXclusive OR)

Le code XOR, illustré par la Figure 4-23, permet de coder un paquet redondant à partir d'un groupe de paquets de même taille en appliquant sur leur contenu l'opération ou exclusif (XOR). Le taux de redondance varie en fonction de la taille du groupe de paquets. Par exemple, si le groupe contient 4 paquets, le taux de redondance est de 25%.

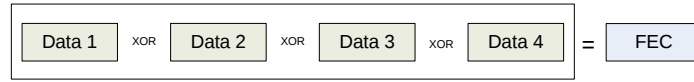


Figure 4-23 : Le code correcteur XOR à une dimension

L'avantage de ce code est sa simplicité qui ne nécessite pas une grande puissance de calcul pour coder et décoder les paquets. Ceci réduit, par conséquent, le délai de transmission de bout-en-bout. Par contre, la capacité de décodage du code XOR est limitée puisqu'il ne peut décoder qu'un seul paquet perdu dans le groupe. Afin d'augmenter cette capacité, certains mécanismes FEC appliquent le codage XOR sur deux dimensions, illustré par la Figure 4-24.

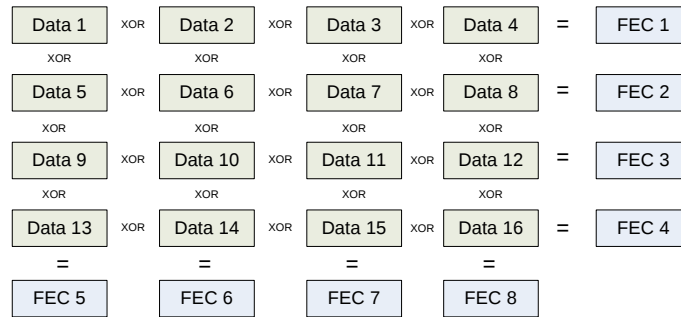


Figure 4-24 : Le code correcteur XOR à deux dimensions

Le forum Pro-MPEG qui réunit différents acteurs de la TV numérique, propose dans la version 2 de la recommandation MPEG CoP (Code of Practice #3 release 2) [208] un mécanisme FEC XOR sur deux dimensions. L'objectif de la recommandation MPEG CoP est de définir le transport des flux MPEG-2 TS sur les réseaux IP. Le MPEG CoP spécifie le nombre de lignes D , le nombre de colonnes L et le nombre de paquets maximal dans une matrice FEC $L \times D$ comme suit :

$$4 \leq D \leq 20 \quad 1 \leq L \leq 20 \quad L * D \leq 100$$

4.3.2.2 Le code RS (Reed Salomon)

Le code RS forme une classe spéciale de code linéaire cyclique non binaire très puissant capable de corriger plusieurs erreurs. Il est construit sur un ensemble fini de symboles, appelé corps ou champ de Galois $CG(q)$. Un code RS (n, k) génère n paquets (avec $n = q - 1$) à partir de k paquets sources en ajoutant h paquets redondants (avec $h = n - k$) comme illustré par la Figure 4-25.

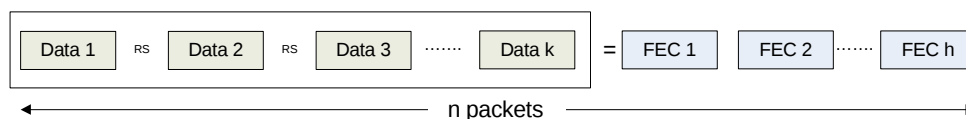


Figure 4-25 : Le code correcteur RS

L'avantage du code RS réside dans sa capacité de décodage. Par contre, son inconvénient majeur est sa complexité algorithmique lorsqu'il est utilisé au niveau logiciel. Cependant, il est

largement utilisé au niveau matériel dans les systèmes de communication DVB-S et sur les supports de stockage numérique DVD.

4.3.2.3 Le code LDPC (Low Density Parity Check)

Le code LDPC (n, k) est un code linéaire qui crée une équation linéaire entre les k paquets sources et les h paquets redondants avec $h = n - k$. L'opération XOR est utilisée pour coder un paquet redondant à partir d'un sous-ensemble des paquets sources. La relation entre les paquets sources et les paquets redondants peut être définie soit par un graphe biparti, soit par une matrice à deux dimensions. Ceci est illustré par la Figure 4-26, la matrice doit être creuse pour pouvoir décoder le maximum de paquets perdus. Les lignes et les colonnes correspondent respectivement aux paquets FEC et aux paquets de données. Lorsque l'élément de la matrice $H(i, j) = 1$, le paquet FEC i contient le paquet de données j .

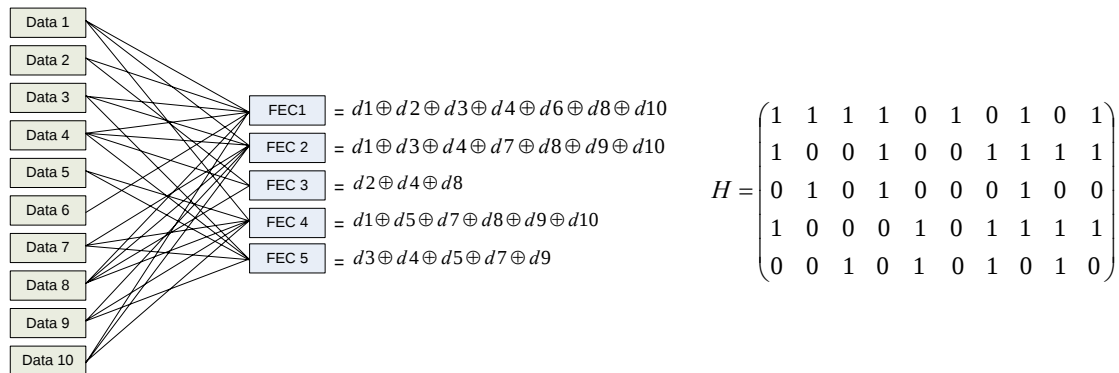


Figure 4-26 : Le code correcteur LDPC

Le code LDPC offre un compromis entre les mauvaises performances du codage XOR et la complexité du codage RS. Il est utilisé, principalement, dans le standard de diffusion par satellite DVB-S2.

4.3.3 L'adaptation conjointe du débit vidéo et du taux de redondance FEC

Lorsqu'un mécanisme FEC de type RS (n, k) est utilisé au niveau paquet entre deux stations communicantes, le récepteur a la capacité de décoder $h = n - k$ paquets perdus. En d'autre terme, le récepteur peut résister contre les pertes avec une probabilité maximale p_{max} qui peut être calculée par l'équation Eq. 4-6.

$$p_{max} = \frac{h}{n} \quad \text{Eq. 4-6}$$

En se basant sur cette équation, le taux de redondance utilisé par le mécanisme FEC au niveau de l'émetteur peut être calculé à partir des taux de perte que subisse le flux durant la transmission. Ces taux de perte peuvent être déterminés facilement au niveau du récepteur et rapportés périodiquement à l'émetteur en utilisant, par exemple, les rapports du récepteur (RR : Receiver Report) dans le protocole RTCP.

$$h = \frac{p \cdot k}{1 - p} \quad \text{Eq. 4-7}$$

L'équation Eq. 4-7 permet de déterminer le taux de redondance dynamiquement en se basant sur les taux de perte instantanés. Plusieurs travaux de recherche utilisent ce type d'approche pour adapter le mécanisme FEC en fonction des conditions de transmission réseau et réduire, par la même, l'*overhead* global (débit supplémentaire) introduit par le mécanisme FEC. Le taux d'*overhead* est donné dans l'équation Eq- 4.8.

$$ov_{FEC} = \frac{h}{k} \quad \text{Eq. 4-8}$$

Cependant, ce type d'adaptation n'est pas efficace puisqu'il ne permet pas d'éviter les pertes mais il réagit une fois que les pertes se sont produites. De plus, ces performances sont liées à la fréquence des rapports de perte. Une fréquence élevée permet à l'émetteur de répondre rapidement à l'apparition de pertes mais elle engendre un *overhead* important dans la voie de retour entre le récepteur et l'émetteur. Par contre, une fréquence faible réduit la réactivité de l'émetteur.

Afin d'avoir un mécanisme FEC adaptatif qui anticipe l'apparition des pertes, il faut utiliser d'autres métriques qui informent sur la probabilité d'apparition de ces pertes. De plus, ces métriques doivent être disponibles au niveau de l'émetteur.

Dans le cadre de notre architecture qui utilise un réseau sans fil 802.11, ces métriques peuvent correspondre à celles qui informent sur l'état du canal puisque ce dernier est lié directement à l'apparition des pertes. Dans la littérature, le SNR est largement utilisé pour déterminer l'état du canal, cependant, nous avons expliqué dans la section 4.2.2.2 qu'il n'est pas possible d'obtenir le SNR en temps réel, que se soit du côté émetteur ou récepteur.

Pour cela, nous reprenons la solution introduite par les auteurs dans [203] qui proposent d'utiliser la puissance du signal SSI en remplacement du SNR (voir section 4.2.2.2). De plus, les auteurs utilisent le SSI des acquittements reçus au niveau de l'émetteur au lieu du SSI des trames reçues par le récepteur puisque le canal de transmission est symétrique.

Ainsi, nous proposons un mécanisme FEC adaptatif qui se base sur le SSI et les taux de perte pour faire varier le taux de redondance. L'objectif de ce mécanisme est d'avoir un flux vidéo résistant aux pertes, mais dont la résistance s'adapte aux conditions de transmission.

Pour cela, l'état du canal est représenté suivant le modèle de Gilbert [209] et Elliott [210]. Dans ce dernier, l'état du canal est modélisé par une chaîne de Markov qui comprend deux états : l'état bon (Good) et l'état mauvais (Bad).

La Figure 4-27 illustre les différents états dans lesquels peut se trouver le système FEC adaptatif ainsi que les différentes transitions entre les états. Ces derniers sont définis sur deux niveaux distincts.

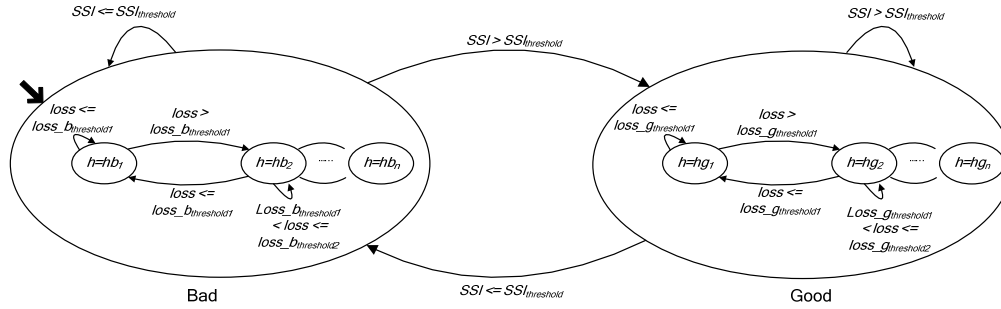


Figure 4-27 : Les différents états du système FEC adaptatif

Le premier niveau définit l'état du système par rapport à l'état du canal. Ce niveau possède deux états, la transition entre ces deux états se base sur la puissance du signal SSI . L'état initial est l'état mauvais, si le SSI est inférieur ou égal à un seuil $SSI_{threshold}$, le système se maintient dans l'état mauvais. Mais si le SSI dépasse le $SSI_{threshold}$, le système passe vers l'état bon et il se maintient dans cet état tant que la condition sur le SSI est vérifiée sinon il repasse vers l'état mauvais.

À l'intérieur de chaque état, le système possède des sous-états qui correspondent au deuxième niveau. Les sous-états représentent aussi une chaîne de Markov finie. La transition entre les états se base sur les taux de perte $loss$ enregistrés durant la transmission. Chaque état utilise un taux de redondance b , $b = \{hb_1, hb_2 \dots hb_n\}$ dans l'état mauvais et $b = \{hg_1, hg_2 \dots hg_n\}$ dans l'état bon. Le taux de redondance peut être nul, ce qui signifie que le mécanisme FEC est désactivé.

Il est évident que la variation du taux de redondance b suivant le modèle proposé ci-dessus permet une utilisation efficace du mécanisme FEC. Ceci réduit globalement l'*overhead* introduit par la FEC comparé à un système qui utilise un taux de redondance fixe quelque soit les conditions de transmission.

Cependant, l'*overhead* reste un problème dans notre système, principalement lorsque ce dernier se trouve dans l'état mauvais. En effet, dans cet état les capacités spectrales du canal sont réduites est l'augmentation du débit à cause du mécanisme FEC peut engendrer d'avantage de perte de paquets. Donc, il est important de maintenir un débit stable pour les flux transmis quelque soit l'état du système.

Pour atteindre cet objectif, nous proposons une variation conjointe du taux de redondance b et du débit vidéo en utilisant la résolution SNR présentée dans la section 4.2.3.2.2. Ceci signifie que le débit vidéo peut diminuer ou augmenter suivant l'*overhead* introduit par la FEC afin de maintenir un débit moyen stable.

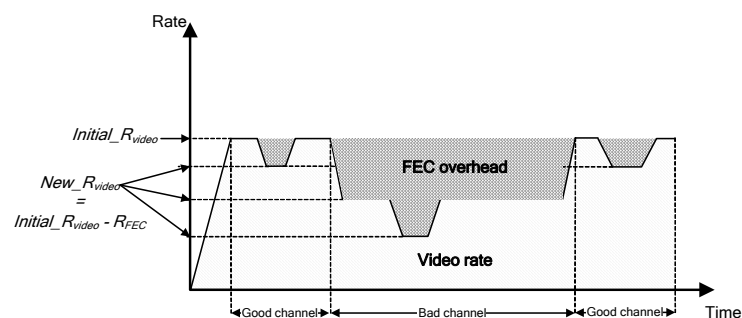


Figure 4-28 : L'adaptation conjointe du taux de redondance FEC et du débit vidéo

Ce procédé est illustré par la Figure 4-28. La vidéo possède un débit initial $Initial_R_{video}$. Lorsque le taux de redondance h est supérieur à zéro, un nouveau débit vidéo New_R_{video} est calculé en réduisant le débit FEC, engendré par le taux de redondance h , du débit vidéo initial $Initial_R_{video}$. Par contre, lorsque le mécanisme FEC est désactivé ($h = 0$) le débit vidéo est réinitialisé avec son débit initial.

En diminuant le débit vidéo, nous diminuons aussi la qualité vidéo, mais le débit global du flux reste stable. De plus, le flux est plus résistant aux pertes qui dégradent plus la qualité vidéo que la diminution du débit. Dans la section suivante, nous détaillons la mise en œuvre de ce mécanisme au niveau de la passerelle d'adaptation XLAG.

4.3.3.1 Mise en œuvre de l'adaptation au niveau du XLAG

La Figure 4-29 illustre les interactions nécessaires entre les modules au niveau du XLAG pour mettre en œuvre l'adaptation conjointe du mécanisme FEC et du débit vidéo. Cette adaptation est exécutée au niveau du module « FEC » et du module « transcoder », mais la prise de décision concernant le changement d'état est mise en œuvre par le module « XLDP ». L'adaptation est exécutée durant la transmission du flux vidéo pour chaque session IPTV active au niveau du XLAVS. Pour cela, le XLDP récupère les deux métriques nécessaires pour le fonctionnement de cette adaptation, à savoir les taux de perte que subisse le flux vidéo et la puissance du signal des acquittements transmis par la station qui reçoit ce flux vidéo.

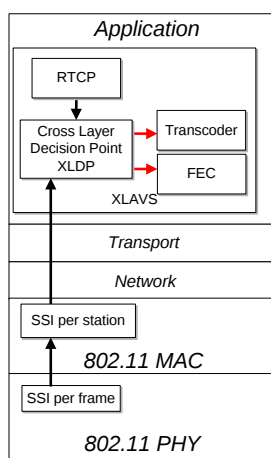


Figure 4-29 : L'interaction entre les modules au niveau du XLAG pour l'adaptation conjointe du mécanisme FEC et du débit vidéo

Les taux de perte sont récupérés du module RTCP qui reçoit les rapports RR (Receiver Report) contenant les taux de perte calculés au niveau du récepteur. Cependant, ce taux doit être calculé avant le décodage FEC pour avoir le taux de perte réel causé par la transmission. La fréquence des rapports dépend du débit du flux vidéo. Suivant le RFC 3550 [92], le débit du flux RTCP ne doit pas dépasser 5% du débit total de la session.

Pour récupérer la puissance du signal de chaque station, deux modules sont utilisés : un module au niveau physique « SSI per frame » et un module au niveau MAC « SSI per station ». Le module « SSI per station » calcule la puissance moyenne du signal pour chaque station. Le calcul de cette moyenne se base sur la puissance du signal de chaque trame « ACK » fournie par le module « SSI per frame ». La moyenne calculée est une moyenne pondérée dans le temps suivant l'équation Eq. 4-9. Elle donne plus d'importance à la puissance du signal des derniers acquittements reçus.

$$WSSI \leftarrow (1 - \lambda') \times SSI + \lambda' \times WSSI \quad \text{Eq. 4-9}$$

L'utilisation de ces deux modules permet d'agréger au niveau MAC la valeur SSI remontée au niveau applicatif. Par la suite, le module XLDP interroge le module « SSI per station » par un appel de fonction interne « *get_ssi (Mac_@)* » pour récupérer la puissance du signal pour une station donnée. La station est identifiée par son adresse MAC. Pour avoir cette adresse, le module XLDP effectue une résolution inverse de l'adresse IP de la station réceptrice au début de la session IPTV.

Ainsi, en fonction de la puissance du signal et des taux de perte, le module « XLDP » détermine son état et détermine, par la même, le taux de redondance FEC et le débit de la vidéo. Ces deux paramètres sont communiqués, respectivement, au module « FEC » et au module « transcoder ».

Le taux de redondance dans chaque état est fixé suivant une politique d'adaptation paramétrable au niveau du système XLAVS.

4.3.4 Tests et évaluations de performance

Pour évaluer notre contribution, nous avons utilisé la plateforme d'expérimentation présentée dans la section 4.2.5. Le XLAG transmet un flux vidéo à la station 1 qui se déplace du point A vers le point B comme illustré par la Figure 4-10. La station 1 utilise le modèle de déplacement décrit dans la section 4.2.5.1, c'est-à-dire, déplacement du point A vers le point B au bout de 60s, arrêt 60s au niveau du point B et retour au point A au bout de 60s. Les caractéristiques techniques du flux vidéo sont données dans le Tableau 4-1, le débit moyen de la vidéo est de 611 Kbits/s, le débit maximum est de 760 Kbits/s et le débit minimum est de 472 Kbits/s. Pour voir l'impact de la dégradation du signal sur la transmission du flux vidéo, le débit physique au niveau du XLAG a été limité à 5.5 Mbits/s.

Le mécanisme FEC implémenté par le nouveau système de transmission XLAVS au niveau du module « FEC » se base sur le codage RS. L'implémentation au niveau RTP est compatible avec le RFC 2733.

Au niveau du lien d'accès 802.11, les cartes WiFi *Atheros*, à travers le driver *MadWiFi*, fournissent deux paramètres qui permettent de déterminer l'état du canal. Le premier est la puissance du Signal *SSI* qui est mesuré en milli-deciBel (dBm) et qui varie dans l'intervalle [-96,0].

La mesure en milli-deciBel s'explique par le fait que la puissance du signal WiFi est très faible inférieure à 1mW qui est égal à 0dBm. L'équation Eq. 4-10 est utilisée pour convertir la puissance du signal du mW en dBm.

$$dBm = \log(mW) \times 10 \quad \text{Eq. 4-10}$$

Ainsi, le signal est au plus haut niveau lorsque le $SSI = 0$ dBm et au plus bas niveau lorsque le $SSI = -96$ dBm.

Le deuxième paramètre, appelé $RSSI$ (Received Signal Strength Indicator), correspond à la qualité du signal. Le $RSSI$ est utilisé par différentes cartes WiFi pour donner une indication sur la qualité du signal. Il ne possède pas d'unité de mesure et il varie dans un intervalle précis $[0, RSSI_{max}]$. Le $RSSI_{max}$ varie en fonction des constructeurs de cartes WiFi et la procédure de calcul du $RSSI$ n'est pas documentée. Pour les cartes *Atheros*, le $RSSI$ varie dans l'intervalle $[0, 94]$, où, $RSSI = 0$ représente le plus bas niveau de qualité et $RSSI = 94$ représente le plus haut niveau.

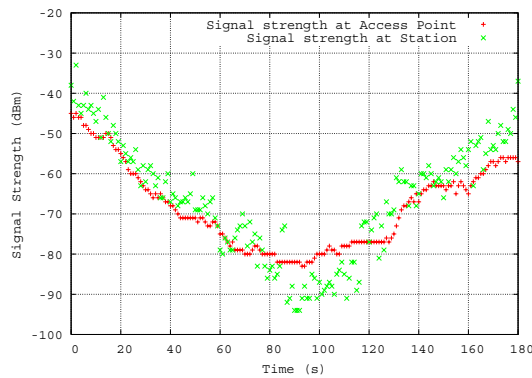


Figure 4-30 : La puissance du signal perçue simultanément au niveau du XLAG et au niveau de la station 1

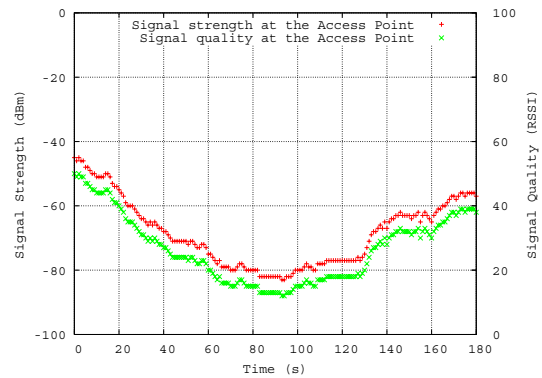


Figure 4-31 : La puissance et la qualité du signal perçues au niveau du XLAG

Avant d'évaluer l'adaptation proposée, nous avons effectué un premier test pour valider la corrélation entre la puissance du signal perçue au niveau du XLAG et au niveau de la station 1. Pour cela, la Figure 4-30 donne le SSI perçu au niveau du XLAG grâce à la réception des acquittements et le SSI perçu au niveau de la station 1 grâce à la réception des trames vidéo. Malgré une légère différence entre les deux niveaux du signal, nous constatons que les $SSIs$ dans les deux extrémités sont fortement corrélés puisqu'ils suivent le même modèle de variation lorsque la station 1 se déplace.

Le deuxième test effectué avait pour objectif de montrer la différence entre la puissance du signal SSI et la qualité du signal $RSSI$ de la station 1 recueillie au niveau du XLAG. La Figure 4-31, illustre dans le même graphe le SSI (axe y gauche) et le $RSSI$ (axe y droit) de la station 1 au niveau du XLAG, tout en long d'un test. Elle montre que les deux paramètres possèdent une variation identique. Ceci permet de conclure que le $RSSI$ est calculé à partir du SSI .

Ainsi, nous avons choisi d'utiliser le taux de $RSSI$ qui représente le rapport entre le $RSSI$ instantané et le $RSSI_{max}$ pour le basculement entre les états du système. Ceci permet d'avoir un taux indépendant du $RSSI_{max}$ qui diffère suivant la carte WiFi utilisée.

4.3.4.1 Évaluation du mécanisme FEC adaptatif

Pour cette évaluation, nous avons comparé notre contribution à deux systèmes, l'un traditionnel qui n'utilise pas de mécanisme FEC et le deuxième utilisant un mécanisme FEC adaptatif basé uniquement sur le taux de perte. Pour cela, nous avons défini trois scénarios et nous avons effectué dix tests dans chaque scénario. Les différents scénarios sont décrits ci-dessous :

- **Scénario 1** : le XLAVS effectue une simple transmission de vidéo sans aucune adaptation
- **Scénario 2** : Le XLAVS utilise un mécanisme FEC adaptatif. L'adaptation de la FEC se base sur les taux de perte. La politique d'adaptation est donnée dans le Tableau 4-5.

Taux de perte (%)	FEC (n, k)	Overhead
= 0 %	Pas de FEC	0%
≤ 5 %	FEC (30, 24)	+ 25%
> 5 %	FEC (30, 21)	+ 42%

Tableau 4-5 : La politique d'adaptation dans le scénario 2

- **Scénario 3** : Le XLAVS utilise notre contribution qui adapte la FEC et le débit vidéo en fonction de la puissance du signal et en fonction des taux de perte. Les états du système sont décrits dans le Tableau 4-6. Le $RSSI_{threshold} = 45\%$.

RSSI (%)	Taux de perte (%)	FEC (n, k)	Taux de réduction du débit vidéo par rapport au débit initial
Good [45%, 100%]	= 0 %	Pas de FEC	0%
	≤ 5 %	FEC (30, 24)	- 25%
	> 5 %	FEC (30, 21)	- 42%
Bad [0%, 45%]	= 0 %	FEC (30, 27)	- 11%
	≤ 5 %	FEC (30, 24)	- 25%
	> 5 %	FEC (30, 21)	- 42%

Tableau 4-6 : La politique d'adaptation dans le scénario 3

Nous comparons ci-dessous les résultats de tests représentatifs obtenus dans chaque scénario. La Figure 4-32 donne la variation du taux RSSI pour tous les scénarios durant un test. Nous constatons clairement la dégradation du taux RSSI lorsque la station 1 s'éloigne du XLAG et son augmentation lorsque la station entame son retour. De plus, la variation est relativement identique pour tous les scénarios. Ceci permet de comparer les résultats obtenus pour chacun d'eux.

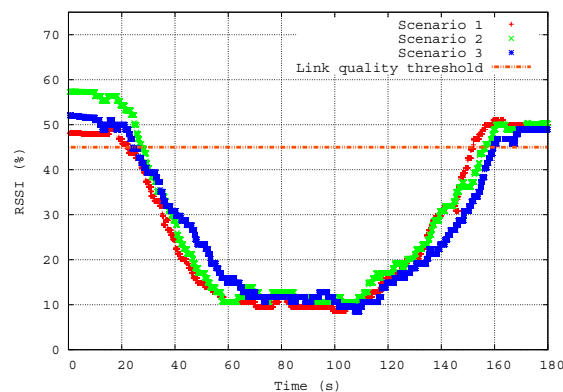


Figure 4-32 : La variation du RSSI pour les scenarios 1, 2 et 3

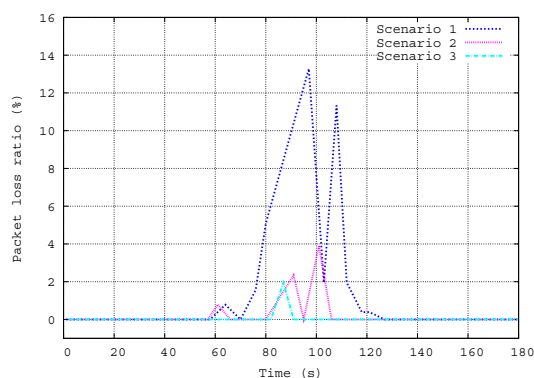


Figure 4-33 : Les taux de perte pour les scenarios 1, 2 et 3 après le décodage FEC

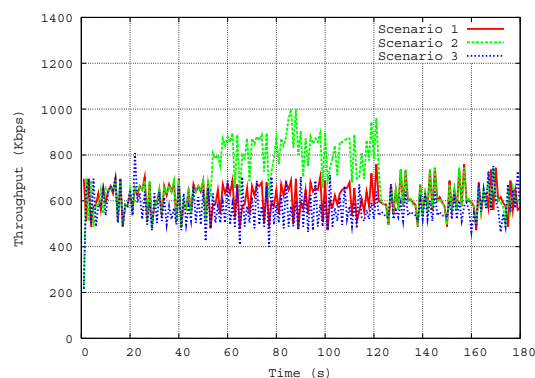


Figure 4-34 : Le débit d'émission des flux video pour les scenarios 1, 2 et 3

Les Figures 4-33 et 4-34 donnent respectivement les taux de perte et le débit d'émission du flux vidéo pour chaque scénario. Les taux de perte pour le scénario 2 et 3 donnés dans la Figure 4-33, sont calculés après le décodage FEC, mais l'adaptation se base sur les taux de perte calculés avant le décodage FEC.

La Figure 4-33 illustre la différence entre les scénarios concernant les taux de perte. Ces pertes se produisent dans l'intervalle [50s, 130s] lorsque la station 1 est le plus loin possible du XLAG. Les taux les plus élevés sont enregistrés dans le scénario 1 puisque ce dernier n'utilise aucune protection contre les pertes. Ces pertes sont moins importantes dans le scénario 2 qui utilise un mécanisme FEC adaptatif en fonction des taux de perte seulement. Cependant, nous remarquons dans ce scénario une oscillation des pertes qui s'explique par le fait que l'utilisation du mécanisme FEC n'est pas stable, il est activé ou désactivé suivant les taux de perte rapportés par les paquets RTCP. Cette situation ne se produit pas dans le scénario 3, qui enregistre le taux de perte le plus bas, parce que le mécanisme FEC ne dépend pas uniquement des taux de perte, mais aussi du taux du RSSI. Ce dernier est dans le bas niveau dans l'intervalle [50s, 130s], ce qui indique au XLAVS que le canal de transmission est dans un état mauvais et que la probabilité d'avoir des pertes est plus élevée. En appliquant la politique d'adaptation décrite dans le Tableau 4-6, le XLAVS maintient la FEC active durant toute cette période. De plus, dans le scénario 3, le débit du flux vidéo reste stable comparé

aux scénarios 1 et 2, comme illustré par la Figure 4-34. Ceci participe aussi à la diminution des taux de perte puisque le canal qui est dans un état mauvais ne subit pas une augmentation du débit.

Pour confirmer ces résultats, la Figure 4-35 donne les taux moyens de perte (avant et après le décodage FEC) dans l'intervalle [60s, 120s] pour les dix tests réalisés dans chaque scénario. Dans le scénario 2, le taux moyen de perte avant le décodage FEC est plus élevé par rapport à ceux des scénarios 1 et 3. Ceci confirme que l'augmentation du débit causé par la FEC engendre plus de pertes. Par contre, le taux moyen de perte avant le décodage FEC dans le scénario 3 est au même niveau que celui du scénario 1 puisque ces deux scénarios possèdent un débit d'émission relativement identique. Pour le taux moyen de perte après le décodage FEC, il est clair que le scénario 3 permet le décodage d'un maximum de paquets perdus par rapport au scénario 2 puisque la FEC est active durant tout l'intervalle [60s, 120s].

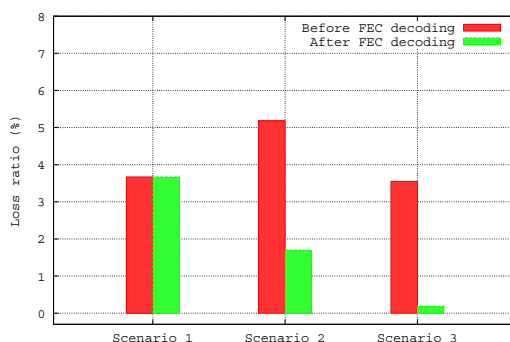


Figure 4-35 : Les taux moyens de perte des flux vidéo pour dix tests dans les scénarios 1, 2 et 3

Par conséquent, les résultats obtenus démontrent que l'adaptation proposée dans le scénario 3 offre de meilleures performances vis-à-vis des taux de perte en anticipant la dégradation du canal de transmission.

Pour voir l'impact de cette adaptation sur la qualité vidéo, nous avons calculé le PSNR et le SSIM des séquences vidéo reçues dans chaque scénario par rapport à la séquence originale. Les résultats sont donnés respectivement dans les Figures 4-36 et 4-37. Les résultats fournis par les deux métriques sont légèrement différents. Avec le PSNR, nous remarquons que la qualité vidéo a diminué dans le scénario 3 par rapport aux scénarios 1 et 2, mais avec le SSIM, la différence n'est pas visible. Ceci est dû à la méthode de calcul de chaque métrique. Cependant, nous constatons, pour les deux métriques, que la dégradation de la qualité vidéo est plus importante avec les pertes de paquets dans les scénarios 1 et 2. Ainsi, le scénario 3 réduit légèrement la qualité globale de la vidéo reçue mais évite, par la même, les dégradations majeures causées par les pertes.

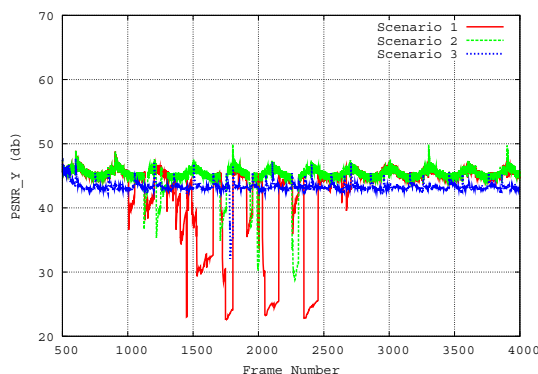


Figure 4-36 : PSNR des flux vidéo pour les scénarios 1, 2 et 3

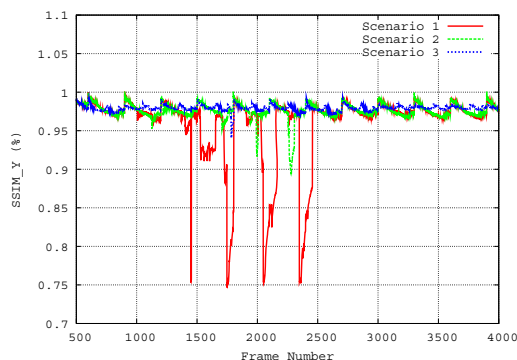


Figure 4-37 : SSIM des flux vidéo pour les scénarios 1, 2 et 3

L'adaptation conjointe du taux de redondance FEC et du débit vidéo en se basant simultanément sur le niveau du signal et des taux de perte permet donc d'améliorer la QoS d'un flux vidéo qui devient plus résistant aux pertes tout en gardant son débit initial.

4.4 Conclusion

En exploitant les approches *Cross-layer* ascendantes, le nouveau système XLAVS se tient au fait des variations qui peuvent se produire au niveau du lien d'accès 802.11 durant la transmission des flux IPTV. Le XLAVS adapte son fonctionnement par rapport à ces variations afin de maintenir la QoS des flux transmis.

Dans ce chapitre, nous avons présenté deux types d'adaptation qui se basent sur l'approche *Cross-layer* ascendantes.

La première adaptation présentée assure une correspondance entre le débit applicatif (le débit vidéo) et le débit physique disponible dans le réseau 802.11. À cet effet, nous avons présenté les différents types d'algorithmes (RCAs) qui assurent la variation du débit physique au niveau MAC 802.11. Ensuite, nous avons présenté deux techniques qui permettent de faire varier le débit vidéo dynamiquement durant la transmission, à savoir (1) la variation de la résolution temporelle et (2) la variation de la résolution SNR basée sur le transcodage. Nous avons montré que la deuxième technique offrait un meilleur contrôle du débit et une meilleure qualité vidéo. Le débit vidéo est décidé suivant le débit physique effectif disponible au niveau du XLAG. Pour cela, nous avons proposé un modèle mathématique pour calculer le débit physique effectif en se basant sur le débit physique utilisé pour chaque station connectée au XLAG. Les tests de performance réalisés ont montré, premièrement, l'impact du RCA sur l'adaptation proposé et, deuxièmement, la capacité de cette adaptation à préserver la QoS des flux vidéo en exploitant judicieusement le débit physique réellement disponible.

La deuxième adaptation présentée dans ce chapitre est l'adaptation conjointe du taux de redondance FEC et du débit vidéo suivant la puissance du signal et les taux de perte. Le mécanisme FEC est appliqué au niveau applicatif sur les paquets vidéo avec un taux de redondance variable. La puissance du signal est utilisée comme un indicateur de l'état du canal afin de prévenir les pertes de paquets. Pour éviter l'augmentation du débit causée par le mécanisme FEC, nous avons proposé une adaptation dynamique du débit vidéo en fonction du taux de redondance FEC. Avec cette

adaptation conjointe, le flux vidéo garde son débit initial tout en étant plus résistant aux pertes. Les résultats des tests effectués démontrent l'avantage de faire collaborer plusieurs mécanismes de QoS en se basant sur plusieurs métriques de performance puisque l'adaptation proposée permet de réduire considérablement les pertes de paquets et évite ainsi la dégradation de la qualité vidéo perçue par le récepteur.

Chapitre 5

Les adaptations *Cross-layer* descendantes exécutées au niveau MAC 802.11

5.1 Introduction

Dans ce chapitre, nous continuons nos investigations concernant les adaptations *Cross-layer* qui peuvent être exécutées par la passerelle d'adaptation XLAG (Cross-Layer Adaptation Gateway) au cours d'une session IPTV. Cette fois-ci, nous nous intéressons aux adaptations *Cross-layer* descendantes (Top-Down) dans lesquelles les couches supérieures configurent dynamiquement les couches inférieures, ou bien, le fonctionnement de ces dernières se base sur des informations des couches supérieures (section 2.3.3). Ce type d'adaptation est nécessaire pour le nouveau système XLAVS qui doit contrôler la QoS des couches sous-jacentes dans l'objectif de maximiser la QoS des flux transmis.

Ce chapitre est organisé en deux contributions principales. La première contribution permet de réduire les pertes de paquet en se basant sur la fragmentation au niveau de la couche MAC 802.11. Pour cela, la taille des trames MAC est adaptée dynamiquement par la couche applicative en fonction des pertes qu'elle subit durant la transmission. Nous commençons cette section par détailler les différents niveaux de fragmentation existants : la fragmentation au niveau applicatif, la fragmentation au niveau réseau et la fragmentation au niveau MAC 802.11. Ensuite, nous présentons une étude comparative de ces différents niveaux de fragmentation (application, réseaux et MAC) afin de déterminer le meilleur niveau qui permet de réduire les pertes durant la transmission. À partir de cette étude, nous introduisons notre système de fragmentation 802.11 adaptative. Enfin, nous présentons les résultats des tests conduits pour évaluer notre proposition.

La deuxième contribution présentée dans ce chapitre permet aux flux vidéo d'avoir des opportunités de transmission (TXOP) adaptatives en fonction de leur débit, ceci en transmettant en rafale (burst) les trames MAC appartenant à la même image vidéo. La couche MAC se base sur des informations de la couche applicative pour déterminer les trames MAC appartenant à la même

image vidéo. Au début, nous présentons le principe de l'envoi en rafale d'un groupe de trames MAC et ses avantages. Ensuite, nous proposons une taille de groupe adaptative qui assure aux flux vidéo un accès au canal proportionnel à leur débit. Enfin, nous terminons cette partie en présentant une évaluation de notre contribution pour démontrer sa capacité à assurer d'avantage de bande passante aux flux vidéo.

5.2 La fragmentation adaptative au niveau MAC 802.11

5.2.1 Motivation

Dans la section 2.2.1.2, nous avons expliqué que le signal de données subit plusieurs dégradations lorsqu'il est transmis dans un canal sans fil. Ces dégradations provoquent deux types d'évanouissement du signal : l'évanouissement rapide et l'évanouissement lent (*fast and slow fading*). Dans cette section, nous nous intéressons à l'évanouissement rapide qui provoque la corruption des trames MAC durant la transmission. Les trames corrompues sont automatiquement supprimées au niveau de la couche MAC ou de la couche physique du récepteur durant la vérification des CRCs. La probabilité de corruption d'une trame dépend principalement de sa taille et du niveau d'interférence du canal de transmission. En effet, la transmission d'une grande trame nécessite plus de temps que la transmission d'une petite trame. Ceci fait augmenter le temps d'exposition aux interférences provoquant, ainsi, une augmentation de la probabilité de corruption.

Afin de réduire cette probabilité, une solution simple peut être utilisée, elle consiste à réduire la taille des trames transmises sur le canal sans fil. Ceci fait diminuer leur temps de transmission et, par la même, leur probabilité de corruption. De plus, en cas de corruption d'une trame de petite taille, la charge supplémentaire générée par la retransmission de cette trame est moins importante, comparée à la retransmission d'une trame de grande taille. Cependant, la diminution de la taille des trames augmente le pourcentage de l'*overhead* qui représente le rapport entre la charge utile d'une trame et les en-têtes ajoutés par les différentes couches de transmission.

Ainsi, la taille d'une trame représente un compromis entre le nombre de paquets perdus (paquets corrompus) et l'*overhead* introduit.

Dans le cadre de notre architecture, la réduction de la taille des trames peut être exploitée pour diminuer les pertes de paquet d'un flux vidéo transmis par le XLAG vers les stations mobiles lorsque le niveau d'interférence dans le canal de transmission est élevé. Toutefois, à travers les couches TCP/IP du XLAG, la réduction de la taille des trames peut s'effectuer sur différentes couches en se basant sur le service de fragmentation. Effectivement, ce service est présent au niveau de la couche application, réseau et MAC. Chaque couche peut fragmenter les paquets ou les données qu'elle reçoit suivant une taille configurable. Il est donc important de déterminer le meilleur niveau de fragmentation qui permet de réduire les pertes de paquets pour les flux vidéo. D'un autre côté, la fragmentation est utile lorsque le système enregistre des pertes de paquet provoquées par les interférences, mais la fragmentation est contraignante lorsqu'il n'y a pas de perte puisqu'elle introduit un *overhead* supplémentaire qui est inutile.

Dans la section suivante, nous évaluons les performances des différents niveaux de fragmentation pour la transmission d'un flux vidéo sur un réseau 802.11.

5.2.2 Comparaison entre les niveaux de fragmentation au niveau du XLAG

L'objectif de cette étude est de déterminer le meilleur niveau de fragmentation qui permet de réduire les pertes de paquets transmis dans un canal 802.11 caractérisé par un niveau élevé d'interférences. Pour cela, des tests ont été conduits sur une plateforme d'expérimentation, mais avant de présenter cette plateforme, nous détaillons ci-dessous le fonctionnement des différents niveaux de fragmentation : applicatif, réseau IP et MAC 802.11.

5.2.2.1 La fragmentation au niveau applicatif

La fragmentation au niveau applicatif est prise en charge par le protocole ALF (Application Level Framing) qui permet aux applications d'implémenter des fonctionnalités traditionnellement localisées sur la couche transport. Au niveau du XLAVS (Cross Layer Adaptive Video Streaming), le fonctionnement du protocole ALF relève des fonctionnalités du protocole RTP (voir section 3.4.4.2).

Parmi les fonctionnalités du protocole RTP, nous trouvons la paquetsisation pour découper les flux audio/vidéo en paquets RTP. L'IETF a publié plusieurs RFC qui définissent la procédure de découpage en paquets RTP pour différents codecs vidéo (voir section 2.2.5.3.1). La procédure comprend le découpage des images vidéo en paquets RTP et le renseignement de quelques champs de l'en-tête RTP (voir l'annexe A.4) en fonction de ce découpage. Les principaux champs qui sont renseignés durant cette procédure sont :

- Padding P (1bit): Informe sur la présence de bit de bourrage s'il est positionné à « 1 ».
- Marker M (1 bit): Son interprétation est définie par l'application. Généralement, il informe sur la fin d'une image s'il est positionné à « 1 ».
- Sequence number (16 bits): Donne le numéro de séquence des paquets RTP.
- TimeStamp (32 bits): Donne l'instant d'échantillonnage du premier octet dans le paquet RTP qui est identique à l'instant d'échantillonnage du premier octet de l'image vidéo. Cet instant doit être dérivé d'une horloge qui augmente de façon monotone et linéaire dans le temps.

La reconstruction des images vidéo se base sur ces champs d'informations. La taille d'un paquet RTP n'est pas fixe, elle peut varier suivant le type du flux. Les flux audio utilisent des paquets de petite taille (100 octets), par contre les flux vidéo utilisent une taille plus grande limitée par le MTU du réseau sous-jacent. Par exemple, la taille d'un paquet RTP est de 1450 octets pour un réseau Ethernet utilisant un MTU de 1500 octets.

Ainsi, la fragmentation au niveau applicatif correspond au découpage d'une image vidéo en plusieurs paquets RTP. La taille des paquets RTP peut être configurée suivant une valeur maximale *RTP_threshold*. Dans la couche transport, aucun regroupement n'est effectué, un paquet RTP est encapsulé dans un paquet UDP quelque soit sa taille.

5.2.2.2 La fragmentation au niveau réseau

Un service de fragmentation est proposé par le protocole IP pour découper les paquets provenant de la couche transport, précisément les paquets UDP, suivant le MTU du réseau. La fragmentation duplique l'en-tête IPv4 [32] de 20 octets sur chaque fragment en changeant les champs suivants :

- Total Length (16 bits) : Donne la taille du fragment IP.
- Flags (MF : More fragments) (1 bit) : Indique la présence d'autres fragments s'il est positionné à «1».
- Fragment Offset (13 bits) : Donne l'emplacement du fragment dans le datagramme d'origine.
- Header Checksum (16 bits) : Le code détecteur d'erreurs doit être recalculé suivant le nouvel en-tête du fragment.

En changeant la valeur du MTU, nous pouvons faire varier la taille des paquets IP qui sont directement encapsulés dans une trame LLC/MAC 802.11. Le protocole IP propose un service de réassemblage qui permet d'assembler un paquet IP à partir ses fragments. Cependant si un fragment est perdu tout le paquet IP est perdu.

5.2.2.3 La fragmentation au niveau MAC 802.11

Le standard 802.11 initial [3] propose un mécanisme de fragmentation/réassemblage au niveau MAC qui permet de réduire la taille des trames transmises par la couche physique 802.11. Le standard définit le « *fragmentationThreshold* » qui représente la taille maximale d'un MPDU (Mac Protocol Data Unit).

Donc, si la taille du MSDU (Mac Service Data Unit) fourni par la couche LLC (Logical Link Control) est supérieure au « *fragmentationThreshold* », le MSDU (la trame LLC) est fragmenté en plusieurs MPDUs avec des tailles inférieures ou égales au « *fragmentationThreshold* ». L'en-tête MAC de 34 octets (voir l'annexe A.2) est dupliqué sur chaque fragment en changeant les champs suivants :

- More Fragments (1 bits) : Indique la présence d'autres fragments s'il est positionné à « 1 ».
- Sequence Control (2 octets) : Il permet de réordonner les fragments et les trames. Il est composé de deux sous-champs : le numéro d'un fragment dans une trame et le numéro de séquence de la trame.
- CRC (Cyclic Redundancy Check) : Permet de vérifier l'intégrité du fragment.

La transmission des fragments est considérée dans le mécanisme d'accès DCF (Distributed Coordination Function) puisque tous les fragments appartenant à la même trame sont transmis en rafale (burst) une fois le canal de transmission détenu. La Figure 5-1 illustre ce principe appelé « *fragment burst* ». Au début, la station suit la procédure standard pour accéder au canal (écoute pendant DIFS et *Backoff time*). Une fois le canal détenu, la station transmet le premier fragment et attend durant un intervalle SIFS l'acquittement du fragment transmis par la station réceptrice. Pour garder l'accès au canal, la station enchaîne la transmission du deuxième fragment après un intervalle

d'attente SIFS. Cette opération est répétée jusqu'à la transmission de tous les fragments ou après les échecs de plusieurs retransmissions d'un fragment.

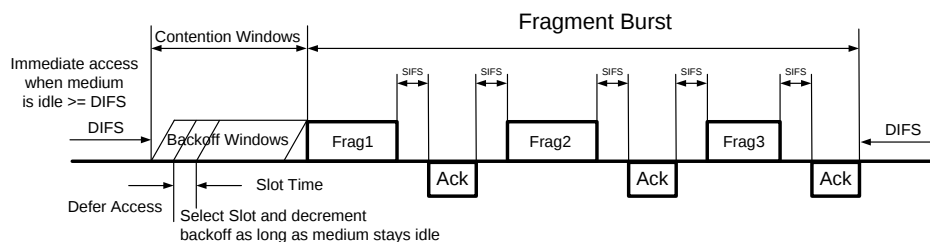


Figure 5-1 : L'envoi en rafale des fragments dans la fragmentation 802.11

5.2.2.4 La plateforme d'expérimentation

Pour comparer les performances des différents niveaux de fragmentation, nous avons exécuté des tests réels en nous basant sur une plateforme d'expérimentation illustrée dans la Figure 5-2. La plateforme comprend deux réseaux 802.11 en mode infrastructure « IPTV WLAN » et « IPERF WLAN ». Les deux réseaux utilisent le même canal de transmission numéro 1. En effet, Afin d'éviter les interférences inter-canaux, les réseaux 802.11 peuvent utiliser trois canaux suffisamment espacés : les canaux numéro 1, 6 et 11 (voir section 2.2.1.1). Le canal 11 étant exploité par le réseau WiFi du laboratoire, nous avons décidé d'utiliser le canal numéro 1 pour nos expérimentations.

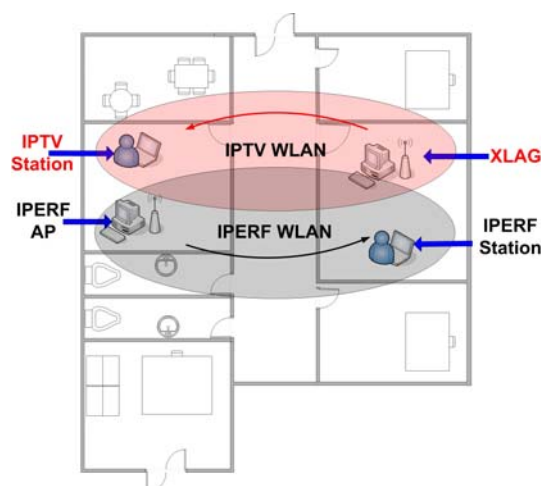


Figure 5-2 : La plateforme d'expérimentation

Le réseau « IPTV WLAN » est composé du XLAG et d'une station « IPTV STA ». D'une manière similaire, le réseau « IPERF WLAN » est composé d'un « IPERF AP » et d'une station « IPERF STA ». La disposition des deux réseaux est inversée comme l'illustre la Figure 5-2. Le XLAG est situé du côté de « IPERF STA » et le « IPERF AP » du côté de « IPTV STA ». Le Tableau 5-1 détaille les caractéristiques techniques des deux réseaux.

	IPTV WLAN	IPERF WLAN
Chipset WiFi/driver	AP/STA : Atheros / Madwifi	AP/STA : Atheros / Madwifi
Norme	802.11g	802.11b
Débit physique	36 Mbits/s	1Mbits/s
Nombre de retransmissions au niveau MAC	5 fois	5 fois
Taille de la file d'attente au niveau MAC	50 trames	50 trames

Tableau 5-1 : Les caractéristiques techniques des réseaux 802.11

Le réseau « IPTV WLAN » permet la transmission d'un flux vidéo du « XLAG » vers le « IPTV STA » et le réseau « IPERF WLAN » permet la transmission d'un flux UDP du « IPERF AP » vers le « IPERF STA ». Pour cela, le réseau « IPERF WLAN » utilise l'outil IPERF [211] qui permet de générer du trafic réseau UDP ou TCP en contrôlant plusieurs paramètres, comme le débit et la taille des paquets. Le trafic généré par le réseau « IPERF WLAN » sur le même canal de transmission provoque des interférences pour le réseau « IPTV WLAN ».

Afin de pouvoir comparer les résultats des tests, la même séquence vidéo est utilisée pour le streaming vidéo. Le Tableau 5-2 résume les caractéristiques de cette séquence vidéo.

Format de Codage	MPEG-4
La taille d'un GOP	12
Structure d'un GOP	IBBPBBPBBPBB
Résolution	CIF : 352x288 pixels
Nombre d'images par seconde	25 fps
Débit moyen/max/min	439/608/328 Kbits/s
Durée de la séquence	240 secondes

Tableau 5-2 : Les caractéristiques de la vidéo de référence

5.2.2.5 Tests et évaluations de performance

Les tests consistent à transmettre simultanément un flux vidéo dans le réseau « IPTV WLAN » et un flux UDP dans le réseau « IPERF WLAN ». Afin de démontrer l'avantage de la fragmentation et de comparer les différents niveaux de fragmentation détaillés dans les sous-sections précédentes, nous avons défini quatre scénarios d'expérimentations. Dans chaque scénario, le XLAG exécute un niveau de fragmentation qui change la taille des trames MAC pour le flux vidéo. Ces scénarios sont présentés ci-dessous :

- **No fragmentation :** Ce scénario utilise des paquets de taille standard sans aucune fragmentation. La taille d'une trame MAC est inférieure ou égale à 1512 octets
- **Application (APP) fragmentation :** Dans ce scénario, les paquets RTP sont réduits pour avoir une taille de trame MAC inférieure ou égale à 512 octets.

- **Network (NET) fragmentation** : Dans ce scénario, le MTU est réduit pour avoir une taille de trame MAC inférieure ou égale à 512 octets.
- **MAC (MAC) fragmentation** : Ce scénario utilise la fragmentation 802.11 au niveau MAC. La taille du fragment est limitée à 512 octets.

La durée d'une expérimentation est de 240 secondes. Durant l'intervalle [60s, 180s], le « IPERF AP » transmet vers le « IPERF STA » un trafic UDP avec un débit fixe de 300 Kbits/s. Les trames MAC dans le réseau « IPERF WLAN » possèdent une taille fixe de 1512 octets. Dans chaque scénario, nous avons exécuté dix tests. Nous nous intéressons pour chaque scénario au taux de perte et au débit vidéo réalisés.

Les Figures 5-3 et 5-4 illustrent un exemple des résultats obtenus tout au long d'un test dans chaque scénario. Elles donnent respectivement les taux de perte RTCP et le débit d'émission du flux vidéo au niveau MAC pour les différents scénarios. Dans la Figure 5-3, nous remarquons, clairement, l'effet des interférences introduit par le réseau « IPERF WLAN » dans l'intervalle [60s, 180s]. Dans la Figure 5-4, nous distinguons la différence du débit vidéo au niveau MAC pour les quatre scénarios.

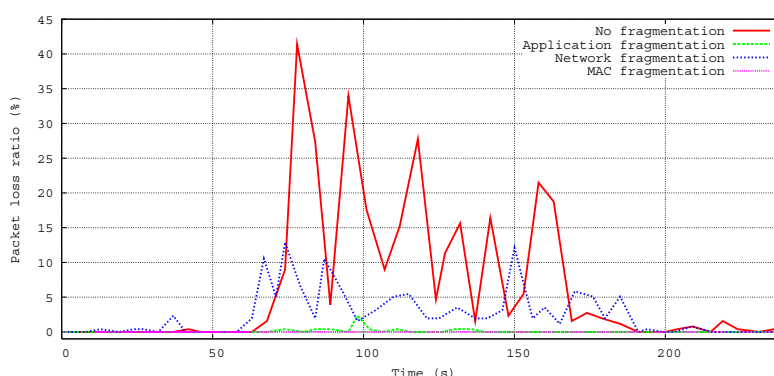


Figure 5-3 : Les taux de perte RTCP d'un test dans chaque scénario

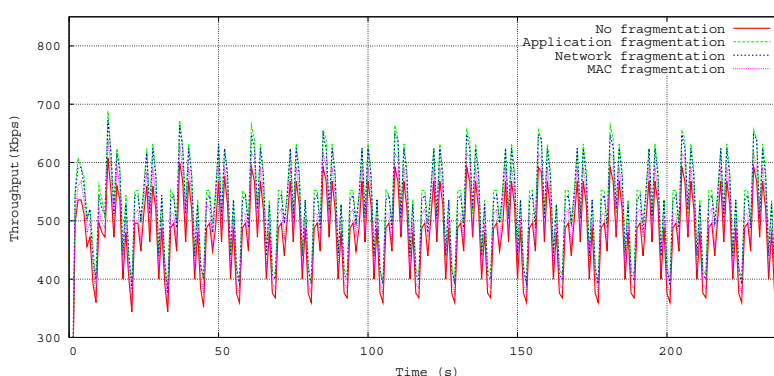


Figure 5-4 : Le débit d'émission du flux vidéo au niveau MAC d'un test dans chaque scénario

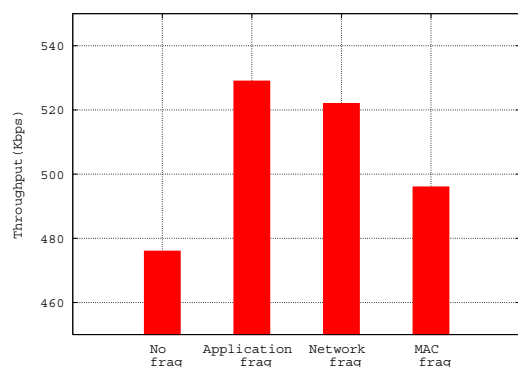


Figure 5-5 : Le débit moyen d'émission des dix tests dans chaque scénario

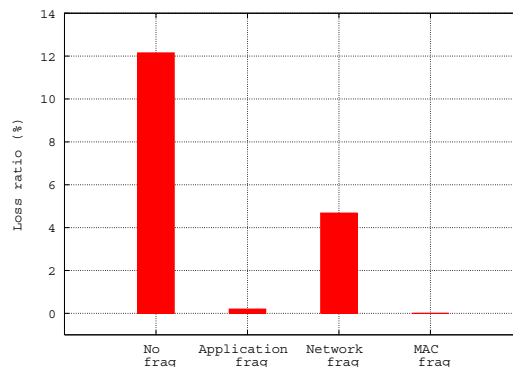


Figure 5-6 : Le taux moyen de perte des dix tests dans chaque scénario

Les Figures 5-5 et 5-6 donnent respectivement le taux moyen de perte et le débit moyen d'émission pour le flux vidéo au niveau MAC des dix tests dans chaque scénario. Ces moyennes ont été calculées durant la période d'interférence [60s, 180s].

Le taux élevé de perte, 12.1 %, est enregistré par le premier scénario « no fragmentation ». Les pertes diminuent avec le troisième scénario « Network fragmentation » pour atteindre une moyenne de 4.6 %. Cette moyenne atteint les 0.2 % dans le deuxième scénario « Application fragmentation » et elle est nulle dans le dernier scénario « MAC fragmentation » qui n'enregistre aucune perte durant cette période d'interférence.

Afin d'expliquer ces résultats, nous rappelons que les interférences provoquent la perte des trames MAC qui peuvent être supprimées à deux niveaux distincts :

- Au niveau de la couche MAC du récepteur : Cette suppression est due à l'altération de quelques bits dans la trame MAC durant la transmission, ce qui provoque sa suppression durant la vérification du champ CRC de la couche MAC ou de la couche Physique.
- Au niveau de la couche MAC de l'émetteur : Cette suppression est causée par une surcharge de la file d'attente MAC au niveau de l'émetteur. En effet, un niveau d'interférence élevé dans un canal peut être interprété dans l'algorithme d'accès DCF comme étant un canal occupé. Ceci fait augmenter le temps d'attente d'une trame à cause des attentes successives DIFS pour la libération du canal. Par conséquent, le taux d'occupation de la file d'attente au niveau MAC augmente et lorsque ce taux dépasse un certain seuil, les trames MAC sont supprimées.

Dans le premier test « No fragmentation », la taille des trames étant de 1512 octets, le taux de perte élevé (12.1%) des paquets peut être expliqué par la suppression conjointe des trames au niveau de la couche MAC de l'émetteur et du récepteur.

Par contre dans les autres scénarios « Application fragmentation », « Network fragmentation » et « MAC fragmentation » la même taille de trame MAC (512 octets) est utilisée mais les moyennes des taux de perte diffèrent. Nous commençons par expliquer la différence des taux de perte entre les deux scénarios « Application fragmentation » et « Network fragmentation ». Ensuite, nous expliquons pourquoi les pertes s'annulent dans le dernier scénario « MAC fragmentation »

Ces deux informations : (1) une taille de trame identique et (2) une moyenne des taux de perte différente pour le « Application fragmentation » et « Network fragmentation » permettent d'affirmer

que la différence dans les taux de perte n'est pas provoquée par la corruption des trames parce que la corruption d'une trame est en relation directe avec sa taille. Ainsi, la différence peut être provoquée uniquement par la suppression de ces trames au niveau de l'émetteur à cause de la surcharge de la file d'attente MAC. Les pertes sont plus élevées dans le scénario « Network fragmentation », comparé au scénario « Application fragmentation » pour deux raisons principales:

- La première raison se résume dans le fait que la fragmentation au niveau réseau « Network fragmentation » augmente le nombre des fragments IP qui sont transmis à la couche MAC. Cette dernière n'arrive pas à transmettre tous ces fragments qui sont stockés dans sa file d'attente. Puisque les files d'attente sont gérées en nombre de trames (50 trames dans notre plateforme de test) et non pas en nombre d'octets. La file d'attente est rapidement surchargée et supprime, par conséquent, des fragments IP. Ce cas de figure ne se produit pas dans le scénario « Application fragmentation » bien que le nombre de paquets RTP soit aussi important. Ces paquets RTP sont stockés dans la file d'attente de la couche IP qui a une taille plus grande. Ceci signifie que la charge que subisse la file d'attente MAC dans le scénario « Network fragmentation » est déplacée vers la file d'attente IP dans le scénario « Application fragmentation », cette dernière arrivant à absorber les paquets RTP.
- La deuxième raison se résume dans le fait que les taux de perte recueillis durant les tests sont les taux de perte des paquets RTP calculés au niveau du récepteur et envoyés dans les rapports RTCP vers l'émetteur. Dans le scénario « Network fragmentation » un paquet RTP peut être fragmenté en plusieurs fragments IP. La perte d'un de ces fragments provoque la perte de tout le paquet RTP. Par contre, dans le scénario « Application fragmentation » c'est la taille du paquet RTP qui est réduite. Ceci évite de détruire des paquets correctement reçus et permet, par la même, de réduire les taux de perte. Par exemple, si une image vidéo est découpée en 3 paquets RTP avec une taille de trame MAC de 1512 octets et 9 paquets avec une trame MAC de 512 octets, la perte d'un fragment dans le scénario « Network fragmentation » génère un taux de perte de 1/3, mais la perte d'un fragment dans le scénario « application fragmentation » génère un taux de perte de 1/9.

Les pertes s'annulent complètement dans le dernier scénario « MAC fragmentation » puisque la fragmentation 802.11 au niveau MAC se base sur le « *fragment burst* » qui transmet tous les fragments en une seule fois sans relâcher le contrôle du canal de transmission. Ceci réduit considérablement le temps d'attente des trames MAC dans la file d'attente, ce qui évite, par la même, sa surcharge et la suppression de ces trames.

Concernant le débit du flux vidéo au niveau de la couche MAC présenté dans la Figure 5-5. Les résultats sont compréhensibles. Plus le niveau de fragmentation est élevé, plus le débit du flux vidéo augmente. Ainsi, le débit est plus élevé dans le scénario « Application fragmentation » parce que les en-têtes RTP (12 octets), UDP (8 octets), IP (20 octets), LLC (8 octets), MAC (34 octets), PHY (24 octets) sont dupliqués dans chaque fragment d'image. Le débit le plus faible est réalisé par la fragmentation MAC puisque seule les en-têtes MAC et PHY sont dupliqués.

Cette étude démontre clairement que la fragmentation au niveau MAC 802.11 offre de meilleures performances de transmission puisque elle évite la suppression des trames MAC aussi

bien au niveau émetteur qu'au niveau récepteur. De plus, elle permet de réduire l'*overhead* des en-têtes parce qu'elle représente le dernier niveau d'encapsulation des données.

5.2.3 La fragmentation adaptative *Cross-layer* au niveau MAC 802.11

Dans cette section, nous détaillons la fragmentation adaptative qui peut être mise en œuvre dans notre nouveau système de streaming XLAVS au niveau des passerelles d'adaptation XLAG. Cette adaptation *Cross-layer* propose le contrôle dynamique de la fragmentation 802.11 par le module XLDP pour chaque station recevant un flux IPTV. Ceci signifie que le XLAG utilise, pour chaque station, une taille de trame spécifique pour lui transmettre le flux IPTV. Ce mécanisme d'adaptation détermine pour chaque station, la taille de la trame idéale qui réduit les pertes causées par les interférences tout en respectant un certain seuil d'*overhead*.

5.2.3.1 Déterminer la taille idéale d'une trame

L'un des défis majeurs dans le mécanisme de fragmentation est de déterminer la taille du fragment qui permet à la fois de minimiser les taux de perte et l'*overhead*. En effet, la taille d'un fragment représente un compromis entre ces deux critères de performance, puisque la diminution des taux de perte tend à faire diminuer la taille d'un fragment et la diminution de l'*overhead* tend à la faire augmenter.

Le calcul des taux de perte des trames durant la transmission est fonction du BER (Bit Error Rate) qui représente la probabilité de corruption d'un bit dans une trame. Cette corruption correspond au décodage incorrect d'un symbole au niveau du signal. Le taux de BER d'un canal de transmission dépend de plusieurs paramètres : la modulation du signal, la puissance du signal SSI (Signal Strength Indicator) et le rapport du signal par rapport aux interférences SNR (Signal-to-Noise Ratio). La Figure 5-7 fournie par les auteurs dans [212] représente la valeur théorique du BER en fonction du SNR pour différentes modulations. Il est important de noter que chaque station dans le réseau expérimente un BER indépendant des autres stations et qui dépend principalement de l'état du canal lors de la transmission.

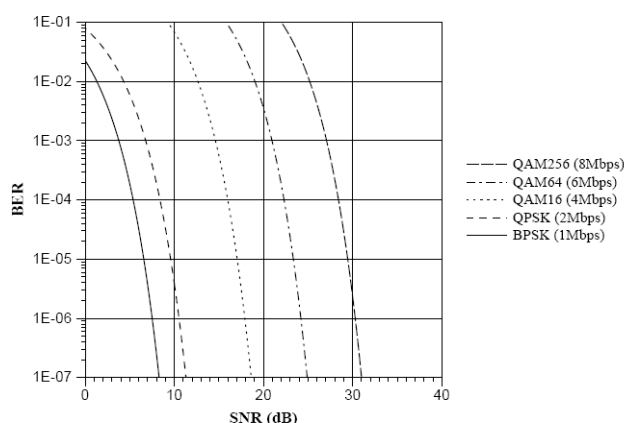


Figure 5-7 : Le BER en fonction du SNR

L'équation Eq. 5-1 donne la probabilité de perte d'une trame p en fonction de sa taille fs et en fonction du BER .

$$p = 1 - (1 - BER)^{fs} \quad \text{Eq. 5-1}$$

$$fs = ps + hs_{mac} + hs_{phy} \quad \text{Eq. 5-2}$$

BER (Bit Error Rate) : Le taux d'erreur bit

fs (frame size) : La taille d'une trame transmise sur le canal 802.11

ps (payload size) : La taille de la charge utile d'une trame

hs_{mac} : La taille d'en-tête de la couche MAC 802.11 (34 octets)

hs_{phy} : La taille d'en-tête de la couche physique 802.11 (24 octets)

À partir de cette équation, nous pouvons remarquer que pour un BER fixe, la probabilité de perte d'une trame est proportionnelle à sa taille. D'un autre côté, l'équation Eq. 5-3 donne le taux d'*overhead* R_{ovh} introduit par la couche MAC et la couche physique 802.11. L'équation Eq. 5-3 démontre que l'*overhead* R_{ovh} est inversement proportionnel à la taille d'une trame fs .

$$R_{ovh} = \frac{hs_{mac} + hs_{phy}}{fs} \quad \text{Eq. 5-3}$$

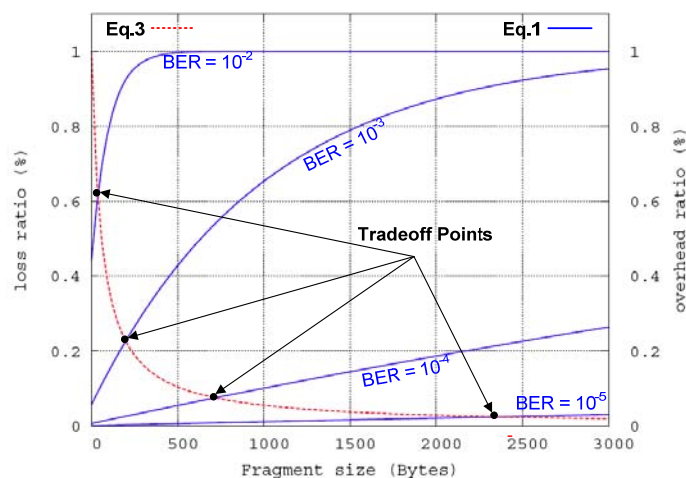


Figure 5-8 : Le taux de perte et le taux d'*overhead* en fonction de la taille d'une trame

Dans la Figure 5-8, nous représentons dans le même graphe les courbes de l'équation Eq. 5-1 (axe y gauche) et Eq. 5-3 (axe y droit). La courbe d'équation Eq. 5-1 (ligne bleue) donne les taux de perte en fonction de la taille des trames pour différentes valeurs de BER : 10^{-2} , 10^{-3} , 10^{-4} , 10^{-5} . La courbe d'équation Eq. 5-3 (ligne rouge) donne le pourcentage de l'*overhead* en fonction des tailles des trames. L'intersection entre les courbes des équations Eq. 5-1 et Eq. 5-3, donne les tailles des trames qui minimisent à la fois les taux de perte et le taux d'*overhead*.

Ainsi, à partir du niveau du BER durant les transmissions et en ayant comme objectif de minimiser l'*overhead*, nous pouvons adapter dynamiquement la taille des trames en utilisant la fragmentation 802.11 au niveau MAC.

Cependant, le principal problème dans un tel système d'adaptation et de déterminer le taux du BER en temps réel en se basant sur les trames reçues. En effet, le taux du BER ne peut être déterminé avec exactitude au niveau du récepteur puisque ce dernier peut uniquement savoir que la trame est corrompue en vérifiant le CRC de l'en-tête de la couche physique et le CRC de la trame complète dans l'en-tête de la couche MAC. Ces CRCs ne permettent pas de savoir le nombre d'erreurs dans une trame.

Cependant, la dérivation du BER à partir d'autres paramètres est possible. Nous pouvons distinguer deux méthodes de dérivation :

- **L'estimation optimiste** : Dans cette méthode, le BER est estimé à partir du SNR et du type de modulation, comme cela est présenté dans la Figure 5-7. Cependant, le niveau d'interférence n'est pas considéré, ce qui engendre une sous estimation du BER réel.
- **L'estimation pessimiste** : Dans cette méthode, le BER est estimé à partir du FER (Frame Error Rate). Ceci engendre une surestimation du BER puisque lorsqu'un paquet est supprimé cela ne signifie pas que tous les bits de la trame sont corrompus.

Ces deux estimations ne fournissent pas un taux exact du BER qui permette de déterminer la taille idéale d'une trame à partir des équations Eq. 5-1 et Eq. 5-3.

5.2.3.2 Déterminer la limite inférieure pour la taille d'une trame

À défaut de trouver la taille idéale des trames, nous proposons de trouver pour chaque station une limite inférieure pour cette taille qui respecte un certain taux d'*overhead*. Effectivement, le taux d'*overhead* Ri_{ovh} maximum qui peut être exploité par une station au niveau du XLAG représente le reste du rapport entre le débit du flux destiné à cette station, au niveau de la couche LLC (Logical Link Control) Ri_{llc} , et le débit physique Ri_{phy} utilisé avec la station. Ceci est illustré par l'équation Eq. 5-4 :

$$Ri_{ovh} = 1 - \frac{Ri_{llc}}{Ri_{phy}} \quad \text{Eq. 5-4}$$

Cependant, dans notre architecture, le XLAG transmet des flux vers plusieurs stations mobiles avec des débits physiques r_i spécifiques pour chaque station (voir section 4.2.4.1). Dans cette configuration, nous avons proposé dans la section 4.2.4.1 une équation pour calculer le débit physique effectif pour chaque station Ri_{phy} à partir de leurs débits physiques. L'équation est reprise ci-dessous :

$$Ri_{phy} = \frac{1}{\sum_{i=1}^{Nb_{sta}} \frac{1}{r_i}} \quad \text{Eq. 5-5}$$

Nb_{sta} : Le nombre de stations connectées au XLAG.

r_i : Le débit physique utilisé pour chaque station.

L'équation Eq. 5-6 est obtenue en utilisant l'équation Eq. 5-5 dans l'équation Eq. 5-4

$$Ri_{ovh} = 1 - Ri_{llc} \sum_{i=1}^{Nb_{sta}} \frac{1}{r_i} \quad \text{Eq. 5-6}$$

Ainsi, la taille minimum d'une trame $fsi_{threshold}$ pour chaque station i peut être calculée par l'équation Eq. 5-7, déduite des équations Eq. 5-3 et Eq. 5-6

$$fsi_{threshold} = \frac{hs_{mac} + hs_{phy}}{1 - Ri_{llc} \sum_{i=1}^{Nb_{sig}} \frac{1}{r_i}} \quad \text{Eq. 5-7}$$

Exemple numérique :

Nous supposons que trois stations sont connectées au XLAG : station 1 (1Mbits/s), station 2 (1 Mbits/s) et station 3 (5 Mbits/s). Nous supposons qu'un flux vidéo est transmis à la station 1. Le débit du flux vidéo au niveau LLC est égal à 450 Kbits/s. En utilisant l'équation Eq.5-7, le $fsi_{threshold}$ de la station 1 est égal à 1000 octets, c'est-à-dire, que la station 1 ne peut pas utiliser une taille de trame inférieure à 1000 octets, sinon l'*overhead* introduit par la fragmentation fera déborder le débit physique disponible.

Dans le cas où le débit physique de la station 3 passe à 54 Mbits/s, le $fsi_{threshold}$ de la station 1 passe à 420 octets puisque le débit physique total a augmenté, ce qui fait diminuer la taille minimale d'un fragment.

5.2.3.3 Le fonctionnement de la fragmentation adaptative au niveau du XLAVS

Dans cette section, nous détaillons comment la limite inférieure pour la taille d'une trame peut être utilisée dans un système de fragmentation adaptative. Cette fragmentation 802.11 adaptative est exécutée durant la transmission des flux TV vers les stations. La Figure 5-9 illustre les échanges d'informations nécessaires pour le fonctionnement de la fragmentation adaptative.

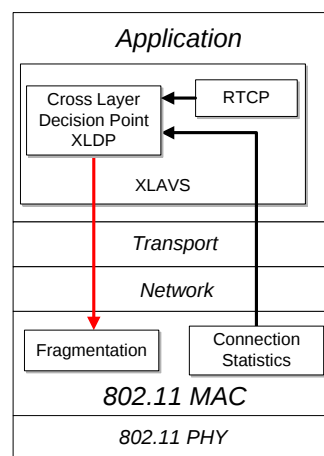


Figure 5-9 : L'échange d'informations pour la fragmentation 802.11 Adaptative

Le module XLDP récupère des paramètres de monitoring au niveau applicatif et au niveau MAC. Ces paramètres sont donnés ci-dessous :

- **Au niveau MAC :**
 - Le nombre de stations qui sont connectées au XLAG.
 - Le débit physique de chaque station.
 - Le débit de chaque flux vidéo au niveau de la couche LLC.

- **Au niveau applicatif :**

- Le taux de perte RTCP de chaque session IPTV ouverte au niveau du XLDP.

À partir des paramètres récupérés au niveau MAC, le XLDP calcule pour chaque station une limite inférieure pour la taille des trames suivant la procédure détaillée dans la section précédente. Ce calcul est répété d'une manière itérative afin de mettre à jour les limites en cas de :

- Connexion et/ou départ des stations.
- Changement des débits physiques des stations.
- Changement des débits des flux vidéo.

Le changement de la taille des trames pour chaque station se base principalement sur les taux de perte RTCP transmis par la station et sur des politiques d'adaptation qui fournissent une correspondance entre les taux de perte et la taille du fragment. Le fonctionnement basique de ces politiques consiste à faire baisser la taille des trames lorsque le taux de perte augmente et à augmenter la taille des trames lorsque le taux de perte diminue. Cependant, la taille du fragment choisie doit toujours respecter la limite inférieure calculée par le module XLDP.

L'activation et la désactivation de la fragmentation ainsi que la taille du fragment sont signalés par le module XLDP à la couche MAC en utilisant un appel de fonction interne.

Afin d'éviter les oscillations de la taille des fragments engendrées par une oscillation des taux de perte RTCP, nous utilisons une moyenne pondérée dans le temps afin de lisser (smooth) le niveau des taux de perte tout en donnant une importance aux taux de perte récents. Cette moyenne wp est calculée par l'équation Eq. 5-8, la valeur λ représente le facteur de pondération.

$$wp \leftarrow (1 - \lambda) \times p_{rtcp} + \lambda \times wp \quad \text{Eq. 5-8}$$

5.2.3.4 Tests et évaluations de performance

Dans cette section, nous présentons une évaluation de performance pour la fragmentation 802.11 adaptative. L'évaluation est limitée à une station pour démontrer uniquement l'importance d'une fragmentation adaptative durant les périodes d'interférences. L'évaluation se base sur des tests conduits sur la plateforme d'expérimentation présentée dans la section 5.2.2.4. Les tests consistent toujours à transmettre simultanément un flux vidéo dans le réseau « IPTV WLAN » et un flux UDP dans le réseau « IPERF WLAN » (voir Figure 5-2).

L'implémentation de la fragmentation adaptative se base sur l'intégration au niveau du XLDP d'un appel de fonction système fourni par le pilote libre *MadWiFi*. Cette fonction représente une interface de configuration qui permet de contrôler dynamiquement la fragmentation au niveau MAC ainsi que la taille des fragments.

Nous avons défini deux scénarios de tests pour comparer la fragmentation adaptative proposée avec un système conventionnel sans aucune adaptation. Dans chaque scénario, nous exécutons dix tests. Les deux scénarios sont détaillés ci-dessous :

- **Scénario 1 :** Dans ce scénario, le XLAVS n'exécute aucune adaptation. La fragmentation 802.11 est désactivée. La taille d'une trame MAC est inférieure ou égale à 1512 octets.

- **Scenario 2 :** Dans ce scenario, le XLAVS exécute une fragmentation adaptative en utilisant la politique d'adaptation détaillée dans le Tableau 5-3.

Taux de perte pondérés (%) ($\lambda = 0.2$)	Taille des trames (octets)
= 0	Aucune fragmentation ≤ 1512
< 5	≤ 1024
≥ 5	≤ 512

Tableau 5-3 : La politique d'adaptation pour la fragmentation adaptative

La durée d'un test est de 240s. Les caractéristiques techniques de la séquence vidéo sont résumées dans le Tableau 5-2. Les deux scénarios sont sujets au même modèle d'interférence présenté ci-dessous :

- [0s, 60s] : Aucune interférence ; aucun trafic UDP.
- [60s, 100s] : Avec interférences ; trafic UDP 100 Kbits/s.
- [100s, 140s] : Avec interférences ; trafic UDP 200 Kbits/s.
- [140s, 180s] : Avec interférences ; trafic UDP 300 Kbits/s.
- [180s, 240s] : Aucune interférence ; aucun trafic UDP.

La Figure 5-10 et la Figure 5-11 donnent respectivement les taux de perte RTCP et le débit d'émission vidéo au niveau MAC pour un test dans chaque scénario. La variation de la taille des trames pour le deuxième scénario est donnée dans la Figure 5-12. Ces résultats montrent la réactivité du système d'adaptation.

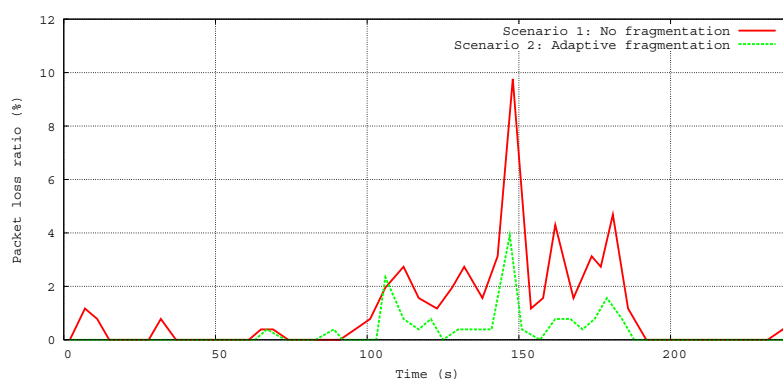


Figure 5-10 : Les taux de perte RTCP d'un test dans les scénarios 1 et 2

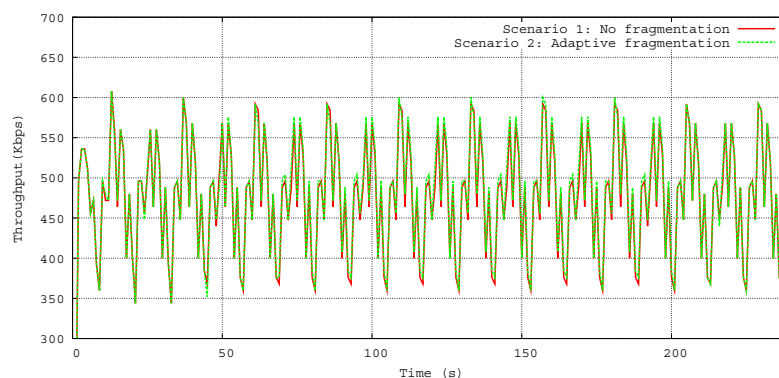


Figure 5-11 : Le débit d'émission du flux vidéo au niveau MAC d'un test dans les scénarios 1 et 2

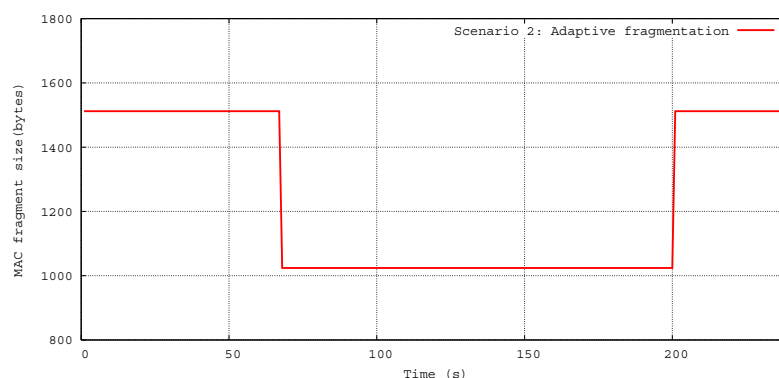


Figure 5-12 : La variation de la taille des trames d'un test dans le scénario 2

Dans la Figure 5-10, nous remarquons clairement la différence entre les deux scénarios concernant les taux de perte. Durant la période d'interférence [60s, 180s], le deuxième scénario utilisant la fragmentation adaptative déclenche dynamiquement la fragmentation en réduisant la taille des trames de 1512 octets à 1024 octets, comme illustré par la Figure 5-12. Ceci réduit considérablement les taux de perte et augmente sensiblement le débit du flux vidéo. À la fin de la période d'interférence]180s, 240s], les taux de perte baissent, Dans le scénario 2 la fragmentation 802.11 est désactivée et la taille des trames est réinitialisée à 1512 octets.

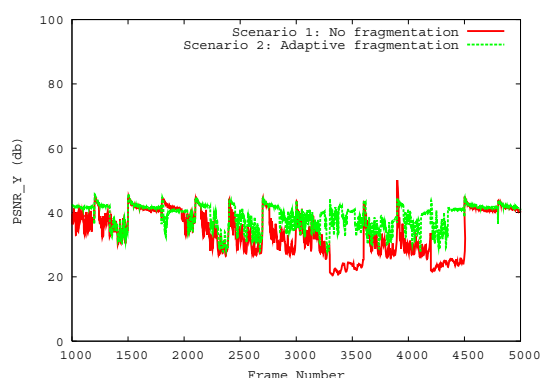


Figure 5-13 : PSNR pour les scénarios 1 et 2 entre l'image #1000 et l'image #5000

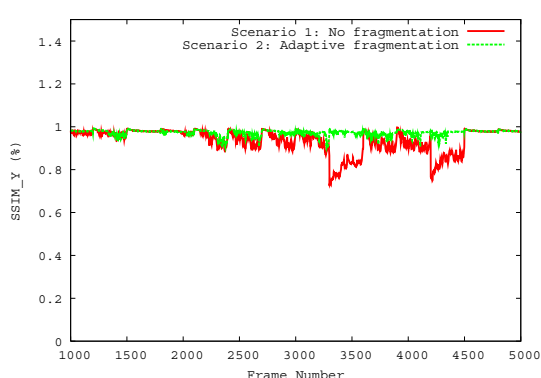


Figure 5-14 : SSIM pour les scénarios 1 et 2 entre l'image #1000 et l'image #5000

Pour mesurer la qualité objective des vidéos reçues par les stations dans les deux scénarios, nous avons calculé durant l'intervalle [40s, 200s] le PSNR et le SSIM donnés respectivement dans la Figure 5-13 et la Figure 5-14. Nous distinguons clairement la différence entre le premier et le deuxième scénario. La fragmentation adaptative dans le scénario 2 réduit les pertes des paquets et augmente, par la même, la qualité de la vidéo reçue par la station.

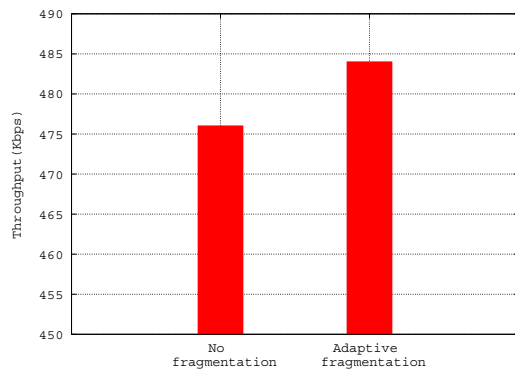


Figure 5-15 : Le débit moyen d'émission des dix tests dans les scénarios 1 et 2

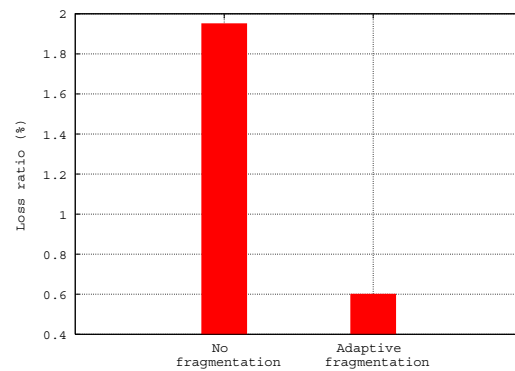


Figure 5-16 : Le taux moyen de perte des dix tests dans les scénarios 1 et 2

Les moyennes des taux de perte et du débit d'émission vidéo obtenus dans les dix tests de chaque scénario sont données respectivement dans la Figure 5-15 et la Figure 5-16. Les moyennes sont calculées durant la période d'interférence. Ces résultats illustrent la capacité de la fragmentation adaptative à diminuer les pertes contre une augmentation sensible du débit d'émission du flux vidéo causée par une augmentation de l'*overhead*.

5.3 Le groupage de trames basé sur l'image vidéo

5.3.1 Motivation

L'une des caractéristiques principales des flux vidéo est leur débit relativement élevé. Nous avons présenté dans la première partie du chapitre précédent une technique qui permet de varier le débit d'un flux vidéo dynamiquement durant la transmission, à savoir la variation de la résolution SNR. Toutefois, cette technique ne permet pas de garantir au flux vidéo un débit plus élevé au niveau du lien d'accès. Dans le cas où le flux vidéo est concurrencé par d'autres flux asynchrones UDP ou TCP, le débit du flux vidéo au niveau du lien d'accès 802.11 sera sérieusement limité à cause du partage équitable du lien d'accès garanti par le mécanisme d'accès DCF.

Afin de garantir plus de débit au niveau du lien d'accès, nous avons vu dans la section 2.2.2.4 que le standard IEEE 802.11e propose l'utilisation du TXOP qui autorise une station à transmettre plusieurs trames successivement sans relâcher l'accès au canal durant un intervalle de temps, appelé opportunité de transmission. Dans l'EDCA, l'intervalle de temps TXOP est fixe pour chaque catégorie d'accès et dans le HCCA, l'intervalle TXOP est spécifique pour chaque station.

Au niveau de notre architecture, le XLAG peut exploiter le TXOP pour garantir plus de bande passante au niveau du lien d'accès 802.11 pour les flux vidéo. Toutefois, le XLAG transmet des vidéos avec des débits différents grâce au module de transcodage. Ainsi, le besoin en débit des flux vidéo est différent et, par conséquent, l'intervalle d'un TXOP pour un flux vidéo doit différer aussi.

Dans cette section, nous proposons une technique pour assurer une certaine correspondance entre le débit vidéo et le nombre de trames envoyées durant un intervalle TXOP. Cette correspondance est dynamique et se base sur le principe du *Cross-layer* qui permet à la couche MAC 802.11 d'accéder à des informations de niveau applicatif. Mais avant de détailler notre contribution, nous présentons dans la section suivante le principe du groupage de trames au niveau MAC qui est à la base des opportunités de transmission TXOP.

5.3.2 Le principe du groupage de trames (frame grouping)

La couche physique et la couche MAC 802.11 introduisent un temps supplémentaire pour la transmission d'une trame de données. Au niveau de la couche physique, ce temps correspond à la transmission de l'en-tête et du préambule physique (PLCP preamble et PLCP header). Tandis qu'au niveau MAC, ce temps correspond aux intervalles de temps IFS et au mécanisme de *Backoff* introduits par le DCF afin d'assurer un accès équitable au canal.

Ainsi, ce temps, appelé t_{overhead} , est calculé par l'équation Eq. 5-9. Il comprend l'intervalle DIFS t_{DIFS} , le temps du backoff t_{backoff} , le temps de transmission de l'en-tête physique de la trame de données $t_{\text{phy_header}}$, le temps de transmission de l'en-tête MAC $t_{\text{MAC_header}}$, l'intervalle SIFS t_{SIFS} , le temps de réception de l'en-tête physique de l'acquittement $t_{\text{ack_PHY_header}}$ et le temps de réception de l'acquittement t_{ack} . Il est important de noter que les trames RTS/CTS ne sont pas considérées ici.

$$t_{\text{overhead}} = t_{\text{DIFS}} + t_{\text{backoff}} + t_{\text{PHY_header}} + t_{\text{MAC_header}} + t_{\text{SIFS}} + t_{\text{ack_PHY_header}} + t_{\text{ack}} \quad \text{Eq. 5-9}$$

Pour un certain débit physique, la valeur du $t_{overhead}$ est fonction du t_{DIFS} et du $t_{backoff}$. La valeur du t_{DIFS} dépend de l'état du canal, tandis que, la valeur du $t_{backoff}$ varie aléatoirement et dépend du succès ou de l'échec des transmissions précédentes.

Afin de réduire le $t_{overhead}$ qui est à l'origine de la dégradation des performances de transmission (débit, délai de transmission, gigue) au niveau du lien d'accès 802.11, plusieurs mécanismes ont été proposés au niveau MAC. Parmi ces mécanismes, nous trouvons la concaténation des trames [213], le *piggyback* des trames [213] et le groupage des trames [214].

La concaténation des trames consiste à regrouper au niveau MAC les MSDUs qui possèdent le même destinataire dans un seul MPDU en introduisant un petit en-tête de concaténation. Ceci permet de partager le $t_{overhead}$ entre plusieurs MSDUs. Cependant, cette technique est limitée par le fait que la corruption des trames dépend principalement de leur taille (voir section 5.2.3.1). Donc, le nombre de MSDUs concaténés reste limité et dépend de l'état du canal.

Le deuxième mécanisme, le *piggyback* des trames, est très utile dans une communication bidirectionnelle active entre deux stations. Il permet à la station réceptrice d'envoyer ses trames de données avec les acquittements qu'elle transmet à l'émetteur. Ce dernier acquitte la trame reçue. Ceci signifie que la station réceptrice économise le temps d'accès au canal $t_{DIFS} + t_{backoff}$. Cette technique n'offre aucun avantage dans une communication unidirectionnelle, par exemple, le streaming vidéo.

Dans le cadre de cette contribution, nous nous intéressons particulièrement au dernier mécanisme, à savoir le groupage de trames. Ce mécanisme permet de partager le temps d'accès au canal ($t_{DIFS} + t_{backoff}$) entre plusieurs trames possédant la même destination. C'est-à-dire que les trames sont transmises en rafale sans relâcher le canal de transmission et chaque trame est acquittée de la part du récepteur comme l'illustre la Figure 5-17.

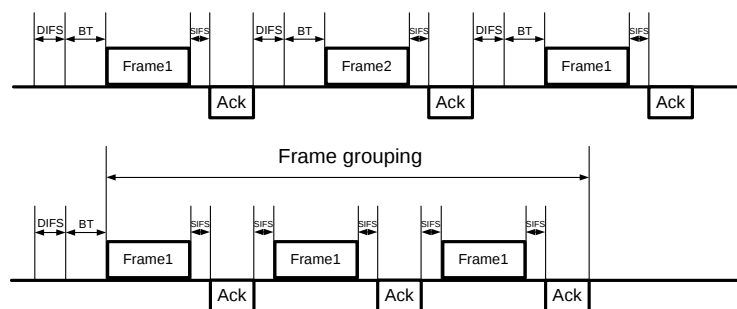


Figure 5-17 : Le groupage de trames comparé à une transmission standard

Ce mécanisme est inspiré principalement du « *fragment burst* » proposé par le standard IEEE 802.11 (voir section 5.2.2.3) et il a été proposé pour augmenter les performances des flux constitués de petites trames. Ce type de flux est pénalisé par le temps d'accès au canal et le regroupement des trames permet d'y remédier. Ce mécanisme possède plusieurs avantages :

- Il évite la corruption des trames en évitant la transmission de trames de grande taille comme cela est proposé dans la concaténation.
- Il permet de partager le temps d'accès au canal entre plusieurs trames.

- Il augmente le débit d'un flux au niveau du lien d'accès puisque le flux est autorisé à transmettre plus de données lorsqu'il a accès au canal.

Le principe de fonctionnement du groupage de trames est repris dans le standard IEEE 802.11e sous le nom d'opportunité de transmissions (TXOP). Dans le mécanisme d'accès EDCA, le TXOP est spécifique à une catégorie d'accès AC. Le Tableau 5-4 donne les valeurs par défaut du TXOP proposées par le standard IEEE 802.11e [31].

classe		AC_BK	AC_BE	AC_VI	AC_VO
TXOP (ms)	FHSS	0	0	6.016 ms	3.264 ms
	DSSS	0	0	3.008 ms	1.504 ms

Tableau 5-4 : Les valeurs par défaut du TXOP dans le IEEE 802.11e

Par contre, dans le mécanisme HCCA, le TXOP est spécifié pour chaque station (QSTA) suivant le TSPEC qu'elle transmet au QAP (voir section 2.2.2.4.2).

Dans la section suivante, nous détaillons notre contribution qui adapte le principe du groupage de trames au besoin de notre nouveau système de streaming XLAVS.

5.3.3 Le principe du groupage de trames MAC basé sur l'image vidéo

Pour la transmission des flux vidéo dans les réseaux d'accès 802.11, dans le cadre de notre architecture, le groupage des trames au niveau MAC peut être très utile pour assurer aux flux vidéo plus de bande passante sur le lien d'accès par rapport à d'autres flux. En effet, les flux vidéo sont caractérisés par un débit élevé qui doit être satisfait au niveau du lien d'accès. Sinon, les trames vidéo seront supprimées à cause de la saturation des files d'attente au niveau de la couche MAC. Cette suppression de trames se traduit par des pertes de paquets qui dégradent sérieusement la qualité vidéo perçue par une station.

La mise en œuvre du groupage des trames au niveau du XLAG doit prendre en considération les caractéristiques du système XLAVS qui a la possibilité de transmettre des flux vidéo avec des débits différents grâce à son module de transcodage. Par conséquent, les besoins en bande passante sont différents pour chaque flux vidéo.

Ainsi, le défi majeur dans l'utilisation du groupage de trames est la délimitation de l'intervalle de temps durant lequel les trames sont transmises en rafales. Cet intervalle doit satisfaire deux exigences principales :

1. Avoir une corrélation avec le débit du flux vidéo.
2. Éviter la monopolisation du canal de transmission.

Pour relever ce défi, nous proposons, tout d'abord, de rendre le mécanisme du groupage de trames indépendant du temps. C'est-à-dire qu'au lieu de délimiter un intervalle de temps durant lequel les trames sont transmises en rafale, nous proposons de délimiter un groupe pour lequel il faut définir une taille en nombre de trames. Ceci est motivé par le besoin de supprimer la dépendance entre le groupage des trames et le débit physique puisque notre objectif est de rendre le groupage des trames dépendant du débit du flux vidéo.

Pour illustrer la dépendance entre le groupage des trames et le débit physique, nous proposons un exemple numérique concret. Par exemple dans le standard IEEE 802.11e, la catégorie d'accès AC_VI définie spécifiquement pour les flux vidéo propose un intervalle TXOP par défaut de 3.008 ms. Avec un débit physique égal à 2Mbits/s et une trame standard égale à 1512 octets, 4 trames peuvent être transmises en rafale. Par contre, avec un débit physique de 11 Mbits/s, 22 trames sont transmises en rafale. Donc en délimitant un groupe par le nombre de trames, nous annulons cette dépendance.

Deuxièmement, nous proposons une taille de groupe dynamique pour chaque flux vidéo. Ce groupe est basé sur une délimitation virtuelle des trames au niveau applicatif qui correspond à une image vidéo. Cette dernière est découpée au niveau ALF (couche applicative) en plusieurs paquets RTP en fonction du MTU du réseau. Le nombre de paquets RTP générés à partir d'une image vidéo est fonction du débit vidéo v_{rate} et du type de l'image pour un codage MPEG p_{type} comme ceci est illustré par l'équation Eq. 5-10.

$$N_{RTP} \propto (p_{type}, v_{rate}) \quad \text{Eq. 5-10}$$

La règle générale stipule que le nombre de paquets RTP est proportionnel au débit vidéo et, pour n'importe quel débit vidéo, la taille d'une image I est plus grande que la taille d'une image P qui est, à son tour, plus grande que la taille d'une image B. Ceci est illustré par le Tableau 5-5 qui donne le nombre de paquets RTP moyen (1450 octets) par type d'image pour différents débits d'un même flux vidéo codé en MPEG-4.

Débit vidéo	Image I	Image P	Image B
500 Kbits/s	7	2	1
1 Mbits/s	12	5	2
1.5 Mbits	16	7	3
2 Mbits/s	20	10	5
2.5 Mbits/s	20	12	6
3 Mbit/s	24	14	7

Tableau 5-5 : Le nombre de paquets RTP moyen par image vidéo

Ainsi en définissant le groupe de trames comme étant une image vidéo, nous satisfaisons la première exigence qui rend dépendant la taille d'un groupe au débit du flux vidéo. De plus, nous assurons une priorité d'occupation du lien d'accès en relation avec la hiérarchie temporelle basée sur le type d'image I, P, B.

Concernant la seconde exigence, à savoir éviter au flux vidéo de monopoliser le canal, ceci ne risque pas de se produire puisque le canal est relâché après la transmission d'une image complète. Pour transmettre l'image suivante, il faut reprendre l'accès au canal suivant le mécanisme d'accès DCF.

La Figure 5-18 illustre notre proposition, appelée VFG (Video Frame grouping) pour le groupage des images vidéo au niveau MAC 802.11.

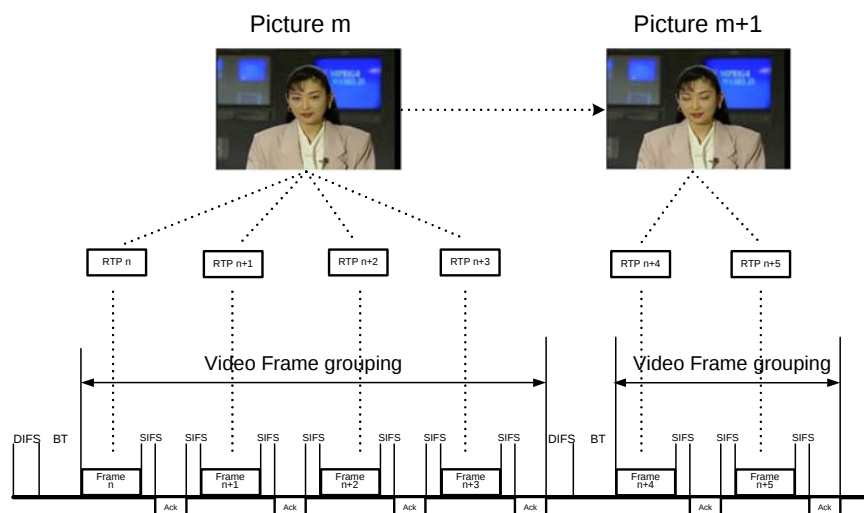


Figure 5-18 : Le groupage des images vidéo

5.3.3.1 Mise en œuvre du VFG au niveau du XLAG

Afin de permettre à la couche MAC de mettre en œuvre le groupage des images vidéo, nous avons introduit un nouveau module au niveau de la couche MAC, appelé VFG. La tâche principale de ce module est de maintenir une connaissance, au niveau MAC, des sessions vidéo au niveau du XLAVS afin de pouvoir identifier les trames appartenant à chaque session mais aussi les trames appartenant à la même image.

L'identification se base sur des informations « in-band » présentes dans les trames MSDU qui sont reçues par la couche MAC. Ces informations correspondent principalement à deux champs de l'en-tête du protocole RTP décrits ci-dessous :

- La source de synchronisation SSRC (Synchronisation Source) : elle permet d'identifier la session vidéo dans chaque session IPTV. Ainsi, les trames qui possèdent un SSRC identique appartiennent à la même session vidéo, puisque aucun mixeur n'est utilisé au niveau du XLAVS.
- Le temps d'échantillonnage (Timestamp) : Il permet d'identifier les paquets RTP appartenant à la même image, puisque l'instant d'échantillonnage du premier octet dans un paquet RTP représente l'instant d'échantillonnage de toute l'image auquel appartient ce paquet.

Ainsi, pour chaque session IPTV créée, ou détruite, au niveau du XLAVS, le module XLDP exécute des appels de fonction qui permettent de signaler explicitement ces deux événements. Les deux appels de fonctions sont les suivants :

- *session_create(ssrc, initial_ts)* : Elle permet de signaler la création d'une session vidéo. Elle possède le SSRC de la session créée et le timestamp initial comme paramètre.
- *session_destroy(ssrc)* : Elle permet de signaler la destruction d'une session vidéo identifiée par son SSRC.

Le module VFG gère à son niveau une table pour sauvegarder toutes les sessions vidéo actives, ainsi que le SSRC et le timestamp de chaque session. Pour chaque entrée dans la table, le

module gère aussi un filtre qui permet de regrouper les trames par session et par image afin d'exécuter le groupage des images vidéo.

5.3.4 Tests et évaluations de performance

L'évaluation du groupage des images vidéo au niveau MAC se base sur des tests réels réalisés sur une plateforme d'expérimentation. Pour cela, le module VFG a été implémenté dans le pilote de carte WiFi libre *MadWifi*. L'implémentation comprend la définition des appels de fonction « *session_create* » et « *session_destroy* », la création de la table qui gère les sessions vidéo actives et la mise en place des filtres qui permettent d'accéder à l'en-tête RTP d'une trame pour exécuter le groupage des trames.

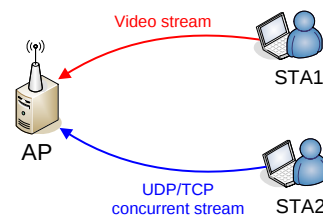


Figure 5-19 : La plateforme d'expérimentation pour le groupage des images vidéo

La plateforme d'expérimentation est illustrée dans la Figure 5-19, elle est constituée d'un réseau 802.11 en mode infrastructure. Le réseau comprend trois machines : un point d'accès (AP) et deux stations (STA1 et STA2). Le mécanisme VFG est implémenté au niveau de la STA1 qui transmet un flux vidéo vers l'AP, tandis que la STA2 transmet vers l'AP un flux concurrent UDP/TCP. Ainsi, les deux stations STA1 et STA2 doivent accéder simultanément au canal pour transmettre leur flux. Ceci crée une concurrence entre les deux stations.

Le plan des tests est découpé en deux parties. Dans la première partie, nous nous intéressons à la capacité de notre mécanisme à s'adapter au débit vidéo et nous comparons les performances du groupage adaptatif VFG par rapport à un groupage fixe. La deuxième partie des tests a pour objectif de montrer la capacité du mécanisme à assurer une allocation de bande passante au niveau du lien, quelque soit le débit du flux concurrent.

5.3.4.1 Évaluation du VFG par rapport à un groupage fixe

Pour démontrer l'avantage du VFG, nous avons conduit des tests avec deux flux vidéo possédant des débits différents afin d'avoir un nombre de paquets par image différent. Le Tableau 5-6 résume les caractéristiques techniques des deux flux vidéo utilisés : vidéo1 et vidéo2.

	Vidéo1	Vidéo2
Format de Codage	MPEG-4	MPEG-4
Résolution	CIF : 352x288 pixels	CIF : 352x288 pixels
Nombre d'images par seconde	25 fps	25 fps
Débit moyen/max/min	860/1064/744 Kbits/s	1252/1472/1120 Kbits/s
Durée de la séquence	120 secondes	120 secondes
La taille d'un GOP	12	12
Structure d'un GOP	IBBPBBPBBPBB	IBBPBBPBBPBB
Nombre moyen de paquets RTP par image I (maximum 1450 octets)	11	14
Nombre moyen de paquets RTP par image P (maximum 1450 octets)	4	6
Nombre moyen de paquets RTP par image B (maximum 1450 octets)	2	3

Tableau 5-6 : les caractéristiques des flux vidéo

À partir du nombre moyen de paquets RTP par image, nous calculons la moyenne du groupage des trames par GOP. Ainsi pour la vidéo1 et la vidéo2, nous avons un groupage respectif de 3,2 et 4,6 trames par GOP.

5.3.4.1.1 Utilisation d'un flux concurrent UDP

Le Tableau 5-7 récapitule les tests réalisés avec un flux concurrent UDP. Dans ces tests, le débit physique du réseau 802.11 a été limité à 2 Mbits/s pour avoir un goulet d'étranglement. Le débit applicatif maximum réalisé avec ce débit est d'environ 1.7 Mbits/s avec une taille de trame de 1512 octets (voir Tableau 4-2). Le débit du flux UDP concurrent a été choisi pour avoir un débit total d'environ 1.8Mbits/s en additionnant le débit du flux UDP et le débit du flux vidéo. Pour chaque test, nous avons fait varier le nombre de trames groupées au niveau MAC pour la STA1 qui transmet le flux vidéo. La durée de chaque test est d'environ 120s, durant laquelle les stations STA1 et STA2 transmettent leurs flux à l'AP.

	Nombre de trames groupées	Vidéo	Débit du flux UDP
Test1	Sans groupage de trames	Vidéo1	900 Kbits/s
Test2	Groupage fixe de 2 trames		
Test3	Groupage fixe de 3 trames		
Test4	Groupage fixe de 4 trames		
Test5	VFG Groupage adaptatif en fonction de l'image		
Test6	Sans groupage de trames	Vidéo2	500 Kbits/s
Test7	Groupage fixe de 3 trames		
Test8	Groupage fixe de 4 trames		
Test9	VFG Groupage adaptatif en fonction de l'image		

Tableau 5-7 : Les caractéristiques des tests

Les résultats des tests s'intéressent aux débits de réception et aux pertes de paquets enregistrées pour chaque flux. Chaque test est répété dix fois et le résultat d'un test représente la moyenne des dix résultats obtenus. Les moyennes ont été calculées sur des échantillons représentatifs inclus dans l'intervalle [10s, 110s] pour éviter les résultats obtenus au début et à la fin des tests.

Les Figures 5-20 et 5-21 donnent respectivement les débits moyens de réception du flux vidéo et du flux UDP pour tous les tests utilisant la vidéo1. Quant aux Figures 5-22 et 5-23, elles donnent respectivement les taux moyens de perte pour le flux vidéo et le flux UDP.

Dans le « Test 1 », le flux vidéo qui utilise une transmission standard, c'est-à-dire sans groupage, n'arrive pas à occuper la bande passante nécessaire au niveau du lien d'accès qui permet de couvrir son débit d'émission. Par conséquent, le flux vidéo subit des pertes importantes, comparé aux flux concurrents UDP. Ces pertes sont accentuées par la variation du débit d'émission du flux vidéo, comparé au débit d'émission du flux UDP. En effet, le flux vidéo original a un débit d'émission qui varie dans l'intervalle [1064 Kbits/s, 744 Kbits/s] avec un écart type de 72 Kbits/s. Le flux UDP, quant à lui, a un débit d'émission stable de 920 Kbits/s avec un écart type de 5 Kbits/s. Durant les débits élevés, le flux vidéo subit des pertes importantes puisque le débit au niveau du lien d'accès est limité.

Dans le « Test 2 », le flux vidéo utilise un groupage fixe de deux trames qui permet au flux d'être plus agressif par rapport au flux UDP. Ceci se traduit par l'augmentation du débit de réception du flux vidéo et par une diminution de ses pertes. Le flux UDP subit l'agressivité du flux vidéo en enregistrant quelques pertes. Cependant, le groupage de deux trames ne permet pas au flux vidéo d'occuper la bande passante nécessaire qui couvre son débit d'émission.

Dans le « Test 3 », le groupage de trois trames permet au flux vidéo d'occuper une bande passante en relation avec son débit d'émission puisque les pertes s'annulent pour le flux vidéo, mais elles augmentent pour le flux UDP pour lequel le débit de réception décroît. Un résultat identique est obtenu avec le « Test 4 » qui regroupe cinq trames. Ceci confirme le fait que le groupage de trois trames dans le « Test3 » représente le meilleur groupage pour le débit de la vidéo, puisque le

groupage de plus de trames ne change pas les performances réalisées par le flux UDP. Il est important de noter que le nombre de paquets moyen par image dans un GOP pour la vidéo1 est de 3,2 paquets.

Enfin, le groupage suivant l'image vidéo dans le « Test 5 » donne également un résultat identique aux résultats de « Test 3 » et « Test 4 ». L'avantage d'utiliser le groupage par image par rapport à un groupage fixe va être démontré avec les résultats des tests utilisant la vidéo2.

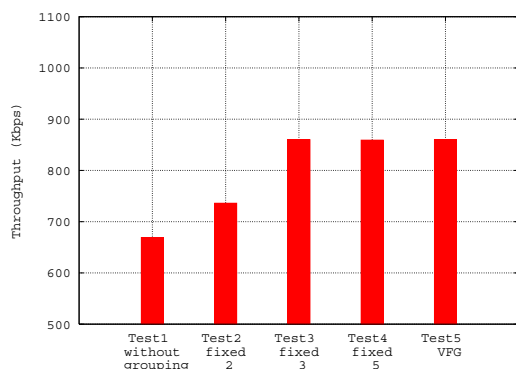


Figure 5-20 : Les débits moyens de réception du flux vidéo1 pour les tests 1, 2, 3, 4, 5

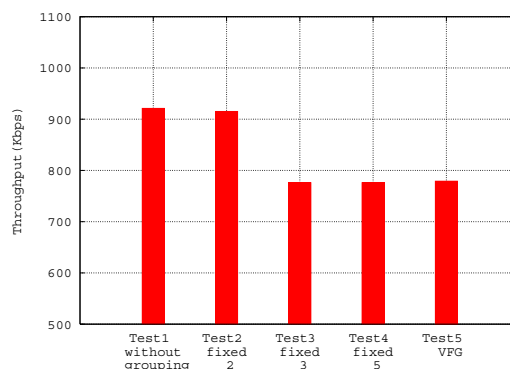


Figure 5-21 : Les débits moyens de réception du flux UDP pour les tests 1, 2, 3, 4, 5

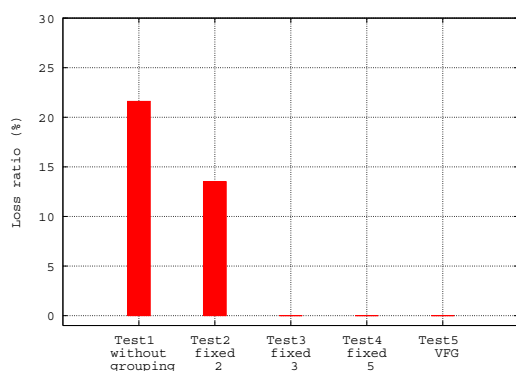


Figure 5-22 : Les taux moyens de perte du flux vidéo1 pour les tests 1, 2, 3, 4, 5

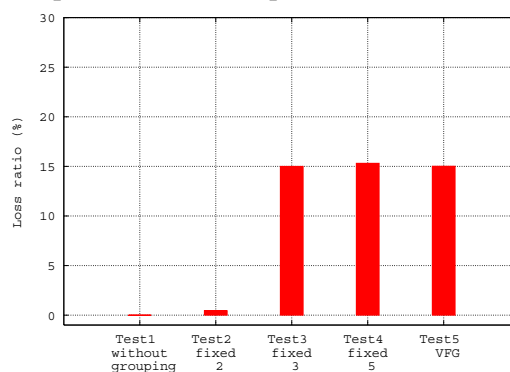


Figure 5-23 : Les taux moyens de perte du flux UDP pour les tests 1, 2, 3, 4, 5

Les Figures 5-24 et 5-25 donnent respectivement les débits moyens de réception du flux vidéo et du flux UDP pour tous les tests utilisant la vidéo2. Les Figures 5-26 et 5-27 donnent respectivement les taux moyens de perte pour le flux vidéo et le flux UDP.

D'une manière générale, les résultats du « Test 6 », « Test 7 », « Test 8 » et « Test 9 » confirment les résultats obtenus dans les tests précédents puisqu'ils suivent la même évolution. Sans groupage de trames, le flux vidéo dans le « Test 6 » subit plusieurs pertes, parce que son débit d'émission excède la bande passante occupée au niveau du lien d'accès. Dans le « Test 7 », les pertes diminuent mais elles ne s'annulent pas, parce que le groupage de trois trames ne suffit pas à couvrir le débit d'émission vidéo. Cependant, avec un groupage de cinq trames qui correspond à la moyenne arrondie du nombre de paquets par image pour un GOP, dans le « Test 8 », les pertes s'annulent complètement. Ce résultat est aussi obtenu avec le groupage des images vidéo dans le « Test 9 ». Il est clair que la diminution des pertes pour le flux vidéo est accompagnée d'une augmentation des pertes pour le flux UDP qui subit l'agressivité grandissante du flux vidéo.

Pour la vidéo1, nous avons vu que le groupage de trois trames permettait d'occuper une bande passante suffisante pour le débit d'émission. Toutefois, pour la vidéo2, le groupage de trois trames n'a pas suffi à annuler les pertes. Ici, nous constatons les limites de l'utilisation d'un groupage de trames fixe pour des vidéos possédant des débits différents.

Par contre, en utilisant le mécanisme proposé, à savoir le groupage des images vidéo VFG, l'occupation de la bande passante s'adapte automatiquement suivant le débit de la vidéo. Ceci est illustré par les résultats obtenus dans le « Test 5 » et « Test 9 ».

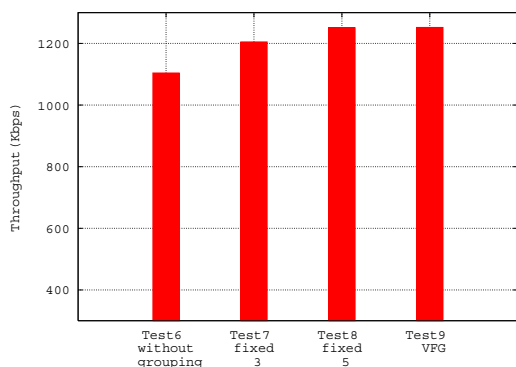


Figure 5-24 : Les débits moyens de réception du flux video2 pour les tests 6, 7, 8, 9

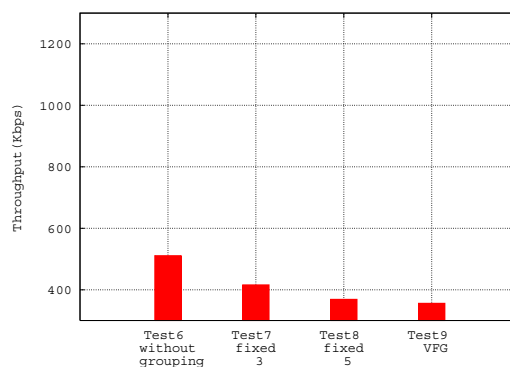


Figure 5-25 : Les débits moyens de réception du flux UDP pour les tests 6, 7, 8, 9

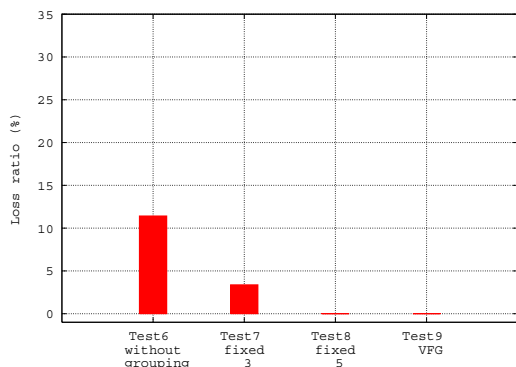


Figure 5-26 : Les taux moyens de perte du flux video2 pour les tests 6, 7, 8, 9

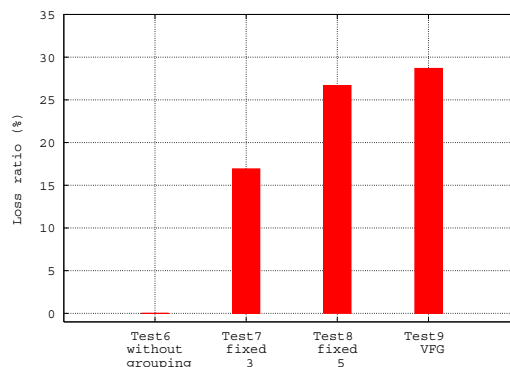


Figure 5-27 : Les taux moyens de perte du flux UDP pour les tests 6, 7, 8, 9

5.3.4.1.2 Utilisation d'un flux concurrent TCP

Afin de voir l'impact du groupage de trames sur un flux TCP, nous avons exécuté les mêmes tests décrits dans le Tableau 5-7 avec un flux concurrent TCP. Ci-dessous, nous donnons le débit moyen et le taux moyen de perte des flux pour chaque test. Il est important de noter que pour le flux TCP, nous n'avons pas de pertes grâce aux retransmissions. Cependant, le débit TCP s'adapte automatiquement à cause du contrôle de flux, et nous nous intéressons principalement au débit réalisé tout au long d'un test.

Pour « Test 1 », « Test 2 », « Test 3 », « Test 4 » et « Test 5 », les Figures 5-28 et 5-29 donnent respectivement les débits moyens de réception pour le flux vidéo et le flux TCP. Les taux moyens de perte du flux vidéo1 sont illustrés dans la Figure 5-30. Concernant « Test 6 », « Test 7 », « Test 8 » et « Test 9 », les débits moyens de réception pour le flux vidéo et le flux TCP sont donnés respectivement dans les Figures 5-31 et 5-32. La Figure 5-33 donne les taux moyens de perte pour le flux vidéo2.

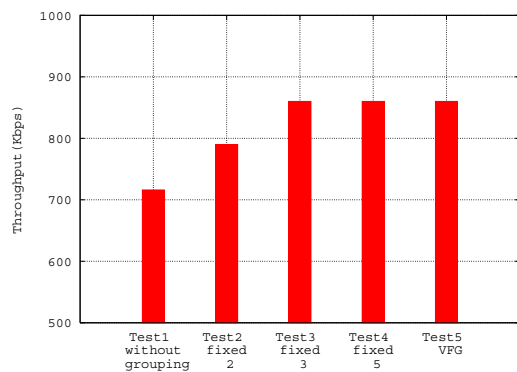


Figure 5-28 : Les débits moyens de réception du flux vidéo1 pour les tests 1, 2, 3, 4, 5

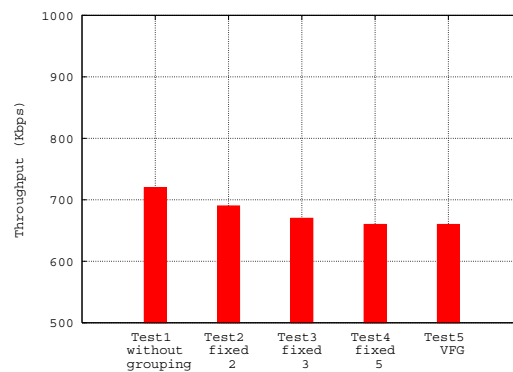


Figure 5-29 : Les débits moyens de réception du flux TCP pour les tests 1, 2, 3, 4, 5

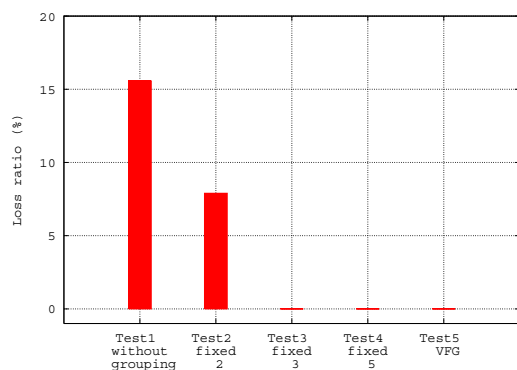


Figure 5-30 : Les taux moyens de perte du flux vidéo1 pour les tests 1, 2, 3, 4, 5

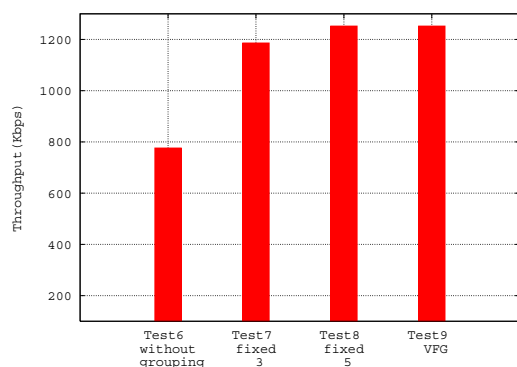


Figure 5-31 : Les débits moyens de réception du flux vidéo1 pour les tests 6, 7, 8, 9

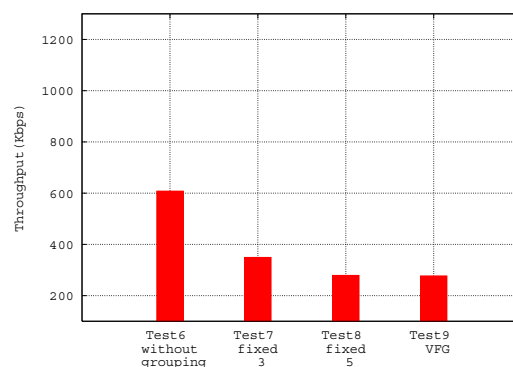


Figure 5-32 : Les débits moyens de réception du flux TCP pour les tests 6, 7, 8, 9

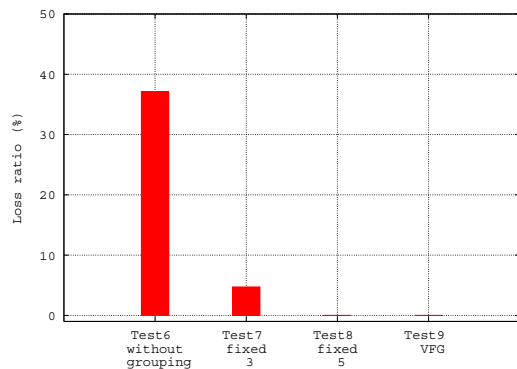


Figure 5-33 : Les taux moyens de perte du flux vidéo1 pour les tests 6, 7, 8, 9

Les résultats obtenus avec le flux TCP sont similaires à ceux obtenus avec le flux UDP. Nous constatons, clairement, que l'augmentation du nombre de trames groupées fait augmenter le débit du flux vidéo et baisser le débit du flux TCP. L'occupation du canal de ce dernier diminue au profit du flux vidéo. Le taux de perte pour le flux vidéo s'annule en regroupant trois trames pour la vidéo1 et cinq trames pour la vidéo2. De plus, le groupage de trois trames ne permet pas à la vidéo2 une occupation du canal qui couvre son débit.

Enfin, le groupage des images vidéo dans le « Test5 » et « Test9 » permet une occupation adaptative du canal qui répond au besoin en débit de la vidéo.

5.3.4.2 La capacité du VFG à assurer la bande passante

Dans cette partie, nous allons tester la capacité du VFG à occuper la bande passante pour un flux vidéo concurrencé par des flux UDP possédant des débits différents. Pour cela, nous exécutons plusieurs tests. Dans chaque test, la STA1 transmet un flux vidéo dont les caractéristiques sont données dans le Tableau 5-8 et la STA2 transmet un flux UDP avec un débit fixe.

Format de Codage	MPEG-4
Résolution	CIF : 352x288 pixels
Nombre d'images par seconde	25 fps
Débit moyen/max/min	2279/2584/1976 Kbits/s
Durée de la séquence	180 secondes
La taille d'un GOP	12
Structure d'un GOP	IBBPBBPBBPBB
Nombre moyen de paquets RTP (maximum 1450 octets) par image I	20
Nombre moyen de paquets RTP (maximum 1450 octets) par image P	11
Nombre moyen de paquets RTP (maximum 1450 octets) par image B	6

Tableau 5-8 : Les caractéristiques du flux vidéo

L'utilisation d'une vidéo différente permet de confirmer le fonctionnement du groupage des images vidéo pour un débit vidéo plus élevé. Le débit du lien d'accès pour les deux stations a été limité à 5.5 Mbits/s, ce qui permet d'avoir un débit applicatif maximum de 4.4 Mbits/s avec une taille de trame de 1512 octets (voir Tableau 4-2). La durée de chaque test est de 180 secondes. Le flux concurrent UDP est transmis durant l'intervalle [30s, 150s].

Tests	Test1	Test 2	Test 3	Test 4	Test 5	Test 6
Débit du flux UDP	500 Kbits/s	1 Mbits/s	1.5 Mbits/s	2 Mbits/s	2.5 Mbits/s	3 Mbits/s

Tableau 5-9 : Le débit du flux UDP pour les différents tests

Les tests présentés dans le Tableau 5-9 sont répétés dix fois pour les deux scénarios décrits ci-dessous :

- **Scénario 1** : La STA1 n'utilise aucun mécanisme de groupage
- **Scénario 2** : La STA1 utilise le mécanisme VFG.

Les Figures 5-34 et 5-35 présentent respectivement les débits moyens de réception pour le flux vidéo et le flux UDP. Les taux moyens de perte des deux flux sont donnés par les Figures 5-36 et 5-37.

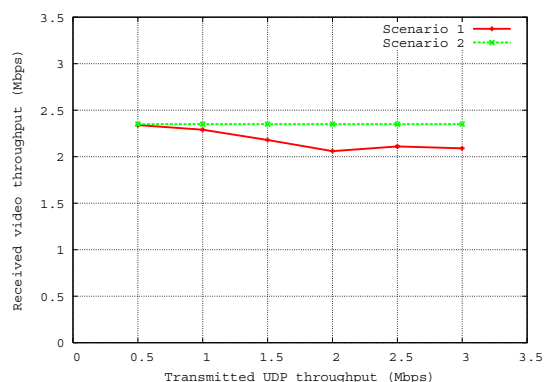


Figure 5-34 : Les débits moyens de réception du flux vidéo pour les scénarios 1 et 2

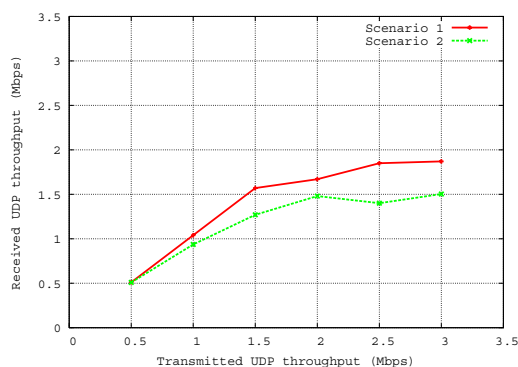


Figure 5-35 : Les débits moyens de réception du flux UDP pour les scénarios 1 et 2

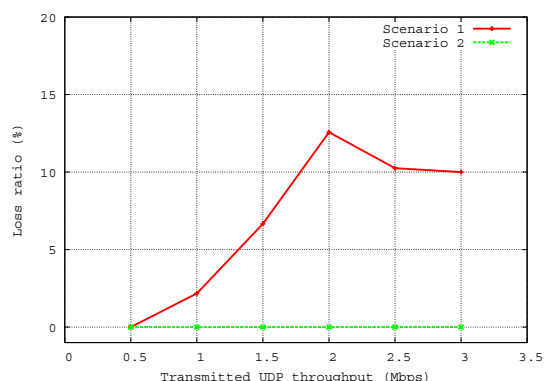


Figure 5-36 : Les taux moyens de perte du flux vidéo pour les scénarios 1 et 2

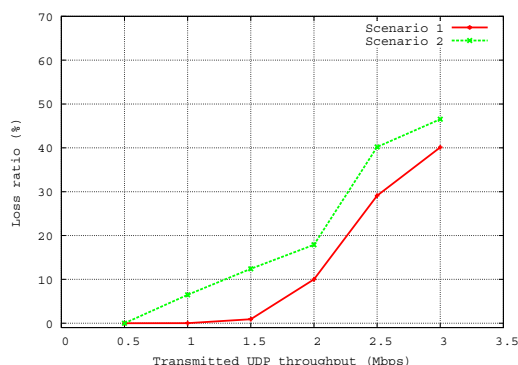


Figure 5-37 : les taux moyens de perte du flux UDP

Pour le « Test1 », les deux flux n'enregistrent aucune perte dans les deux scénarios puisque le débit du lien d'accès est supérieur au débit d'émission des deux flux. Cependant, dans le « Test2 », bien que le débit du lien d'accès absorbe théoriquement le débit d'émission moyen des deux flux, le flux vidéo enregistre des pertes de paquets dans le scénario 1. Ceci s'explique par la variation importante du débit d'émission du flux vidéo, qui varie entre 1976 Kbits/s et 2584 Kbits/s avec un écart type de 143 Kbits/s. Lorsque le débit vidéo est élevé, le débit du lien d'accès n'arrive pas à satisfaire ce besoin. Ceci se traduit par une suppression des trames vidéo au niveau de la couche MAC. Cette situation ne se produit pas avec le scénario 2 puisque le flux vidéo est plus agressif dans l'occupation du canal. Cette agressivité affecte la transmission du flux UDP, qui commence à enregistrer des pertes de paquets.

Cette tendance se confirme avec les résultats du « Test 3 », « Test4 », « Test5 » et « Test6 », durant lesquels le flux vidéo n'arrive pas à occuper assez de bande passante dans le scénario 1, ce qui provoque des pertes de plus en plus importantes, qui se stabilisent à partir du « Test 4 ». La stabilité est due au partage équitable du canal entre les deux stations dans le scénario 1. Donc chaque station possède un débit physique théorique garanti (dans nos tests 2.25 Mbits/s pour chaque station) qui ne peut pas être affecté par l'augmentation du débit du flux concurrent. Ceci explique le fait que, dans le scénario 1, l'écart entre le débit de réception des deux flux n'est pas important à partir du « Test 4 » (voir les Figures 5-34 et 5-35).

Toutefois, dans le scénario 2, ce partage équitable n'existe plus puisque la STA1 utilise le groupage des images vidéo, ce qui permet au flux vidéo d'occuper le canal suivant son débit. Ceci se traduit par la stabilité du débit de réception du flux vidéo et la réduction du débit du flux UDP qui enregistre des taux de perte élevés, comparé au scénario 1.

Ces résultats sont confirmés avec un flux concurrent TCP. Le Tableau 5-10 résume le débit de réception moyen réalisé par le flux vidéo et le flux TCP ainsi que le taux moyen de perte enregistré par le flux vidéo dans les deux scénarios. Dans le scénario 1, il n'y a pas de grandes différences entre le débit de réception de chaque flux, ce qui explique les pertes de paquets pour le flux vidéo. Par contre, dans le scénario 2, le groupage des images vidéo permet au flux vidéo d'avoir accès plus longtemps au canal, ce qui fait augmenter son débit et fait baisser, par la même, le débit du flux TCP.

	Scénario 1		Scénario 2	
	Flux TCP	Flux vidéo	Flux TCP	Flux vidéo
Taux de perte (%)		25.1%		0%
Débit (Mbit/s)	1.88 Mbits/s	1.74 Mbits/s	1.28 Mbits/s	2.35 Mbits/s

Tableau 5-10 : Le débit moyen de réception et le taux moyen de perte avec un flux concurrent TCP

5.4 Conclusion

Les adaptations *Cross-layer* descendantes exécutées durant la transmission sont aussi importantes pour améliorer la QoS des flux vidéo en termes de taux de perte et de débit de transmission. Dans ce chapitre, nous avons présenté deux contributions qui démontrent cette importance.

La première contribution propose une adaptation dynamique de la fragmentation 802.11 en fonction des conditions de transmission d'un flux vidéo. L'objectif de cette adaptation est de réduire les pertes de paquets causées par les interférences en réduisant la taille des trames transmises sur le canal sans fil. Pour cela, nous avons présenté une évaluation des différents niveaux de fragmentation à travers les couches réseaux et nous avons montré que la fragmentation MAC 802.11 offre de meilleures performances de transmission en termes de perte de paquets et d'*overhead*. Ensuite, nous avons proposé un modèle mathématique qui détermine la taille minimale d'un fragment pour chaque station présente dans le réseau en se basant sur une estimation de l'*overhead* autorisé. Cette méthodologie est exploitée dans la fragmentation adaptative qui autorise le module

XLDP à contrôler dynamiquement la fragmentation 802.11 au niveau MAC pour chaque station. Les tests de performance réalisés sur une plateforme d'expérimentation ont illustré la capacité de réaction d'un tel mécanisme afin de réduire les pertes de paquets durant les périodes d'interférence.

La deuxième contribution représente un exemple concret d'une exploitation des informations des couches supérieures au niveau des couches inférieures pour améliorer les performances de transmission. L'objectif de cette contribution est d'assurer au flux vidéo une occupation du canal proportionnel à son débit. Pour atteindre cet objectif, nous avons proposé un mécanisme de groupage des images vidéo au niveau MAC, appelé VFG. Ce dernier permet d'avoir des opportunités de transmission adaptatives, basées sur le nombre de trames. Ce nombre est délimité par une image vidéo. La couche MAC se base sur des informations « in-band » pour déterminer les trames appartenant à la même session et à la même image vidéo. Les évaluations de ce mécanisme ont montré, tout d'abord, l'avantage d'avoir un mécanisme de groupage adaptatif en relation avec le débit de la vidéo, et ensuite, la capacité du VFG à occuper le canal pour satisfaire le débit vidéo. Les résultats montrent, clairement, que le mécanisme VFG favorise les flux vidéo en minimisant les taux de perte et en garantissant le débit au niveau du lien d'accès.

Chapitre 6

Conclusion et Perspectives

Actuellement, les réseaux 802.11 sont considérés comme une alternative sérieuse aux réseaux filaires pour les réseaux d'accès. En effet, les nombreux avantages qu'offre cette technologie (rapidité de déploiement, réduction des coûts d'installation, mobilité, etc.) lui ont permis de s'imposer rapidement sur le marché des réseaux. Ce monopole est conforté par l'augmentation du débit avec les standards récents et l'intégration d'une interface de communication 802.11 dans un large panel d'équipements (webcam, disque dure externe, casque audio, etc.) et de terminaux (téléphone mobile, pda, laptop, etc.). Dans le contexte de la convergence des réseaux et des services vers la technologie IP, les réseaux d'accès 802.11 n'échappent pas à cette mouvance et sont de plus en plus sollicités pour la transmission de services IP multimédia très exigeants en QoS.

Si la QoS dans le cœur du réseau IP a été largement explorée durant ces dernières années, principalement par l'IETF, la QoS dans les réseaux d'accès reste un chantier ouvert qui commence à peine à être investi par la communauté de la recherche. La QoS dans ce dernier segment doit être assurée non seulement au niveau IP, mais à travers toutes les couches de l'architecture réseau. Pour cela, les interactions inter-couches, dites *Cross-layer*, sont nécessaires pour permettre aux couches d'échanger des métriques de performance afin d'adapter leur fonctionnement en conséquence. Ainsi, l'adaptation est le mot clé dans la conception et l'implémentation des nouveaux protocoles et des nouvelles applications. Ces derniers doivent être de plus en plus flexibles et ils doivent interagir avec leur environnement, caractérisé par une hétérogénéité grandissante.

Dans cette thèse, nous avons exploré les mécanismes d'adaptation *Cross-layer* pour la transmission des services vidéo sur les réseaux d'accès 802.11. L'objectif principal des adaptations proposées est de préserver la QoS des flux vidéo en instaurant une collaboration inter-couches fonctionnelle et utile. Dans ce qui suit, nous résumons les principales contributions qui ont été détaillées dans cette thèse et, enfin, nous introduisons quelques perspectives intéressantes qui peuvent être explorées dans de futurs travaux.

6.1 Principales contributions

Les adaptations *Cross-layer* exposées dans cette thèse sont organisées en deux grandes catégories : Les adaptations ascendantes, exécutées au niveau applicatif et les adaptations descendantes, exécutées au niveau MAC. Ces adaptations sont mises en œuvre dans le cadre d'une architecture globale qui permet la transmission des flux DVB-T dans les réseaux IP en général et dans les réseaux 802.11 en particulier. Les principales contributions sont :

- **Une architecture pour interconnecter les réseaux DVB-T et IP/802.11 :** Découpée en trois segments, l'architecture proposée permet un accès universel au service DVB-T à travers les réseaux 802.11 pour des terminaux hétérogènes et mobiles. Le réseau DVB-T est relié au cœur du réseau IP à travers l'entité TVSP (TV Stream Provider) qui effectue une correspondance entre les standards DVB-T et IP afin de transformer les services DVB-T en service IPTV. Les flux IPTV sont transmis dans les réseaux d'accès 802.11 à travers une passerelle d'adaptation XLAG (Cross Layer Adaptation Gateway) qui embarque un nouveau système de transmission audio/vidéo adaptatif, appelé XLAVS (Cross Layer Adaptive Video Streaming). Ce dernier communique activement avec son environnement afin de personnaliser le flux TV suivant le terminal de réception et maintenir la QoS des flux vidéo en mettant en œuvre plusieurs adaptations *Cross-layer*. Pour valider le fonctionnement de cette architecture, un prototype a été réalisé en se basant sur des bibliothèques et des outils libres.
- **Une adaptation dynamique du débit vidéo en fonction du débit physique :** La première adaptation fait partie des techniques *Cross-layer* ascendantes. Elle permet au XLAVS d'adapter dynamiquement le débit des flux vidéo transmis suivant la variation du débit physique 802.11. Pour cela, nous avons proposé une technique qui permet de contrôler dynamiquement le débit d'un flux vidéo en se basant sur la résolution SNR. Nous avons également proposé un modèle qui permet de calculer le débit physique réellement disponible lorsque le XLAG utilise un débit physique différent pour chaque station. L'utilité de cette adaptation a été démontrée par des expérimentations réelles réalisées sur le prototype conçu.
- **Un système FEC adaptatif accompagné d'une adaptation dynamique du débit vidéo :** Cette adaptation appartient aussi aux techniques *Cross-layer* ascendantes en faisant collaborer plusieurs techniques de QoS (la FEC adaptative et le débit vidéo adaptatif) et plusieurs métriques de performance (la puissance du signal physique et les taux de perte applicatifs). L'objectif de cette adaptation est d'avoir un flux vidéo résistant aux pertes avec un débit stable. Les différentes expérimentations réalisées ont démontré que la puissance du signal représentait un meilleur indicateur de l'état du canal en anticipant l'apparition des pertes et que la variation du débit vidéo permettait de préserver la capacité spectrale du canal lorsque la FEC est utilisée.
- **Une fragmentation 802.11 adaptative :** Dans cette adaptation nous avons exploré l'apport des techniques *Cross-layer* descendantes dans la préservation de la QoS des flux vidéo. À cet effet, nous avons proposé un contrôle dynamique de la fragmentation 802.11 au niveau MAC par le nouveau système de transmission XLAVS. Ce contrôle permet de déterminer la taille des trames qui offre un compromis entre les taux de perte applicatifs et l'*overhead* introduit par la fragmentation. L'expérimentation de cette adaptation a montré la capacité de cette dernière à réduire les pertes durant les périodes d'interférence.

- **Le groupage des images vidéo au niveau MAC (VFG : Video Frame Grouping) :** Ce mécanisme s'inscrit aussi parmi les techniques *Cross-layer* descendantes. Mais dans ce cas, la couche MAC utilise des informations applicatives appartenant au flux vidéo afin de lui permettre un accès au canal proportionnel à son débit. Ceci permet au flux vidéo d'avoir plus de bande passante au niveau physique, comparé aux autres flux concurrents. Les expérimentations de ce mécanisme ont montré la capacité du VFG à fournir au flux vidéo une occupation adaptative du canal en fonction de son débit.

6.2 Perspectives et nouveaux défis

Les travaux exposés dans cette thèse n'apportent pas de réponses à toutes les problématiques engendrées par la transmission des services vidéo dans les réseaux 802.11. Plusieurs perspectives peuvent être données pour améliorer l'architecture IPTV et les adaptations *cross-layer* proposées. Nous détaillons ci-dessous ces perspectives.

- **La scalabilité des mécanismes *Cross-layer* proposés :** Les adaptations *Cross-layer* présentées dans cette thèse ont été expérimentées en utilisant une plateforme d'expérimentation très limitée (au maximum trois stations). Ceci pose inévitablement le problème de la scalabilité de ces adaptations sur une plateforme plus large contenant plusieurs stations. La scalabilité peut être étudiée en augmentant le nombre de stations ou bien en effectuant des simulations.
- **La coexistence des adaptations *Cross-layer* :** Dans cette thèse, nous avons identifié, étudié et implémenté quatre adaptations plus au moins indépendantes. Il est important d'étudier la coexistence de toutes ces adaptations sur une seule plateforme. Nous avons vu dans le chapitre 4, l'avantage de faire collaborer plusieurs mécanismes de QoS (l'adaptation FEC et l'adaptation du débit vidéo). Ce concept doit être étendu afin de permettre au module XLDP d'exécuter des adaptations cohérentes suivant l'état du système. Ceci engendre le problème du mapping entre les métriques et les paramètres de configuration puisque ces derniers sont récupérés sur différentes couches et possèdent des unités de mesures complètement différentes.
- **L'introduction de l'apprentissage dans l'adaptation :** Les politiques d'adaptation utilisées au niveau du XLDP sont simplistes et elles peuvent être modifiées par un administrateur externe. Durant ces dernières années, nous avons assisté à l'apparition de nouveaux concepts tels que l'autonomie et l'apprentissage qui permettent aux entités réseaux de s'auto-configurer, s'auto-organiser et s'auto-gérer suivant leur environnement. Ces deux nouveaux concepts peuvent être exploités au niveau du XLDP pour modifier dynamiquement les politiques d'adaptation si celles-ci n'arrivent pas à améliorer la QoS.
- **La proposition d'un nouvel algorithme pour le contrôle du débit physique :** Nous avons présenté dans le chapitre 4, les algorithmes de variation du débit physique RCA (Rate Control Algorithm). Nous avons expliqué comment le fonctionnement de ces algorithmes influait sur la transmission des flux vidéo. Nous projetons de proposer un RCA adapter aux caractéristiques des flux multimédia.

- Le développement de nouvelles métriques pour la qualité vidéo sans référence :**
 L'objectif ultime des adaptations *Cross-layer* est de préserver la qualité de la vidéo perçue par l'utilisateur, côté récepteur. La métrique de performance considérée dans la majorité de ces adaptations est le taux de perte applicatif. Nous avons vu que les pertes de paquets dégrade la qualité vidéo mais la dégradation dépend principalement du type d'images auxquelles appartiennent les paquets perdus. En effet, si la transmission provoque un taux de perte de 5%, la qualité vidéo ne sera pas identique si les paquets perdus appartiennent aux images I, aux images P ou aux images B. Pour cela, il devient nécessaire d'avoir des métriques qui puissent informer sur la qualité vidéo perçue par le récepteur en temps réel et sans utiliser la vidéo de référence, comme cela est fait avec les métriques traditionnelles (PSNR, SSIM, etc.).
- L'optimisation du transcodage et l'utilisation du codage hiérarchique SVC :** Le nouveau système de transmission XLAVS transcode en temps réel les flux audio/vidéo afin de personnaliser le flux IPTV aux caractéristiques du terminal et aux préférences de l'utilisateur. Pour transcoder un flux vidéo d'un format de codage vers un autre format, les images vidéo sont décodées complètement à partir de leur format initial, c'est-à-dire MPEG-2, ensuite ré-encodées suivant le nouveau format. En sachant que les architectures de codage MPEG (MPEG-1, MPEG-2, MPEG-4, H.264) ne diffèrent pas beaucoup, il serait intéressant de trouver des éléments communs entre les codecs afin d'éviter de décoder complètement les images. Ceci permet d'économiser le calcul de certaines informations, par exemple, les vecteurs de mouvements. D'un autre côté, l'apparition des premières implémentations du codage hiérarchique SVC, comme le JSVM (Joint Scalable Video Model) [215], nous oblige à considérer ce nouveau codec dans notre architecture puisqu'il permet une adaptation dynamique du flux vidéo sur trois dimensions (spaciale, temporelle et SNR) sans aucun transcodage. Cependant, le transcodage reste nécessaire pour passer d'un format de codage vers un autre.

Références

- [1] G. Pujolle, "Les réseaux", 5ème édition, ISBN : 2-212-11437-0, Eyrolles, septembre 2004.
- [2] Andrew Tanenbaum, "Réseaux", 4ème édition, ISBN : 2744070017, Pearson Education, mai 2003.
- [3] IEEE 802.11, IEEE Standards for Information Technology -- Specific Requirements -- Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications, Edition (ISO/IEC 8802-11: 1999), 1999.
- [4] IEEE 802.11a, (8802-11:1999/Amd 1:2000(E)), IEEE Standard for Information technology—Specific requirements—Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications—Amendment 1: High-speed Physical Layer in the 5 GHz band, 1999.
- [5] IEEE 802.11b Supplement to 802.11-1999, Wireless LAN MAC and PHY specifications: Higher speed Physical Layer (PHY) extension in the 2.4 GHz band, 1999.
- [6] IEEE 802.11g, IEEE Standard for Information technology—Specific requirements—Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications—Amendment 4: Further Higher-Speed Physical Layer Extension in the 2.4 GHz Band, 2003.
- [7] IEEE P802.11n_D3.00, Approved Draft Standard for Information Technology-Telecommunications and information exchange between systems--Local and metropolitan area networks--Specific requirements-- Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications: Amendment 4: Enhancements for Higher Throughput, Sep 2007.
- [8] David Tse and Pramod Viswanath, "Fundamentals of Wireless Communication", ISBN-13: 978-0521845274, Cambridge University Press, mai 2005.
- [9] Quin Yan, "A Study of Transmit and Receive Antenna Diversity Techniques for Communication Systems", presented to the Graduate and Research Committee of Lehigh University in candidacy for degree of Doctor of Philosophy in Electrical Engineering, Lehigh University, April 2001.
- [10] J. N. Laneman, et al., "Comparing application- and physical-layer approaches to diversity on wireless channels", in Proc of IEEE ICC'03, vol.4, pp. 2678 - 2682, May 2003.
- [11] J. N. Laneman, et al., "Source-channel diversity approaches for multimedia communication", in IEEE Transactions on Information Theory, vol. 51, no. 10, pp. 3518-3539, October 2005.
- [12] M.Yuksel, E.Erkip, "Multiple-Antenna Cooperative Wireless Systems: A Diversity-Multiplexing Tradeoff Perspective", IEEE Transactions on Information Theory, Volume 53, Issue 10, Page(s):3371 - 3393, octobre 2007.
- [13] William Stallings, Jean-Alain Hernandez, René Joly, "Réseaux et communication sans fil", 2ème édition, ISBN-13: 978-2744070853, Pearson Education, mai 2005.
- [14] J. Deng and R.S. Chang, "A priority scheme for IEEE 802.11 DCF access method", in IEICE Transactions on Communication, vol. E82-B, no. 1, pp. 96-102, January 1999.
- [15] W. Pattara-Atikom, S. Banerjee, and P. Krishnamurthy, "Starvation Prevention and Quality of Service in Wireless LANs", The 5th International Symposium on Wireless Personal Multimedia Communications, Volume 3, Page(s): 1078 - 1082, Oct. 2002.
- [16] I. Aad , C. Castelluccia, "Differentiation mechanisms for IEEE 802.11", in Proc. of IEEE INFOCOM'01, vol. 1, pp.209–218, Anchorage, Alaska, April 2001.
- [17] Y. Ge and J. Hou, "An Analytical Model for Service Differentiation in IEEE 802.11," IEEE ICC '03, vol. 2, pp. 1157–62, May 2003.
- [18] L. Romdhani, et al., "Adaptive EDCAF: enhanced service differentiation for IEEE 802.11 wireless ad-hoc networks", in Proc. of IEEE WCNC'03, vol. 2, pp. 1373-1378, New Orleans, Louisiana, March 2003.
- [19] M. Malli, et al., "Adaptive fair channel allocation for QoS enhancement in IEEE 802.11 wireless LAN", in Proc. of IEEE ICC'04, vol. 6, pp. 347 –3475, Paris, France, July 2004.
- [20] Z.J. Haas and J.Deng, "On optimizing the backoff interval for random access schemes", in IEEE Transactions on Communications, vol. 51, no. 12, pp. 2081-2090, December 2003.

- [21] A. Banchs and X. Pérez, "Providing Throughput Guarantees in IEEE 802.11 Wireless LAN," IEEE WCNC '02, vol.1, pp. 130–38, 2002.
- [22] A. Banchs and X. Pérez, "Distributed Weighted Fair Queuing in 802.11 Wireless LAN," IEEE ICC '02, vol. 5, pp. 3121–27, Apr. 2002.
- [23] N. H. Vaidya, P. Bahl, and S. Gupta, "Distributed Fair Scheduling in a Wireless LAN," Proc. ACM MOBICOM 2000, pp. 167–78, Aug. 2000.
- [24] K. Liu et al., "A Reservation-based Multiple Access Protocol with Collision Avoidance for Wireless Multihop Ad Hoc Networks" IEEE ICC '03, vol. 2, pp. 1119–23, May 2003.
- [25] A. Lindgren, A. Almquist, and O. Schelén, "Evaluation of Quality of Service Schemes for IEEE 802.11 Wireless LANs," Proceedings on 26th Annual IEEE Conference on Local Computer Networks, Page(s):348 - 351, 2001.
- [26] M. Barry, A. T. Campbell, and A. Veres, "Distributed Control Algorithms for Service Differentiation in Wireless Packet Networks" Proc. IEEE INFOCOM, 2001.
- [27] S. Valaee and B. Li, "Distributed Call Admission Control in Wireless Ad Hoc Networks," Proc. IEEE VTC 2002, Vancouver, BC, Canada, Sept. 24–28, 2002.
- [28] S. H. Shah, K. Chen, and K. Nahrstedt, "Dynamic Bandwidth Management for Single-Hop Ad Hoc Wireless Networks," Proc. IEEE Int'l. Conf. Perv. Comp. and Commun. 2003.
- [29] M. Kazantzidis, M. Gerla, and S.-J. Lee, "Permissible Throughput Network Feedback for Adaptive Multimedia in AODV MANETs," IEEE ICC '01, vol. 5, pp. 1352–56, June 2001.
- [30] K. Liu et al., "A Reservation-based Multiple Access Protocol with Collision Avoidance for Wireless Multihop Ad Hoc Networks," IEEE ICC '03, vol. 2, pp. 1119–23, May 2003.
- [31] IEEE 802.11e, IEEE Standard for Information technology—Telecommunications and information exchange between systems—Local and metropolitan area networks—Specific requirements Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications: Amendment 8: Medium Access Control (MAC) Quality of Service Enhancements, 2005.
- [32] Jon Postel, "RFC 791 : INTERNET PROTOCOL", Request for Comments, IETF, sep. 1981.
- [33] S. Deering, R. Hinden, "Internet Protocol, Version 6 (IPv6) Specification", Request for Comments, IETF, Dec. 1998.
- [34] R. Braden, D. Clark, S. Shenker, "RFC 1633: Integrated Services in the Internet Architecture: an Overview", Request for Comments, IETF, Jun. 1994.
- [35] S. Shenker, J. Wroclawski, "RFC 2215: General Characterization Parameters for Integrated Service Network Elements", Sep. 1997.
- [36] J. Wroclawski, "RFC: 2211 Specification of the Controlled-Load Network Element Service", Request for Comments, IETF, Sep. 1997.
- [37] S. Shenker, C. Partridge, R. Guérin "RFC 2212 : Specification of Guaranteed Quality of Service", Request for Comments, IETF, Sep. 1997.
- [38] R. Braden, L. Zhang, S. Berson, S. Herzog, S. Jamin, "RFC 2205 : Resource ReSerVation Protocol (RSVP)-- Version 1 Functional Specification", Request for Comments, IETF, Sep. 1997.
- [39] A. Mankin, F. Baker, B. Braden, S. Bradner, M. O'Dell, A. Romanow, A. Weinrib, L. Zhang "RFC 2208 : Resource ReSerVation Protocol (RSVP) Version 1 Applicability Statement Some Guidelines on Deployment", Request for Comments, IETF, Sep. 1997.
- [40] R. Braden, L. Zhang, "RFC 2009 : Resource ReSerVation Protocol (RSVP) -- Version 1 Message Processing Rules", Request for Comments, IETF, Sep. 1997.
- [41] S. Blake, D. Black, M. Carlson, E. Davies, Z. Wang, W. Weiss, "RFC 2475: An Architecture for Differentiated Services", Request for Comments, IETF, Dec. 1998.
- [42] K. Nichols, S. Blake, F. Baker, D. Black, "Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers", Request for Comments, IETF, Dec. 1998.
- [43] V. Jacobson, K. Nichols, K. Poduri, "RFC 2598: An Expedited Forwarding PHB", Request for Comments, IETF, Jun. 1999.
- [44] J. Heinanen, F. Baker, W. Weiss, J. Wroclawski, "RFC 2597: Assured Forwarding PHB Group", Request for Comments, IETF, Jun. 1999.

- [45] E. Rosen, A. Viswanathan, R. Callon, "RFC 3031: Multiprotocol Label Switching Architecture", Request for Comments, IETF, Jan. 2001.
- [46] L. Andersson, P. Doolan, N. Feldman, A. Fredette, B. Thomas, "RFC 3036: LDP Specification", Request for Comments, IETF, Jan. 2001.
- [47] D. Awduche, L. Berger, D. Gan, T. Li, V. Srinivasan, G. Swallow, "RSVP-TE: Extensions to RSVP for LSP Tunnels", Request for Comments, IETF, Dec. 2001.
- [48] Wei Sun, P. Bhaniramka, R. Jain, "Quality of service using traffic engineering over MPLS: an analysis", p. 238, 25th Annual IEEE International Conference on Local Computer Networks, 2000.
- [49] Ji-Young Lim, Ki-Joon Chae, "Differentiated Link Based QoS Routing Algorithms for Multimedia Traffic in MPLS Networks", p. 587, 15th International Conference on Information Networking, 2001.
- [50] Na Lin, Hongman Qi, "A QoS Model of Next Generation Network based on MPLS", pp. 915-919, IFIP International Conference on Network and Parallel Computing Workshops, 2007.
- [51] S.I. Maniatis, E.G. Nikolouzou, I.S. Venieris, "End-to-end QoS specification issues in the converged all-IP wired and wireless environment", IEEE Communications Magazine, Volume 42, Issue 6, Page(s): 80 - 86, Jun. 2004.
- [52] J. Martin, A. Nilsson, "On service level agreements for IP networks", Proc in the 21st Annual Joint Conference of the IEEE Computer and Communications Societies INFOCOM 2002, Volume 2, Issue , Page(s): 855 - 863, 2002.
- [53] N. Thi Mai Trang, N. Boukhatem, G. Pujolle, «COPS-SLS Usage for dynamic Policy-Based QoS Management over Heterogenous IP Networks », Proc in IEEE Network, Volume 17, Issue 3, Page(s): 44 - 50, May-June 2003.
- [54] D. Durham, J. Boyle, R. Cohen, S. Herzog, R. Rajan, A. Sastry, " RFC 2748 : The COPS (Common Open Policy Service) Protocol", Request for Comments, IETF, Janvier 2000.
- [55] TEQUILA Project, «Traffic Engineering for Quality of Service in the Internet, at Large Scale», accessible via : <http://www.ist-tequila.org>
- [56] TEQUILA Project, «Final Architecture, Protocol and Algorithm Specification », Livrable Tequila D 3.4 Part B, Octobre 2003.
- [57] J.C. Chen et al., "Design and implementation of Dynamic Service Negotiation Protocol (DSNP)," Elsevier J. Computer Communications, Jun. 2006.
- [58] J. Manner, G. Karagiannis, A. McDonald, "NSLP for Quality-of-Service Signaling", draft-ietf-nsis-qos-nslp-16.txt, IETF, Feb 2008.
- [59] R. Hancock, G. Karagiannis, J. Loughney, S. Van den Bosch, "RFC 4080 : Next Steps in Signaling (NSIS): Framework", Request for Comments, IETF, June 2000.
- [60] R. Yavatkar, D. Hoffman, Y. Bernet, F. Baker, M. Speer, "SBM (Subnet Bandwidth Manager): A Protocol for RSVP-based Admission Control over IEEE 802-style networks", Request for Comments, IETF, May 2000.
- [61] "Information technology - Telecommunications and information exchange between systems - Local and metropolitan area networks - Common specifications - Part 3: Media Access Control (MAC) Bridges: Revision (Incorporating IEEE P802.1p: Traffic Class Expediting and Dynamic Multicast Filtering)", ISO/IEC Final CD 15802-3 IEEE P802.1D/D15, November 24, 1997.
- [62] Jon Postel, "RFC 768 : User Datagram Protocol", Request for Comments, IETF, aug. 1980.
- [63] Jon Postel, "RFC 793 : Transmission Control Protocol", Request for Comments, IETF, sep. 1981.
- [64] Charles Krasic, Kang Li, Jonathan Walpole, "The Case for Streaming Multimedia with TCP", Proceedings of the 8th International Workshop on Interactive Distributed Multimedia System, Vol. 2158, Pages: 213 - 218, 2001.
- [65] Dapeng Wu, Yiwei Thoms Hou, Ya-Qin Zhang, "Transporting real-time video over the Internet: challenges and approaches", Proceedings of the IEEE, Volume 88, Issue 12, Page(s):1855 - 1877, Dec 2000.
- [66] Dapeng Wu, Y.T.Hou, Wenwu Zhu, Ya-Qin Zhang, J.M.Peha, "Streaming video over the Internet: approaches and directions", IEEE Transactions on Circuits and Systems for Video Technology, Volume: 11, Issue: 3, On page(s): 282-300, Mar 2001.
- [67] D.Bansal, H.Balakrishnan, "Binomial congestion control algorithms", Proceedings in IEEE Twentieth Annual Joint Conference of the IEEE Computer and Communications Societies, INFOCOM 2001, Volume 2, Page(s):631 - 640 vol.2, 2001.
- [68] D. Wu, Y. T. Hou, W. Zhu, H.-J. Lee, T. Chiang, Y.-Q. Zhang, and H. J. Chao, "On end-to-end architecture for transporting MPEG-4 video over the Internet," IEEE Transaction on Circuits Syst. Video Technol., vol. 10, pp. 923-941, Sept. 2000.

- [69] T. Turletti and C. Huitema, "Videoconferencing on the Internet," *IEEE/ACM Transaction on Networking*, vol. 4, pp. 340–351, June 1996.
- [70] J. -C. Bolot, T. Turletti, and I. Wakeman, "Scalable feedback control for multicast video distribution in the Internet," in *Proc. ACM SIGCOMM'94*, pp. 58–67, London U.K, Sept. 1994.
- [71] S. Floyd and K. Fall, "Promoting the use of end-to-end congestion control in the Internet," *IEEE/ACM Transaction on Networking*, vol. 7, pp. 458–472, Aug. 1999.
- [72] M. Handley, S. Floyd, J. Padhye, J. Widmer, "RFC 3448 : TCP Friendly Rate Control (TFRC): Protocol Specification", Request for Comments, IETF, January 2003.
- [73] S. Floyd, E. Kohler, "RFC 4828 : TCP Friendly Rate Control (TFRC): The Small-Packet (SP) Variant", Request for Comments, IETF, April 2007.
- [74] S. McCanne, V. Jacobson, and M. Vetterli, "Receiver-driven layered multicast" in *Proc. ACM SIGCOMM'96*, pp. 117–130, Aug. 1996.
- [75] L. Vicisano, L. Rizzo, and J. Crowcroft, "TCP-like congestion control for layered multicast data transfer," in *Proc. IEEE INFOCOM'98*, vol. 3, pp. 996–1003, Mar. 1998.
- [76] S. Y. Cheung, M. Ammar, and X. Li, "On the use of destination set grouping to improve fairness in multicast video distribution," in *Proc. IEEE INFOCOM'96*, San Francisco, CA, pp. 553–560, Mar. 1996.
- [77] R. Stewart, "RFC 4960 : Stream Control Transmission Protocol", Request for Comments, IETF, September 2007.
- [78] R. Stewart, M. Ramalho, Q. Xie, M. Tuexen, P. Conrad, "RFC 3758 : Stream Control Transmission Protocol (SCTP) Partial Reliability Extension", Request for Comments, IETF, May 2004.
- [79] A. Argyriou, V. Madiseti, "Streaming H.264/AVC video over the Internet", *Proc In First IEEE Consumer Communications and Networking Conference, CCNC 2004*, On page(s): 169- 174, Jan. 2004.
- [80] Li Wang, Ken'ichi Kawanishi, Yoshikuni Onozato, "Simulation-Based Optimization on MPEG-4 over SCTP Multi-streaming with Differentiated Retransmission Policy in Lossy Link", *The 2nd IEEE Asia-Pacific Service Computing Conference (APSCC 2007)*, pp. 164-171, 2007.
- [81] Sang Tae Kim, Seok Joo Koh, Yong Jin Kim, "Performance of SCTP for IPTV Applications", *The 9th International Conference on Advanced Communication Technology*, Volume: 3, On page(s): 2176-2180, Feb. 2007.
- [82] E. Kohler, S. Floyd, "RFC 4340: Datagram Congestion Control Protocol (DCCP)", Request for Comments, IETF, March 2006.
- [83] S. Takeuchi, H. Koga, K. Iida, Y. Kadobayashi, S. Yamaguchi, "Performance evaluations of DCCP for bursty traffic in real-time applications", *Proc In The Symposium on Applications and the Internet*, On page(s): 142- 149, Jan-Feb 2006.
- [84] M. Burak Gorkemli, Reha Civanlar, "SVC Coded Video Streaming over DCCP", *The 8th IEEE International Symposium on Multimedia*, Page(s): 437 - 441, Dec 2006.
- [85] S. Shakkotai, T. Rappaport, P. Karlsson, "Cross-layer design for wireless networks", *IEEE Communications Magazine*, vol. 41, no. 10, pp. 74-80, October 2003.
- [86] ITU-T Recommendation H.323 (11/96), "Visual telephone systems and equipment for local area networks which provide a non guaranteed quality of service", Nov. 1996.
- [87] ITU-T Recommendation H.323 Draft v4 (11/2000), "Packet-based multimedia communications systems", Nov. 2000.
- [88] J. Rosenberg, H. Schulzrinne, G. Camarillo, A. Johnston, J. Peterson, R. Sparks, M. Handley, E. Schooler, "RFC 3261: SIP : Session Initiation Protocol", Request for Comments, IETF, June 2002.
- [89] H. Schulzrinne, A. Rao, R. Lanphier, "RFC 2326 : Real Time Streaming Protocol (RTSP)", Request for Comments, IETF, April 1998.
- [90] M. Handley, V. Jacobson, C. Perkins, "RFC 4566 : SDP : Session Description Protocol", Request for Comments, IETF, July 2006.
- [91] M. Handley, C. Perkins, E. Whelan, "RFC: 2974 Session Announcement Protocol", Request for Comments, IETF, Oct. 2000.
- [92] H. Schulzrinne, S. Casner, R. Frederick, V. Jacobson, "RFC 3550 : RTP : A Transport Protocol for Real-Time Applications", Request for Comments, IETF, July. 2003.
- [93] ITU-T Recommendation H.261: "Video codec for audiovisual services at p x 64 kbit/s", March 1993.

- [94] ITU-T Recommendation H.263 : "Video coding for low bitrate communication", March 1996.
- [95] ISO/IEC JTC1 IS 11172 (MPEG-1), "Coding of moving picture and coding of continuous audio for digital storage media up to 1.5 Mbps", 1992.
- [96] ISO/IEC JTC1 IS 13818 (MPEG-2), "Generic coding of moving pictures and associated audio", 1994.
- [97] ISO/IEC JTC1 IS 14386 (MPEG-4), "Generic Coding of Moving Pictures and Associated Audio", 2000.
- [98] ISO/IEC 14496-10 AVC or ITU-T Rec. H.264, September 2003.
- [99] ISO/IEC JTC 1/SC 29/WG 11, Joint Draft 5 : Scalable Video Coding", Bangkok, Jan. 2006.
- [100] P. Amon and J. Pandel, "Evaluation of adaptive and reliable video transmission technologies", In Proc. of the 13th Packet Video Workshop, Nantes France, 2003.
- [101] N. Färber and B. Girod, "Robust H.263 compatible video transmission for mobile access to video servers," in Proc. IEEE International Conference on Image Processing (ICIP 1997), vol. 2, pp. 73–76, Oct. 1997.
- [102] G. J. Conklin, G. S. Greenbaum, and K. O. Lillevold, "Video coding for streaming media delivery on the Internet," IEEE Transactions on Circuits and Systems for Video Technology, vol. 11, no. 3, pp. 269–281, Mar. 2001.
- [103] M. Karczewicz and R. Kurceren, "The SP- and SI-frames design for H.264/AVC," IEEE Transactions on Circuits and Systems for Video Technology, vol. 13, no. 7, pp. 637–644, July 2003.
- [104] A. Vetro, C. Christopoulos, Huifang Sun, "Video transcoding architectures and techniques: an overview", IEEE Signal Processing Magazine, Volume 20, Issue 2, Page(s):18 - 29, March 2003.
- [105] J. Xin, C.-W. Lin, M.-T. Sun, "Digital Video Transcoding", Proceedings of the IEEE, Volume: 93 , Issue: 1, On page(s): 84 - 97, Jan. 2005.
- [106] I. Ahmad, Xiaohui Wei, Yu Sun, Ya-Qin Zhang, "Video transcoding : an overview of various techniques and research issues", IEEE Transactions on Multimedia, Volume 7, Issue 5, Page(s):793 - 804, Oct. 2005.
- [107] Peng Chen; Jeongyeon Lim; Bumshik Lee; Munchul Kim; Sangjin Hahm; Byungsun Kim; Keunsik Lee; Keunsoo Park, "A network-adaptive SVC Streaming Architecture", The 9th International Conference on Advanced Communication Technology, Volume 2, Page(s):955 - 960, Feb. 2007.
- [108] Z. Avramova, D. De Vleeschauwer, K. Spaey, Wittevrongel, S.; Bruneel, H.; Blondia, C., "Comparison of simulcast and scalable video coding in terms of the required capacity in an IPTV network", Packet Video 2007, Page(s):113 - 122, Nov. 2007.
- [109] Real-Time System for Adaptive Video Streaming Based on SVC Wien, M.; Cazoulat, R.; Graffunder, A.; Hutter, A.; Amon, P.; Circuits and Systems for Video Technology, IEEE Transactions on Volume 17, Issue 9, Page(s):1227 - 1237, Sept. 2007.
- [110] M. van der Schaar, D.S. Turaga, "Multiple description scalable coding using wavelet-based motion compensated temporal filtering", Proc on International Conference on Image Processing (ICIP 2003), Volume 3, Page(s):III - 489-92, Sept 2003.
- [111] Zhe Wei, Canhui Cai, Kai-Kuang Ma, "A Novel H.264-based Multiple Description Video Coding Via Polyphase Transform and Partial Prediction", International Symposium on Intelligent Signal Processing and Communications (ISPACS '06), Page(s):151 - 154, Dec. 2006.
- [112] Zhe Wei, Canhui Cai, Kai-Kuang Ma, "H.264-based Multiple Description Video Coder and Its DSP Implementation", IEEE International Conference on Image Processing, Page(s):3253 - 3256, Oct. 2006.
- [113] J. Apostolopoulos, T. Wong, Wai-tian Tan, S. Wee, "On multiple description streaming with content delivery networks", Proceedings in IEEE INFOCOM 2002, Twenty-First Annual Joint Conference of the IEEE Computer and Communications Societies, Volume 3, Page(s):1736 - 1745, June 2002.
- [114] Multiple description coding for Internet video streaming Pereira, M.; Antonini, M.; Barlaud, M.; Image Processing, 2003. ICIP 2003. Proceedings. 2003 International Conference on Volume 3, Page(s):III - 281-4, Sept. 2003.
- [115] Multiple Description Coding for Video Delivery Wang, Y.; Reibman, A.R.; Lin, S.; Proceedings of the IEEE, Volume 93, Issue 1, Page(s):57 - 70, Jan. 2005.
- [116] Joohee Kim; Mersereau, R.M.; Altunbasak, Y., "Distributed video streaming using multiple description coding and unequal error protection", IEEE Transactions on Image Processing, Volume 14, Issue 7, Page(s):849 - 861, July 2005.

- [117] Jong-Tzy Wang, Pao-Chi Chang, "Error-propagation prevention technique for real-time video transmission over ATM networks", IEEE Transactions on Circuits and Systems for Video Technology, Volume 9, Issue 3, Pages: 513 - 523, Apr. 1999.
- [118] M. Claypool and Y. Zhu, "Using interleaving to ameliorate the effects of packet loss in a video stream", in Proc. of the International Workshop on Multimedia Network Systems and Applications, Providence, Rhode Island, May 2003.
- [119] J. van der Meer, D. Mackie, V. Swaminathan, D. Singer, P. Gentric, "RFC 3640 : RTP Payload Format for Transport of MPEG-4 Elementary Streams", Request for Comments, IETF, November 2003.
- [120] C. Papadopoulos and G. M. Parulkar, "Retransmission-based Error Control for Continuous Media Applications", Proc. 6th International Workshop on Network and Operating Systems Support for Digital Audio and Video (NOSSDAV), April 1996.
- [121] D.Loguinov, H.Radha, "On retransmission schemes for real-time streaming in the Internet", Proceedings In IEEE INFOCOM 2001, Twentieth Annual Joint Conference of the IEEE Computer and Communications Societies, Volume 3, Page(s):1310 - 1319, 2001.
- [122] Injong Rhee, Srinath R. Joshi, "FEC-based Loss Recovery for Interactive Video Transmission," icmcs, p. 9250, 1999 IEEE International Conference on Multimedia Computing and Systems (ICMCS'99) - Volume 1, 1999
- [123] Cai, J.; Qian Zhang; Wenwu Zhu; Chen, C.W., "An FEC-based error control scheme for wireless MPEG-4 videotransmission", Proc In IEEE Wireless Communications and Networking Conference (WCNC 2000), Volume 3, Page(s):1243 - 1247, 2000.
- [124] J. Rosenberg, and H. Schulzrinne, "RFC 2733 : RTP payload format for generic forward error correction", Request for Comments, IETF, December 1999.
- [125] Jianfei Cai, Chang Wen Chen, "FEC-based video streaming over packet loss networks withpre-interleaving", Proc In International Conference on Information Technology: Coding and Computing, On page(s): 10-14, Las Vegas, NV, USA, Apr 2001.
- [126] Abdelhamid Nafaa, Toufik Ahmed, and Ahmed Mehaoua, "Unequal and Interleaved FEC for Wireless MPEG-4 Video Multicast", on Proceedings of IEEE ICC'04, the IEEE International Conference on Communications 2004, vol.3, pp. 1431-1435, Paris, June 20th 2004.
- [127] Wai-Tian Tan, A.Zakhor, "Video multicast using layered FEC and scalable compression", IEEE Transactions on Circuits and Systems for Video Technology, Volume: 11, Issue: 3, On page(s): 373-386, Mar 2001.
- [128] I.Bouazizi, M.Gunes, "Distortion-optimized FEC for unequal error protection in MPEG-4 video delivery", Proc In Ninth International Symposium on Computers and Communications (ISCC 2004), Volume 2, Page(s): 615 - 620, June-July 2004.
- [129] N.Thomos, S.Argyropoulos, N.V.Boulgouris, M.G.Strintzis, "Robust Transmission of H.264/AVC Video using Adaptive Slice Grouping and Unequal Error Protection", Proc In IEEE International Conference on Multimedia and Expo, Page(s):593 - 596, July 2006.
- [130] Chih-Chin Liu; Shu-Chin Su Chen, "Providing unequal reliability for transmitting layered videostreams over wireless networks by multi-ARQ schemes", Proc In International Conference on Image Processing (ICIP 99), Volume 3, Page(s):100 - 104, 1999.
- [131] Fen Hou, Pin-Han Ho, Yongbing Zhang, "Performance analysis of differentiated ARQ scheme for video transmission over wireless networks", Proceedings of the 1st ACM workshop on Wireless multimedia networking and performance modeling, Pages: 1 - 7, Montreal Quebec Canada, 2005.
- [132] Qinqing Zhang, S.A. Kassam, "Hybrid ARQ with selective combining for video transmission overwireless channels", Proc In International Conference on Image Processing, Volume 2, Page(s):692 - 695, Oct 1997.
- [133] Jianwei Wen; Qionghai Dai; Yihui Jin, "Channel-adaptive hybrid ARQ/FEC for robust video transmission over 3G", IEEE International Conference on Multimedia and Expo (ICME 2005), Page(s): 4, July 2005.
- [134] D.G.Sachs, I.Kozintsev, M.Yeung, D.L. Jones, "Hybrid ARQ for robust video streaming over wireless LANs", Proc In International Conference on Coding and Computing, On page(s): 317-321, Las Vegas, NV, USA, Apr 2001.
- [135] Trista Pei-chun Chen, Tsuhan Chen, "Second-generation error concealment for video transport over error prone channels", Proc In International Conference on Image Processing, Volume: 1, On page(s): 1-25- 1-28, 2002.
- [136] Yao Wang, S.Wenger, Jiantao Wen, A.K.Katsaggelos, "Error resilient video coding techniques", Proc In IEEE Signal Processing Magazine, Volume 17, Issue 4, Page(s):61 - 82, Jul 2000.
- [137] D. Hoffman, G. Fernando, and V. Goyal, "RFC 2038 : RTP payload format for MPEG1/MPEG2 video", Request for Comments, IETF, October 1996.

- [138] Y. Kikuchi, et al., "RFC 3016 : RTP payload format for MPEG-4 audio/visual streams", Request for Comments, IETF, November 2000.
- [139] S. Wenger, et al., "RFC 3984 : RTP payload format for H.264 video", Request for Comments, IETF, February 2005.
- [140] K. Fujimoto, S. Ata, and M. Murata, "Statistical analysis of packet delays in the Internet and its application to playout control for streaming applications", IEICE Transactions on Communications, vol. E84-B, pp. 1504-1512, June 2001.
- [141] E. G. Steinbach, N. Farber, and B. Girod, "Adaptive Playout for Low Latency Video Streaming," Proc. International Conference on Image Processing (ICIP-01), Oct. 2001.
- [142] M. Kalman, E. Steinbach, and B. Girod, "Adaptive playout for real-time media streaming, in Proc. IEEE Int. Symp. on Circ. and Syst. May 2002.
- [143] V. Srivastava and M. Motani, "Cross-layer design : a survey and the road ahead", IEEE Communications Magazine, vol.43, no.12, pp.112-119, December 2005.
- [144] M. Van Der Schaar, et al., "Cross-layer wireless multimedia transmission : challenges, principles, and new paradigms", IEEE Wireless Communications Magazine, vol. 12, no. 4, pp. 50-58, August 2005.
- [145] Nasser Sedaghati-Mokhtari, Mahdi Nazm Bojnordi, Nasser Yazdani, "Cross-Layer Design:A New Paradigm", Proc In International Symposium on Communications and Information Technologies (ISCIT '06), On page(s): 183-188, Bangkok, Sept 2006.
- [146] Q. Wang, M. A. Abu-Rgheff, "Cross-Layer Signalling for Next-Generation Wireless Systems," Proc In IEEE Wireless Communication and Networking , New Orleans, Mar. 2003.
- [147] V.T.Raisinghani, S.Iyer, "Cross-layer feedback architecture for mobile device protocol stacks", IEEE Communications Magazine, Volume 44, Issue 1, Page(s): 85 - 92, Jan 2006.
- [148] R.Winter, J.H.Schiller, N.Nikaein, C.Bonnet, "CrossTalk : cross-layer decision support based on global knowledge", in IEEE Communications Magazine, ISSN: 0163-6804, pages 93- 99, Volume: 44, Issue: 1, Jan 2006.
- [149] S.Khan, Y.Peng, E.Steinbach, M.Sgroi, W.Kellerer, "Application-driven cross-layer optimization for video streaming over wireless networks", In IEEE Communications Magazine, ISSN: 0163-6804, pages 122- 130, Volume: 44, Issue: 1, Jan 2006.
- [150] Hai Jiang; Weihua Zhuang; Xuemin Shen, "Cross-layer design for resource allocation in 3G wireless networks and beyond", IEEE Communications Magazine, Volume 43, Issue 12, Page(s): 120 - 126, Dec 2005.
- [151] A.Sali, A.Widiawan, S.Thilakawardana, R.Tafazolli, B.G.Evans, "Cross-Layer Design Approach for Multicast Scheduling over Satellite Networks", 2nd International Symposium on Wireless Communication Systems, ISBN: 0-7803-9206-X, page(s): 701- 705, Sept 2005.
- [152] E.Setton, Taesang Yoo, Xiaoqing Zhu, A.Goldsmith, B.Girod, "Cross-layer design of ad hoc networks for real-time video streaming", IEEE Communications Magazine, Volume 12, Issue 4, Page(s): 59 - 65, Aug 2005.
- [153] Y. Shan and A. Zakhori, "Cross layer techniques for adaptive video streaming over wireless networks", IEEE ICME, vol. 1, pp. 277-280, 2002.
- [154] W.Eberle, B.Bougard, S.Pollin, F.Catthoor, "From myth to methodology: cross-layer design for energy-efficient wireless communication", Proceedings In 42nd Design Automation Conference, Page(s): 303 - 308, Jun 2005.
- [155] C. Schurgers et al., "Modulation scaling for energy aware communication systems," Proc. ISLPED, Huntington Beach, CA, Aug. 2001.
- [156] E. Uysal-Biyikoglu, "Energy-Efficient Packet Transmission over a Wireless Link", ACM/IEEE Trans. Netw. vol. 10, no. 4, Aug. 2002.
- [157] P. A. Chou and Z. Miao, "Rate-Distortion Optimized Streaming of Packetized Media" Microsoft Research tech rep MSR-TR-2001-35, Feb 2001.
- [158] E. Setton, X. Zhu, and B. Girod, "Congestion-Optimized Scheduling of Video over Wireless Ad Hoc Networks," IEEE International Symposium on Volume 4, Page(s): 3531-3534, May 2005.
- [159] A.Ksentini, M.Naimi, A.Gueroui, "Toward an improvement of H.264 video transmission over IEEE 802.11e through a cross-layer architecture", in IEEE Communications Magazine, ISSN: 0163-6804, pages 107- 114, Volume: 44, Issue: 1, Jan 2006.
- [160] Q. Zhang, F.Yang, W.Zhu, "Cross-Layer QoS Support for Multimedia Delivery over Wireless Internet", EURASIP Journal on Applied Signal Processing, vol: 2005:2, pp. 207-219, 2005.

- [161] H. Balakrishnan, V. Padmanabhan, S. Seshan, and R. Katz, "A comparison of mechanisms for improving TCP performance over wireless links", IEEE/ACM Trans. Networking, vol. 5, no. 6, pp. 756-769, 1997.
- [162] L. Benini et al., "A Survey of Design Techniques for System-Level Dynamic Power Mgmt," IEEE Trans. VLSI, vol. 8, no. 3, June 2000.
- [163] T. Simunic et al., "Energy Efficient Design of Portable Wireless Systems" Proc. ISLPED, pp. 49-54, Italy, Aug, 2000.
- [164] R. Zheng and R. Kravets, "On-demand power management for ad hoc networks" Proc. IEEE INFUCOM, San Francisco, USA, 2003.
- [165] A. Acquaviva, T. Simunic, et al., "Remote power control of wireless network interfaces," J. Embedded Comp., vol. 3, 2004.
- [166] Qiong Li, M.van der Schaar, "Providing adaptive QoS to layered video over wireless local area networks through real-time retry limit adaptation", IEEE Transactions on Multimedia, Volume 6, Issue 2, Page(s): 278 - 290, Digital Object Identifier 10.1109/TMM.2003.822792, April 2004.
- [167] Lai-U Choi, W.Kellerer, E.Steinbach, "Cross layer optimization for wireless multi-user video streaming", International Conference on Image Processing 2004 (ICIP'04), Volume: 3, pages: 2047- 2050, October 2004.
- [168] Lai-U Choi, Wolfgang Kellerer, Eckehard Steinbach, "on Cross-Layer Design for Streaming Video Delivery in Multiuser Wireless Environments", EURASIP Journal on Wireless Communications and Networking, Volume 2006, Article ID 60349, pages 1-10, May 2006.
- [169] V.T. Raisinghani, B.S. Iyer, Cross-layer design optimizations in wireless protocol stacks, Computer Comm. Volume 27, pages 720-724, October 2004.
- [170] V. T. Raisinghani and S. Iyer, "ECLAIR: An Efficient Cross-Layer Architecture for Wireless Protocol Stacks," World Wireless Cong., San Francisco, CA, May 2004.
- [171] 4MORE project website <http://4more.av.it.pt/>
- [172] PHOENIX project website <http://www.ist-phoenix.org/>
- [173] ENTHRONED II project website <http://www.ist-enthroned.org/>
- [174] ITU-T Rec. Y.2001, "General Overview of NGN", Dec 2004
- [175] ITU-T REC. Y.2011, "General Principles and General Reference Model of Next Generation Network", Oct 2004
- [176] ETSI ES 282 001, "NGN Functional Architecture", Aug 2005
- [177] Technical Specification Group Services and System Aspects, "IP Multimedia Subsystem (IMS)", Stage 2, V5.15.0, TS 23.228, 3rd Generation Partnership Project, 2006.
- [178] S.-F.Chang, A.Vetro, "Video Adaptation: Concepts, Technologies, and Open Issues", In Proc of the IEEE, Vol. 93, Issue 1, Page(s):148 - 158, Jan. 2005.
- [179] ISO/IEC JTC1/SC29/WG11N6828, "MPEG-7 Overview (version 10)", Oct 2004
- [180] ISO/IEC TR 21000-1:2004, "Information technology -- Multimedia framework (MPEG-21) -- Part 1: Vision, Technologies and Strategy", 2004.
- [181] ISO/IEC 21000-7:2007, Information technology -- Multimedia framework (MPEG-21) -- Part 7: Digital Item Adaptation, 2007.
- [182] ITU-T FG IPTV-R-0014, 2nd FG IPTV meeting, Busan, Korea 16- 20 October 2006.
- [183] DVB - Digital Video Broadcasting Web Site <http://www.dvb.org>
- [184] ETSI - European Telecommunications Standards Institute Web Site <http://www.etsi.org>
- [185] B. Cain, S. Deering, I. Kouvelas, B. Fenner, A. Thyagarajan, "RFC 3376 : Internet Group Management Protocol, Version 3", Request for Comments, IETF, Oct 2002.
- [186] R. Vida, L. Costa, " RFC 3810: Multicast Listener Discovery Version 2 (MLDv2) for IPv6", Request for Comments, IETF, Jun 2004.
- [187] D. Waitzman, C. Partridge, S. Deering, " RFC 1075: Distance Vector Multicast Routing Protocol", Request for Comments, IETF, Nov 1988.
- [188] J. Moy Proteon, "RFC 1584: Multicast Extensions to OSPF", Request for Comments, IETF, Mar 1994.

- [189] T. Pusateri, "RFC 4602: Protocol Independent Multicast - Sparse Mode (PIM-SM)", Request for Comments, IETF, Aug 2006.
- [190] D. Thaler, "RFC 3913: Border Gateway Multicast Protocol (BGMP): Protocol Specification", Request for Comments, IETF, Set 2004.
- [191] VideoLan project web site <http://www.videolan.org>
- [192] Fedora Project web site <http://fedoraproject.org>
- [193] WinTv web site http://www.hauppauge.co.uk/new-fr/site/products/prods_nova_external.html
- [194] libdvbpsi library web site <http://www.videolan.org/developers/libdvbpsi.html>
- [195] 3COM web site <http://www.3com.com>
- [196] Atheros web site <http://www.atheros.com>
- [197] MadWiFi open source project web site <http://madwifi.org>
- [198] FFMPEG open source project web site <http://ffmpeg.sourceforge.net/index.php>
- [199] LiveMedia open source project web site <http://www.live555.com/liveMedia>
- [200] Haratcherev, J. Taal, K. Langendoen, R. Lagendijk and H. Sips, "Automatic IEEE 802.11 rate control for streaming applications", *Wireless Communications and Mobile Computing*, Vol 5, pp.412-437, 2005.
- [201] Kamerman and L. Monteban. WaveLAN-II : "A high-performance wireless LAN for the unlicensed band", AT&T Bell Laboratories Technical Journal, pages 118-133, 1997.
- [202] M. Lacage, M. H. Manshaei, and T. Turletti, "IEEE 802.11 rate adaptation : a practical approach," in *Proceedings of the 7th Symposium on Modeling, Analysis and Simulation of Wireless and Mobile Systems (MSWiM '04)*, pp. 126-134, Venice, Italy, October 2004.
- [203] Jd.P.Pavon, Sunghyun Choi, "Link adaptation strategy for IEEE 802.11 WLAN via received signal strength measurement", *IEEE International Conference on Communications (ICC '03)*, Volume: 2, On page(s): 1108-1113, Anchorage, AK, USA, 2003.
- [204] G. Holland, N. Vaidya and P. Bahl, "A Rate-Adaptive MAC Protocol for Multi-Hop Wireless Networks", in *Proceedings of the 7th annual international conference on Mobile computing and networking (ACM MOBICOM'07)*, Pages: 236 - 251, Rome, July 2001.
- [205] J. C. Bicket, "Bit-rate Selection in Wireless Networks", M.S Thesis, MIT, February 2005.
- [206] B.Furht and O. Marqure, "Chapter 41 in *The Handbook of Video Databases: Design and Applications*", ed.CRC Press, pp. 1041-1078, September 2003.
- [207] Z. Wang, A. C. Bovik, H. R. Sheikh and E. P. Simoncelli, "Image quality assessment: From error visibility to structural similarity," *IEEE Transactions on Image Processing*, vol. 13, no. 4, pp. 600-612, Apr. 2004.
- [208] Pro-MPEG Forum: "Transmission of Professional MPEG-2 Transport Streams over IP Networks", Pro-MPEG Code of Practice #3 Release 2, July 2004.
- [209] E. N. Gilbert, "Capacity of a Burst-noise Channel", *BSTJ*, Vol. 39, pp. 1253-1265, September 1960.
- [210] E. O. Elliott, "Estimates of Error Rates for Codes on Burst-Noise Channels", *BSTJ*, Vol. 42, pp. 1977-1997, September 1963.
- [211] Iperf web site <http://dast.nlanr.net/Projects/Iperf>
- [212] G.Holland, N.Vaidya, P.Bahl, "A rate-adaptive MAC protocol for multi-Hop wireless networks", *Proceedings of the 7th annual international conference on Mobile computing and networking*, pages: 236-251, Rome Italy, 2001.
- [213] Yang Xiao, "IEEE 802.11 performance enhancement via concatenation and piggyback mechanisms", *IEEE Transactions on Wireless Communications*, Volume 4, Issue 5, Page(s): 2182 - 2192, sept 2005.
- [214] J.Tourrilhes, "Packet frame grouping: improving IP multimedia performance over CSMA/CA", *IEEE International Conference on Universal Personal Communications ICUPC'98*, Volume 2, Page(s):1345 - 1349 vol.2, Oct 1998.
- [215] Julien Reichel, Heiko Schwarz, and Mathias Wien, Joint Scalable Video Model (JSVM) 10, Joint Video Team, Doc. JVT-W202, San Jose, CA, USA, April 2007.

A Annexe

A.1 L'en-tête de la couche PHY 802.11

La Figure A-1 présente l'en-tête de la couche physique défini par le standard IEEE 802.11 [3].

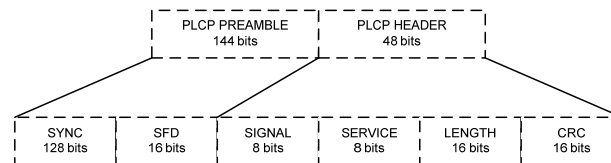


Figure A-1 : L'en-tête physique 802.11

- **SYNC** : Il correspond aux bits de synchronisation entre l'émetteur et le récepteur
- **SFD (Start Frame Delimiter)** : Il indique le début des paramètres de l'en-tête physique
- **SIGNAL** : Il informe sur la modulation utilisée pour la transmission (et la réception) de la trame.
- **SERVICE** : Ce champ est utilisé pour supporter des extensions de débits plus élevés.
- **LENGTH** : Il indique le temps en microsecondes nécessaire pour la transmission de la trame. Ce temps est calculé à partir de la taille de la trame.
- **CRC (Cyclic Redundancy Check)** : Il correspond à la somme de contrôle sur l'en-tête physique.

A.2 L'en-tête de la couche MAC 802.11

La Figure A-2 présente l'en-tête de la couche MAC 802.11 défini dans le standard 802.11 [3]

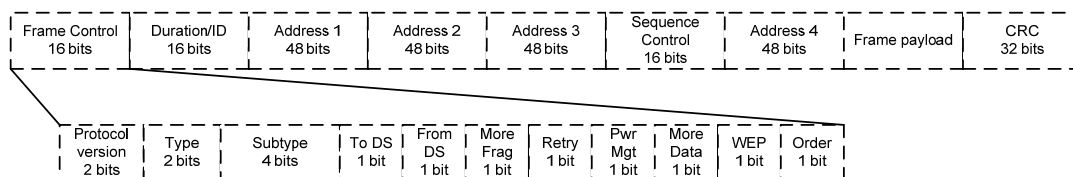


Figure A-2 : L'en-tête MAC 802.11

- **Frame Control** : Ce champ contient les informations suivantes
 - **Protocol Version** : Il indique la version du standard 802.11
 - **Type et Subtype** : Il indique le type de la trame (RTS, CTS, ACK, etc.).
 - **To DS** : Il est positionné à '1' lorsque la trame est destinée au point d'accès. Dans les autres cas, il est positionné à '0'.

- **From DS** : Il est positionné à '1' lorsque la trame provient du point d'accès. Dans les autres cas, il est positionné à '0'.
- **More Fragment** : Il indique la présence d'autres fragments à recevoir lorsqu'il est positionné à '1'. Dans les autres cas, il est positionné à '0'.
- **Retry** : Lorsqu'il est positionné à '1', il indique que la trame est une retransmission. Il permet au récepteur de détecter les trames dupliquées quand l'acquittement est perdu.
- **Pwr Mgt (Power Management)** : Lorsqu'il est positionné à '1', il indique que la station émettrice entre en mode de gestion d'énergie.
- **More data** : Ce bit est utilisé par l'AP en mode de gestion d'énergie pour indiquer à une station que d'autres trames sont en attente d'être transmises.
- **WEP** : Ce bit indique que la trame est cryptée par l'algorithme de chiffrement WEP.
- **Order** : Ce bit indique que les trames sont envoyées suivant la classe de service strictement ordonnée (Strictly-ordered).
- **Duration/ID** : Ce champ possède deux interprétations suivant le type du message. Il peut indiquer l'identifiant d'une station ou le temps d'occupation du canal par la trame.
- **Address 1/2/3/4** : Ces champs donnent les adresses MAC des entités qui prennent part dans une communication. Leurs utilisations dépendent principalement des champs « To DS » et « From DS ». Le Tableau A-1 résume cette utilisation suivant les valeurs de ces deux champs.

To DS	From DS	Address 1	Address 2	Address 3	Address 4
0	0	Destination	Source	Point d'accès	Non renseigné
0	1	Destination	Point d'accès	Source	Non renseigné
1	0	Point d'accès	Source	Destination	Non renseigné
1	1	Emetteur	Récepteur	Destination	Source

Tableau A-1 : L'utilisation des champs « Address 1/2/3/4 » dans l'en-tête MAC 802.11

- **Sequence Control** : Ce champ permet de distinguer les fragments appartenant à une même trame. Il est constitué de deux champs : Le numéro du fragment et le numéro de séquence.
- **CRC** : Il correspond à la somme de contrôle sur la trame complète.

A.3 L'en-tête d'un paquet TS

La Figure A-3 illustre l'en-tête des paquets TS (Transport Stream) défini par le standard MPEG-2 [96].

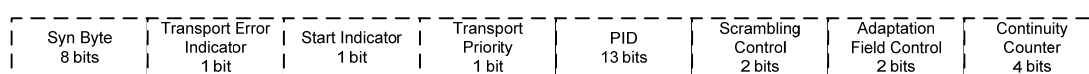


Figure A-3 : L'en-tête d'un paquet TS

- **Syn Byte** : Il possède une valeur fixe 0x47. Il indique le début d'un paquet TS
- **Transport Error Indicator** : Il est positionné à '1' si le paquet est susceptible de contenir des erreurs.
- **Start Indicator** : Il indique le début d'un PES (Packetized Elementary Stream) ou d'un PSI (Program Specific Information) s'il est positionné à '1'.
- **Transport Priority** : Il indique que le paquet TS est prioritaire s'il est positionné '1'.
- **PID** : Il correspond à l'identifiant du paquet TS
- **Scrambling Control** : Il informe sur le cryptage du paquet TS.
- **Adaptation Field Control** : Il indique la présence d'un champ additionnel « Adaptation Field » à la fin de l'en-tête s'il est positionné à '1'.
- **Continuity Counter** : Il correspond au numéro de séquence du paquet TS possédant le PID. Il permet de détecter les paquets perdus, dupliqués ou hors séquence.

A.4 L'en-tête du protocole RTP

La Figure A-4 illustre l'en-tête du protocole RTP défini par le RFC 3550 [92].

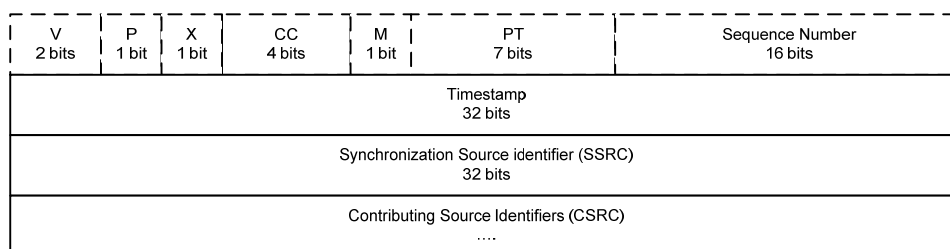


Figure A-4 : L'en-tête du protocole RTP

- **V (Version)** : Il identifie la version du protocole RTP.
- **P (Padding)** : Lorsqu'il est positionné à '1', il indique que le paquet RTP contient des octets de bourrage qui n'appartiennent pas aux données lorsqu'il est positionné à '1'.
- **X (eXtention)** : Il indique la présence d'une extension d'en-tête lorsqu'il est positionné à '1'.
- **CC (CSRC count)** : Il informe sur le nombre des identificateurs CSRC à la fin de l'en-tête.
- **M (Marker)** : L'interprétation de ce bit est définie par l'application à travers des profils particuliers.
- **PT (Payload Type)** : Il identifie le format des données transportées par les paquets RTP.
- **Sequence Number** : Il correspond au numéro de séquence du paquet RTP.
- **Timestamp** : Il donne l'instant d'échantillonnage du premier octet dans le paquet RTP. Cet instant doit être dérivé d'une horloge qui augmente de façon monotone et linéaire dans le temps.

- **Synchronization Source Identifier (SSRC)** : Il identifie la source de synchronisation. Il est choisi aléatoirement et il doit être unique.
- **Contributing Source Identifiers (CSRC)** : La liste des CSRCs identifie les sources SSRC qui ont contribué à l'obtention des données transportées par les paquets RTP.

Liste des abréviations

AC	Access Categories
ACC	Algorithme de Contrôle de Congestion
AIFS	Arbitrary IFS
AIMD	Additive-Increase / Multiplicative-Decrease
AP	Access Point
API	Application programming Interface
APP	Application
ARQ	Automatic Repeat reQuest
BER	Bit Error rate
BT	Backoff Time
CRC	Cyclic Redundancy Check
CSMA/CA	Carrier Sence Multiple Access/ Collision Avoidance
CSMA/CD	Carrier Sence Multiple Access/ Collision Detection
CW	Contention Windows
DBPSK	Differential Binary Phase Shift Keying
DCCP	Datagram Congestion Control Protocol
DCF	Distributed Coordination Function
DIA	Digital Item Adaptation
DiffServ	Differentiated services
DIFS	Distributed IFS
DQPSK	Differential Quadrature PSK
DSCP	DiffServ Code Point
DSSS	Direct Sequence Spread Spectrum
DVB-T	Digital Video Broadcast – Terrestrial
DVB-S	Digital Video Broadcast – Satellite
ECN	Explicit Congestion Notification
EDCA	Enhanced Distributed Channel Access
ETSI	European Telecommunications Standards Institute
FEC	Forward Error Correction
FHSS	Frequency Hopping Spread Spectrum
GOP	Groupe of Pictures
IEEE	Institute of Electrical and Electronics Engineers
IETF	Internet Engineering Task Force
IFS	Inter Frame Spacing
IGMP	Internet Group Multicast Protocol
IMS	Industrial, Scientific, and Medical
IMS	Internet Multimedia Subsystem
IntServ	Integrated services
IP	Intenet Protocol
ITU-T	International Telecommunication Union - Telecommunications
JVT	Joint Video Team
LLC	Logical Link Control
LoD	Live on Demand
MAC	Media Access Control
MPEG	Moving Pictures Expert Group
MPLS	Multi Protocol Label Switching
MTU	Maximum Transit Unit

NGN	Next Generation Network
PHY	Physical
PID	Packet IDentifier
PIM	Protocol Independent Multicast
PLCP	Physical Layer Convergence Protocol
PMT	Program Map Table
PSI	Program Specific Information
PSNR	Peak Signal-to-Noise Ratio
QoS	Quality of Service
RCA	Rate Control Algorithm
RFC	Request For Comment
RS	Reed Salomon
RSSI	Received Signal Strength Indicator
RSVP	ReServation Protocol
RTP/RTCP	Real Time Protocol/Real Time Control Protocol
RTSP	Real Time Streaming Protocol
RTT	Round Tripe Time
SAP	Session Announcement Protocol
SCTP	Stream Control Transmission Protocol
SDP	Session Description Protocol
SDT	Service Description Table
SIFS	Short IFS
SIR	Signal-to-Interference Ratio
SNR	Signal-to-Noise Ratio
SSI	Signal Strength Indicator
SSIM	Structural SIMilarity
STA	Station
SVC	Scalable Video Coding
TCP	Transmission Control Protocol
TFRC	TCP-friendly Rate Control
TGN	Task Group N
TS	Traffic Stream
TS	Transport Stream
TSPEC	Traffic Specification
TVSP	TV stream provider
TXOP	Transmission Opportunity
UCD	Universal Constraint Description
UDP	Usage Datagram Protocol
UED	Usage Environment Description
UMA	Universal Multimedia Access
VCEG	Video Compression Expert Group
VFG	Video Frame grouping
VoD	Video on Demand
WLAN	Wireless Local Area Network
XLAG	Cross-Layer Adaptation Gateway
XLAVS	Cross Layer Adaptive Video Streaming
XLDP	Cross Layer Decision Point