

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
CHAPITRE 1 LES RÉSEAUX WIMAX	4
1.1 La normalisation du WiMAX	4
1.1.1 La description.....	4
1.1.2 Les normes	5
1.1.3 Le principe de fonctionnement	6
1.1.4 Le positionnement du WiMAX parmi les autres réseaux sans fil	7
1.2 Caractéristiques techniques du WiMAX	8
1.2.1 La topologie	8
1.2.2 La couche physique « PHY ».....	10
1.2.3 La couche MAC.....	15
1.2.4 La Qualité de service	19
1.2.5 Conclusion	22
CHAPITRE 2 ÉTAT DE L'ART.....	23
2.1 Introduction.....	23
2.2 Survol des algorithmes d'ordonnancement existants.....	23
2.3 Synthèse et limites des solutions existantes.....	36
CHAPITRE 3 ALGORITHME DE COURTOISIE	38
3.1 Introduction.....	38
3.2 Description de l'algorithme de courtoisie.....	39
3.3 Conditions d'application.....	41
3.3.1 Condition 1.....	41
3.3.2 Condition 2.....	42
3.3.3 Condition 3.....	43
3.3.4 Condition 4.....	43
3.4 Analyse mathématique.....	44
3.4.1 Cas de deux files d'attente	44
3.4.2 Cas de n files d'attente	59
3.5 Importance de l'analyse mathématique	63
3.6 Structure de l'algorithme de courtoisie.....	63
3.7 Conclusion	68
CHAPITRE 4 SIMULATIONS ET RÉSULTATS	70
4.1 Introduction.....	70
4.2 Approche adaptée dans les simulations	70
4.2.1 Description du modèle PQ considéré.....	72

4.2.2	Description du modèle WFQ considéré.....	73
4.3	Scénarios de tests Matlab.....	74
4.3.1	Scénario 1 : scénario de référence	74
4.3.2	Scénario 2: l'effet de la diminution de λ_1	83
4.3.3	Scénario 3: Étude de l'effet de l'augmentation de λ_1 , λ_2 et $D_{\max_rtPS}$	91
4.3.4	Scénario 4: Étude de l'effet de l'augmentation de λ_2	98
4.3.5	Scénario 5: Étude de l'effet de λ_1 et λ_2 et de la taille de l'échantillon	106
4.3.6	Scénario 6 : Étude de l'impact de l'augmentation de R_1	114
4.4	Apport de l'algorithme de courtoisie	123
4.5	Conclusion	126
CONCLUSION		128
RECOMMANDATIONS		131
LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES.....		133

LISTE DES TABLEAUX

	Page
Tableau 1.1	Les principales normes du standard 802.16.....6
Tableau 4.1	de la simulation, relatifs au scénario de référence75
Tableau 4.2	Résultats numériques pour le scénario 178
Tableau 4.3	Résultats numériques pour le scénario 2.....84
Tableau 4.4	Résultats numériques pour le scénario 3.....92
Tableau 4.5	Résultats numériques correspondants au scénario 4.....99
Tableau 4.6	Résultats numériques pour le scénario 5.....107
Tableau 4.7	Résultats numériques relatifs pour le scénario 6.....115
Tableau 4.8	Apport de CPQ en fonction de taux de trafic nrtPS.....124

LISTE DES FIGURES

	Page
Figure 1.1	Exemple d'un réseau WiMAX avec les deux variantes fixe et mobile.5
Figure 1.2	Positionnement de la couverture du WIMAX parmi les réseaux sans fil. ...8
Figure 1.3	Topologie PMP (Point à MultiPoint).9
Figure 1.4	Topologie Maillée.9
Figure 1.5	Les couches du standard IEEE 802.16.11
Figure 1.6	Couverture et modulation dans le WIMAX.13
Figure 1.7	Multiplexage TDD et FDD dans les réseaux 802.16.14
Figure 1.8	Format du MAC PDU.17
Figure 1.9	En-tête MAC générique.18
Figure 1.10	En-tête MAC de demande de bande passante.18
Figure 2.1	Modèle analytique relatif à l'algorithme FEQ.24
Figure 2.2	Discrimination des paquets au niveau de la BS pour le lien descendant. ...26
Figure 2.3	Algorithmes d'ordonnancement Channel-aware QoS.30
Figure 2.4	Délai d'attente dans le lien commun à UGS et nrtPS.31
Figure 2.5	Taux de perte de paquets dans le lien commun à UGS et nrtPS.31
Figure 3.1	Système de file d'attente M/G/1 courtoisie : deux queues.45
Figure 3.2	Diagramme d'état du système de files d'attente M/G/1 avec PQ.47
Figure 3.3	Système de file d'attente M/G/1 avec courtoisie (n files d'attente).....58
Figure 3.4	Algorithme de courtoisie pour le réseau WiMAX : Deux classes de service.64
Figure 3.5	Fonction Calcul_Priorité.66

Figure 3.6	Fonction Calcul _Seuils.	67
Figure 4.1	Exemple d'un modèle de Priority Queuing.	73
Figure 4.2	Exemple d'un système M/G/1 pour un réseau WiMAX fixe : deux queues.	74
Figure 4.3	Longueur de la file d'attente de la VoIP vs. le temps.	79
Figure 4.4	Longueur de la file d'attente de la VoIP vs. le temps.	80
Figure 4.5	Longueur de la file d'attente de données vs. le temps.	81
Figure 4.6	Délai d'attente dans la Queue rtPS vs. le temps.	82
Figure 4.7	Délai d'attente dans la file d'attente nrtPS vs. le temps.	83
Figure 4.8	Longueur de la file d'attente de la voix pour le scénario 2.	85
Figure 4.9	Longueur de la file d'attente nrtPS pour le scénario 2.	86
Figure 4.10	Délai d'attente dans la file rtPS pour le scénario 2.	87
Figure 4.11	Délai d'attente dans la file nrtPS (scénario 2).	88
Figure 4.12	Délai d'attente dans la file nrtPS : Temps = [0.02, 0.03] sec.	89
Figure 4.13	Pertes de paquets nrtPS pour le scénario 2.	90
Figure 4.14	Longueur de la file d'attente rtPS pour le scénario 3.	93
Figure 4.15	Longueur de la file de données pour le scénario 3.	94
Figure 4.16	Délais d'attente dans la file rtPS pour le scénario 3.	95
Figure 4.17	Délai d'attente dans la file nrtPS pour le scénario 3.	97
Figure 4.18	Longueur de la file d'attente rtPS (scénario 4).	100
Figure 4.19	Longueur de la file d'attente nrtPS pour le scénario 4.	101
Figure 4.20	Délai d'attente des paquets rtPS (scénario 4).	103
Figure 4.21	Délai d'attente dans la file nrtPS (scénario 4).	104
Figure 4.22	Délai d'attente dans la file nrtPS : Figure partielle.	104
Figure 4.23	Pertes de paquets nrtPS (scénario 4).	105

Figure 4.24	Longueur de la file d'attente rtPS (scénario 5).	109
Figure 4.25	longueur de la file d'attente nrtPS (scénario 5).	110
Figure 4.26	longueur de la file d'attente nrtPS : résultats partiels.	110
Figure 4.27	Délai d'attente dans la file de la voix (scénario 5).	111
Figure 4.28	Délai d'attente dans la file nrtPS (scénario 5).	113
Figure 4.29	Pertes de paquets de la classe nrtPS (scénario 5).	114
Figure 4.30	Longueur de la file d'attente rtPS (scénario 6).	116
Figure 4.31	Longueur de la file d'attente rtPS : Figure partielle.	117
Figure 4.32	Longueur de la file d'attente nrtPS (scénario 6).	118
Figure 4.33	Longueur de la file d'attente nrtPS : Portion du graphe global.	119
Figure 4.34	Délai d'attente dans la file rtPS (scénario 6).	120
Figure 4.35	Délai d'attente dans la file rtPS (scénario 6).	120
Figure 4.36	Délai d'attente dans la file nrtPS (scénario6).	121
Figure 4.37	Délai d'attente dans la file nrtPS : Figure partielle.	122
Figure 4.38	Nombre de paquets perdus vs. Le temps : R=2 et R=6.	123
Figure 4.39	Apport de l'algorithme de courtoisie.	125

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

16 QAM	Quadrature Amplitude Modulation à 16 états
64 QAM	Quadrature Amplitude Modulation à 64 états
BE	Best Effort
BR	Bandwidth Request
BS	Base Station
CBR	Constant Bit Rates
CDMA	Code Division Multiple Access
CI	CRC Indicator
CID	Connection Identifier
CINR	Carrier to Interference and Noise Ratio
CRC	Contrôle par Redondance Cyclique
CPQ	Courtesy Priority Queuing
DFPQ	Deficit Fair Priority Queue
DRR	Deficit Round Robin
Dmax_rtPS	Délai d'attente maximal des paquets rtPS
EC	Encryptions Control
EDD	Earliest Due Date
EDF	Earliest Deadline First
EDGE	Enhanced Data GSM Environment
EKS	EncryptionKeySequence

ertPS	Extended real-time Polling Services
FDD	Frequency Division Duplex
FIFO	Fisrt In First Out
FQ	Fair Queuing
GPRS	General Packet Radio Service
GSM	Global System for Mobile communications
HCS	Header Control Sequence
HT	Header Type
IWFQ	Ideal Weighted Fair Queuing
LEN	LENgth
LOS	Line Of Sight
MAC	Media Access Control
MAC-PDU	Media Access Control Packet Data Unit
MRR	Minimum Reserved Rate
mSIR	maximum Signal-to-Interference Ratio
NLOS	Non Line Of Sight
nrtPS	non real-time Polling Services
PHY	Couche PHYsique
PMP	Point-Multipoint
PQ	Priority Queuing
OFDM	Orthogonal Frequency-Division Multiple
OFDMA	Orthogonal Frequency-Division Multiple Access
OSQ	Opportunistic Stable Queue

QoS	Qualité de service
QPSK	Quadrature Phase Shift Keying
RED	Random Early Detection
RS	Reed Solomon
RSV	Reserved = 0 ou 1
rtPS	real-time Polling Services
SC2	Single carrier
SCDS	Service Classe Downlink Scheduling
SDU	Service Data Unit
SS	Subscriber Station
SSCS	Service Specific Convergence Sublayer
TDD	Time Division Duplex
TRS	Temporary Removal Scheduler
UGS	Unsolicited Grant Service
UMTS	Universal Mobile Telecommunications System
VoIP	Voice over IP
VPN	Virtual Private Network
WFQ	Weighted Fair Queuing
WiFi	Wireless Fidelity
WiMAX	Worldwide Interoperability for Microwave Access
WLAN	Wireless Local Area Network
WMAN	Wireless Metropolitan Network

WPAN	Wireless Personal Area Network
WRR	Weighted Round Robin
WWAN	Wireless Wide Area Network

INTRODUCTION

Le besoin de communication et d'échange d'informations augmente au fil des jours, autant en milieux urbains qu'en zones rurales. Le recours à l'Internet est le premier moyen qui assure un transfert de données rapide et satisfaisant, d'autant plus, s'il s'agit d'une connexion filaire. Pour des contraintes géographiques et économiques, cette dernière technologie ne peut être déployée partout dans le monde. Une solution à ce problème a heureusement pu voir le jour, il s'agit du WIMAX « Worldwide Interoperability for Microwave Access » avec ses deux variantes fixe et mobile dont les normes sont IEEE802.16d et IEEE802.16e, pour les types fixe et mobile respectivement.

Les réseaux IEEE802.16, une technologie en pleine expansion, promettent un apport très encourageant pour le monde des télécommunications, particulièrement pour les réseaux sans fil. Leur débit est assez important autant qu'en état fixe qu'en déplacement, ainsi que pour la portée qui peut atteindre les 50 km. Il est vrai que la bande passante attribuée à la topologie Point-Multipoint est coûteuse, car elle est soumise à des restrictions réglementaires, mais la variante maillée permet de contourner ce problème, reste que la présence des interférences dans cette topologie ne l'a pas mis à l'abri des critiques.

Quoique les réseaux IEEE802.16 sont conçus principalement pour la desserte d'Internet haut débit dans les zones dépourvues d'accès aux réseaux filaires, son déploiement se retrouve également dans les zones urbaines, avec ses différents services, notamment la VoIP, le WEB, le FTP, la vidéo-conférence et la messagerie, chacun appartient à une classe de trafic distincte et exige des conditions différentes en termes de besoin en bande passante et de Qualité de Service (QoS).

Afin de satisfaire les divers types d'applications, la norme a défini quatre classes de qualité de service, à savoir UGS, rtPS, nrtPS et BE. Une cinquième a été conçue particulièrement pour la version mobile, dont le nom est ertPS. Cette politique a permis d'appliquer la

différenciation de services. Certains trafics ont des critères de QoS plus importants. Ils ont donc une plus haute priorité, tandis que d'autres ont moins d'exigence. Les tolérances des différents services ne sont pas forcément identiques. Particulièrement, les trafics temps réel ne tolèrent pas les délais et la gigue, les trafics non temps réel ne sont pas affectés par ces deux facteurs, mais restent sensibles aux taux de pertes de paquets.

La garantie de la QoS devient donc capitale dans les systèmes IEEE802.16, afin d'assurer leur performance. Par conséquent, la gestion de la bande passante reste un défi pour eux, surtout en présence de plusieurs types de connexions, à savoir les appels courants, les nouveaux appels, ainsi que les handoff pour la version mobile. Le WiMAX devrait ainsi appliquer les mécanismes de contrôle d'admission des appels afin d'optimiser la gestion des ressources de la bande passante. Les tâches accomplies par ces mécanismes sont de plus en plus complexes quand les clients WiMAX sont des réseaux WLAN, dans ce cas la détermination de la quantité de la bande passante exigée va dépendre non seulement du type d'applications correspondantes, mais aussi du nombre d'utilisateurs associés à ce WLAN.

La coexistence de plusieurs types de trafics dans un même réseau 802.16 requiert l'application d'une politique de gestion de file d'attente qui permettra le partage des ressources de la bande passante selon les exigences de chaque type d'application. Bien que le standard 802.16 définisse les classes de services assurant la différenciation des types de trafic selon la priorité et les exigences de chacun, il ne décrit pas un système d'ordonnancement pour les liens montant (uplink) et descendant (downlink). Plusieurs recherches ont été menées dans ce sens, dans le but de concevoir un système performant de gestion de trafic dans les réseaux 802.16. Toutefois, ces processus dévoilent quelques faiblesses, vu que la majorité favorise le trafic rtPS par rapport au nrtPS.

L'objectif de notre travail est de garantir une meilleure QoS de service pour les réseaux WiMAX en présence de plusieurs types de trafic, en assurant une gestion équitable des ressources de la bande passante et en optimisant le service du trafic moins prioritaire, sans affecter la QoS des services de haute priorité. L'apport de notre travail consiste à réduire le

délai d'attente et le taux de perte de paquets des données, quand celles-ci sont servies conjointement avec la VoIP.

Le présent document est organisé de telle sorte que le chapitre 1 va résumer la norme et les caractéristiques techniques du WiMAX, le chapitre suivant va présenter l'état de l'art, dans lequel nous allons présenter les solutions les plus importantes qui traitent le problème d'ordonnancement dans les réseaux sans fil en général et dans les réseaux WiMAX en particulier. Par la suite, le chapitre 3 va exposer notre solution, qui sera validée par simulation dans le chapitre 4. La partie conclusion et recommandations va conclure ce travail.

CHAPITRE 1

LES RÉSEAUX WIMAX

1.1 La normalisation du WiMAX

1.1.1 La description

Le WIMAX (Worldwide Interportability Microwave Access) est un réseau hertzien, haut débit, large bande qui couvre un rayon de plusieurs kilomètres. Il est normalisé par l'organisme IEEE sous la norme 802.16. Son objectif principal est de fournir une connexion Internet haut débit aux zones dépourvues d'accès aux réseaux filaires à cause des contraintes économiques ou géographiques. De nos jours, il existe deux types de réseaux WiMAX, à savoir :

Le WIMAX fixe : Dont la norme est la IEEE 802.16d. Il a été conçu pour un usage fixe avec une petite antenne d'abonné placée sur un point d'une certaine hauteur, tel qu'un toit, de la même manière qu'une antenne TV, ou directement sur le PC. Ce type de réseau opère dans une bande de fréquence allant de 2 à 11 GHz. Son débit théorique est de l'ordre de 70 Mb/s et son rayon de couverture est de 50 km.

Le WIMAX mobile : Sa norme est la 802.16e. Son objectif est d'autoriser les abonnés mobiles, une communication continue, en basculant d'une antenne émettrice à une autre, donc d'une cellule à une autre. Ce réseau opère dans les bandes de fréquences allant de 2 à 6 GHz et permet de préserver la connexion lorsque l'on est en déplacement dans la couverture du réseau avec une vitesse allant jusqu'à 150 km/h dans des conditions idéales qui se résument par l'absence d'obstacles.

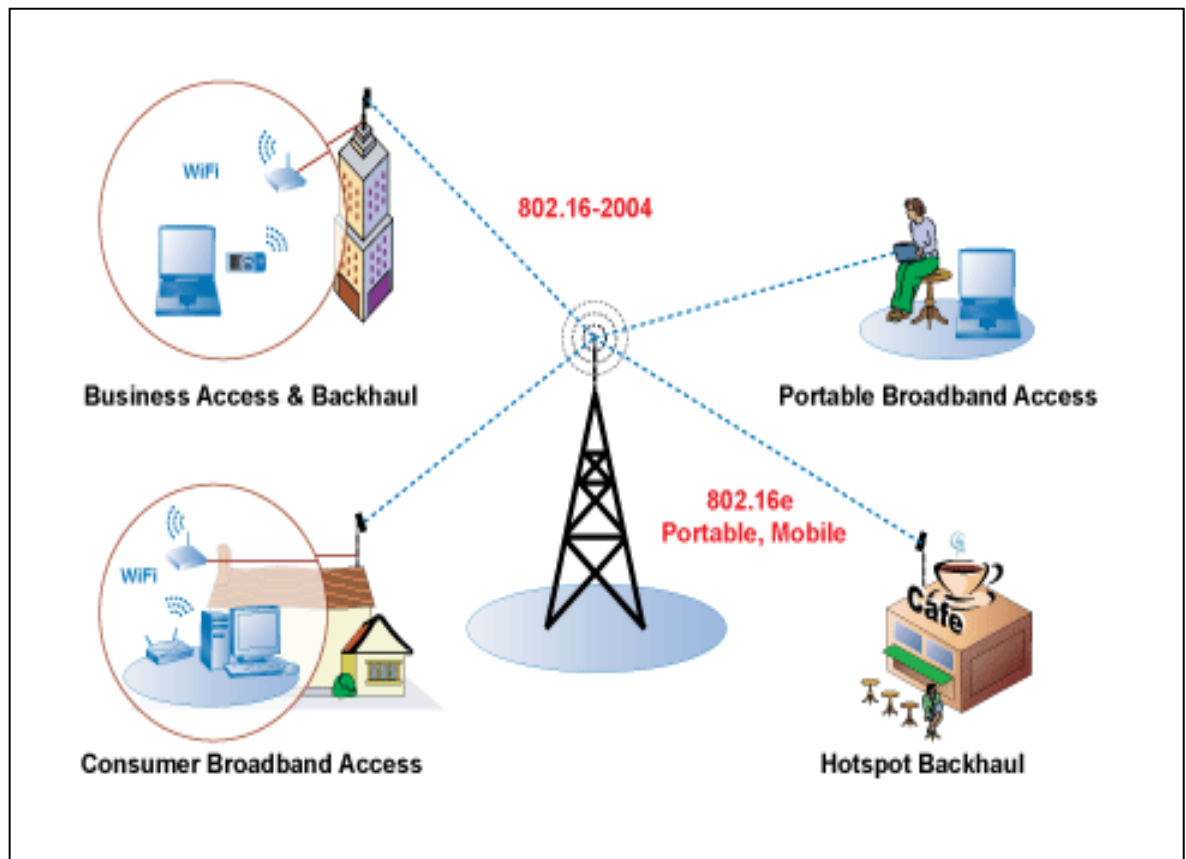


Figure 1.1 Exemple d'un réseau WiMAX avec les deux variantes fixe et mobile.

Tirée de <http://molecularvoices.molecular.com/2007/wimax-redefining-connectivity>

Figure 1.1 illustre un exemple d'un réseau WiMAX avec ses deux variantes, à savoir fixe et mobile. Tel que le montre la figure, ce réseau se compose principalement d'une station de base, qui joue le rôle d'un nœud émetteur, et des stations réceptrices qui jouent le rôle de clients WiMAX. Nous allons ultérieurement présenter le principe de fonctionnement d'un tel type de réseau.

1.1.2 Les normes

Les normes du WIMAX sont en évolution continue. Le Tableau 1.1 résume les normes les plus importantes. Dans ce tableau, on repère les deux normes principales, à savoir la 802.16d

appelée aussi 802.16-2004, et la 802.16 e. La première concerne la version fixe du WiMAX et la seconde concerne le type mobile.

Tableau 1.1 Les principales normes du standard 802.16

Standard	Description
IEEE 802.16	Réseaux sans fil métropolitains opérant sur les bandes de fréquence supérieures à 10 GHZ
IEEE 802.16a	Amendement du standard 802.16 pour les fréquences de 2 à 11 GHZ
IEEE 802.16c	Concerne les réseaux utilisant les fréquences allant de 10 à 66 GHZ
IEEE 802.16d (WiMAX fixe)	Révision des standards 802.16, 802.16a et 802.16c
IEEE 802.16e (WiMAX mobile)	Introduction de la solution de mobilité des clients du réseau WiMAX au standard

1.1.3 Le principe de fonctionnement

Le principe de fonctionnement du WiMAX est simple : Une station émettrice (station de base) émet des ondes radio (hertziennes), dans la bande de fréquence de 2,5 GHZ (3,5 GHZ en Europe), qui sera captée par plusieurs antennes d'abonnés, ainsi que par d'autres stations WiMAX conçues pour jouer le rôle de relais.

Dans un système WiMAX, la station de base est connectée au réseau public en utilisant la fibre optique, le câble, la liaison à ondes radio, ou n'importe autre connexion point à point haute vitesse, connu sous le nom de *backhaul*. Dans certains cas comme dans les réseaux de topologie maillée, la connexion PMP (Point à Multipoint) est aussi utilisée comme *backhaul*, (Pareek 2006)

La station de base sert ses abonnés en utilisant des connexions point à multipoints. Elle utilise la couche MAC pour leur allouer des liens montants, entre les stations clients SSs et la BS, et des liens descendants, reliant la BS aux SSs, selon les besoins en bande passante. Les SSs peuvent représenter un seul utilisateur, comme elles peuvent former un réseau sans fil ou filaire.

Le mode Ligne Of Sight (LOS) utilise une antenne qui pointe directement la station de base du WiMAX. Dans ce cas, les hautes fréquences sont utilisées. Celles-ci peuvent atteindre les 66 MHz où il y a moins d'interférences et plus de bande passante (Pareek 2006).

Un signal émis par une station de base peut franchir de petits obstacles, tels que les maisons et les arbres. On parlera alors d'une communication NLOS (Non Line Of Sight). Une communication NLOS connaît une diminution en termes de débit, qui pourra atteindre les 20Mbits/s. En outre, les grands obstacles tels que les collines et les grands immeubles ne peuvent malheureusement pas être franchis par les signaux WiMAX. Les connexions LOS sont donc plus puissantes et plus stables que les connexions NLOS, qui connaissent un taux d'erreur plus élevé.

1.1.4 Le positionnement du WiMAX parmi les autres réseaux sans fil

Les réseaux sans fil ne sont pas tous d'un seul et même type, ils sont classés par catégorie, par rapport à un ensemble de caractéristiques communes, tel que le débit, la portée et la bande de fréquence dans laquelle ils opèrent.

La Figure 1.2 représente la position de la couverture du WIMAX parmi les autres catégories des réseaux sans fil, spécialement par rapport au IEEE802.11 et au IEEE802.20. Le WIMAX, étant classé comme un réseau sans fil métropolitain, WMAN, se positionne entre les réseaux étendus (IEEE802.20) et WLAN (IEEE 802.11). Notons que les cercles représentent les couvertures des différents réseaux sans fil.

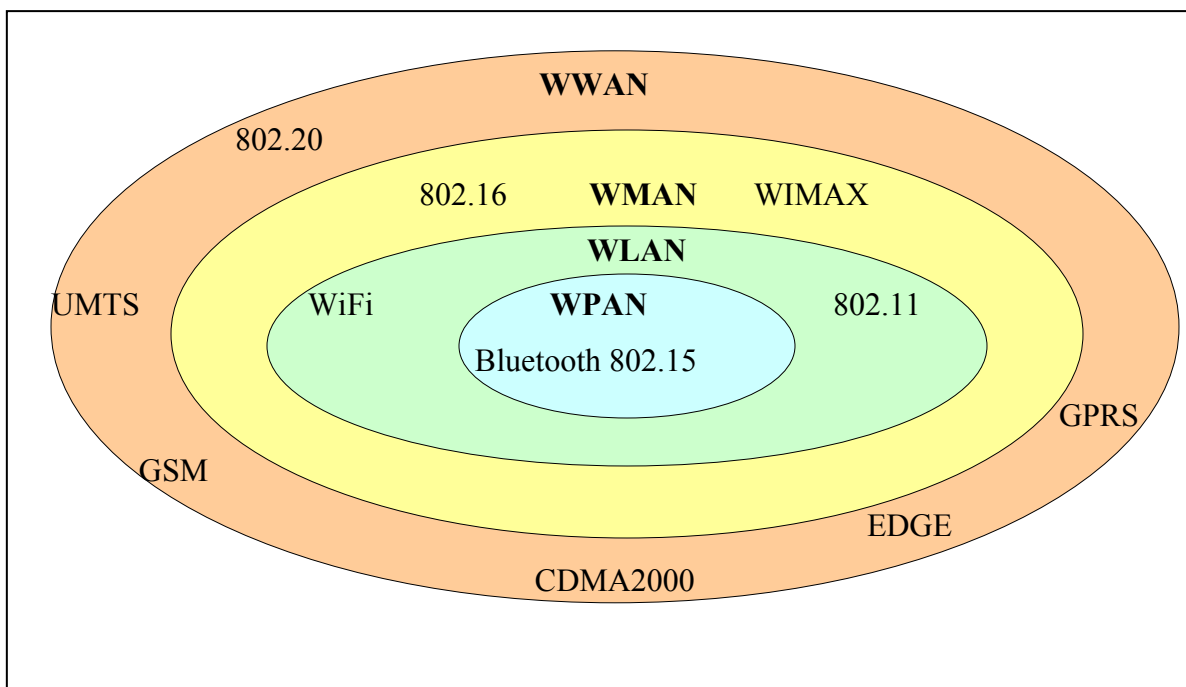


Figure 1.2 Positionnement de la couverture du WIMAX parmi les réseaux sans fil.

1.2 Caractéristiques techniques du WiMAX

1.2.1 La topologie

Deux topologies peuvent être définies pour un réseau WIMAX : La topologie en étoile (Figure 1.3) ou Point-MultiPoints (PMP), et la topologie maillée (Figure 1.4).

La différence principale entre les deux modes, est que dans la topologie PMP le trafic ne peut avoir lieu qu'entre la station de base (BS) et ses stations réceptrices (SSs), alors que dans la topologie maillé, les SSs peuvent également échanger de l'information entre elles.

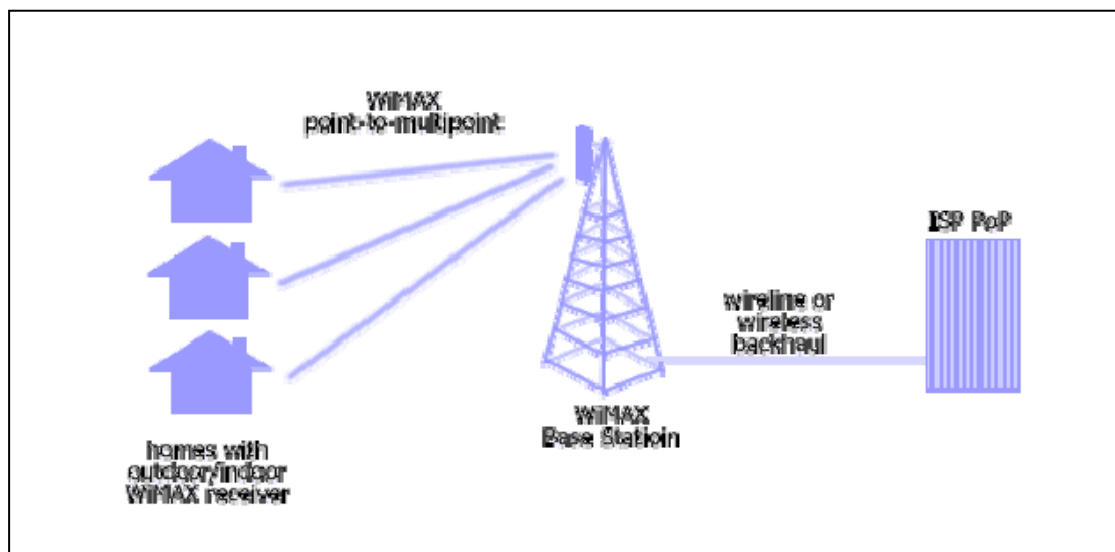


Figure 1.3 Topologie PMP (Point à MultiPoint).

Tirée de <http://www.conniq.com/images/PMP-WiMAX-residential.gif>

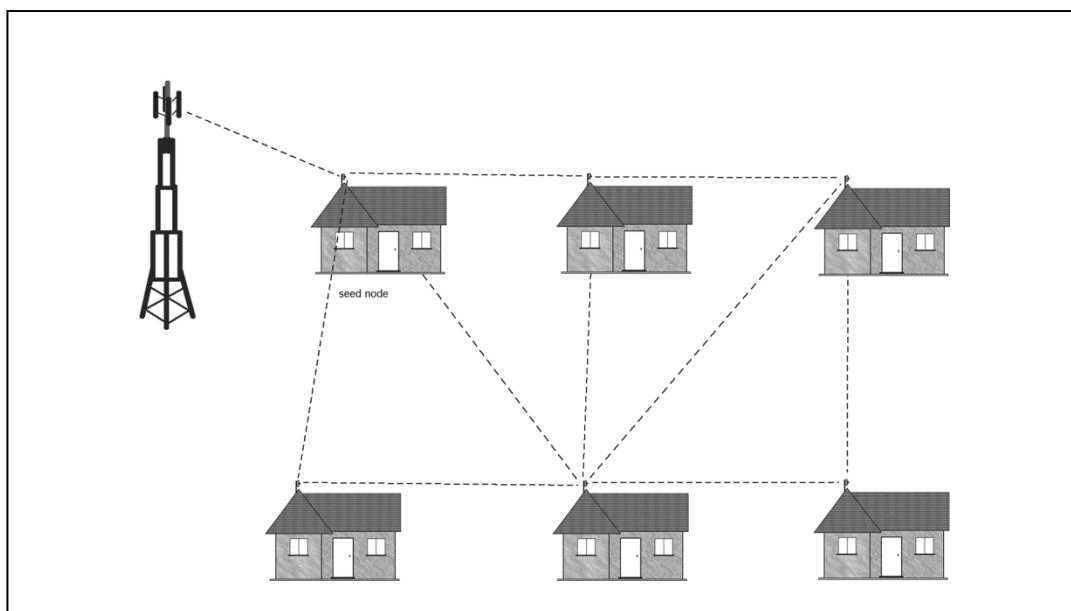


Figure 1.4 Topologie Maillée.

Tirée du site <http://www.emeraldinsight.com/fig/2720100203002.png>



Selon (Nuaymi 2007), dans la topologie maillée, chaque station peut créer sa propre communication avec n'importe quelle autre station dans le réseau, sans qu'elle soit obligatoirement une station de base. Ainsi, la couverture de cette dernière peut devenir plus importante, selon le nombre de sauts permis menant vers la dernière station réceptrice du réseau.

Il existe deux types de réseau maillé : le réseau maillé complet (Full Mesh), où chaque nœud a une liaison avec tous les autres nœuds du même réseau, et le réseau maillé partiel (Partial Mesh) où quelques nœuds seulement ont des liaisons avec tous les autres nœuds du réseau, alors que les autres nœuds ont seulement un ou deux liaisons dans tout le système, (Pareek 2006).

1.2.2 La couche physique « PHY »

La couche physique, PHY du réseau WiMAX (Figure 1.5), est responsable de l'établissement des connexions physiques entre les deux parties qui veulent communiquer, souvent dans les deux sens : lien descendant et lien montant. Elle s'occupe aussi de la transmission des séquences de bits, de la définition du type du signal utilisé, de la modulation et de la démodulation, ainsi que de la puissance de transmission (Nuaymi 2007).

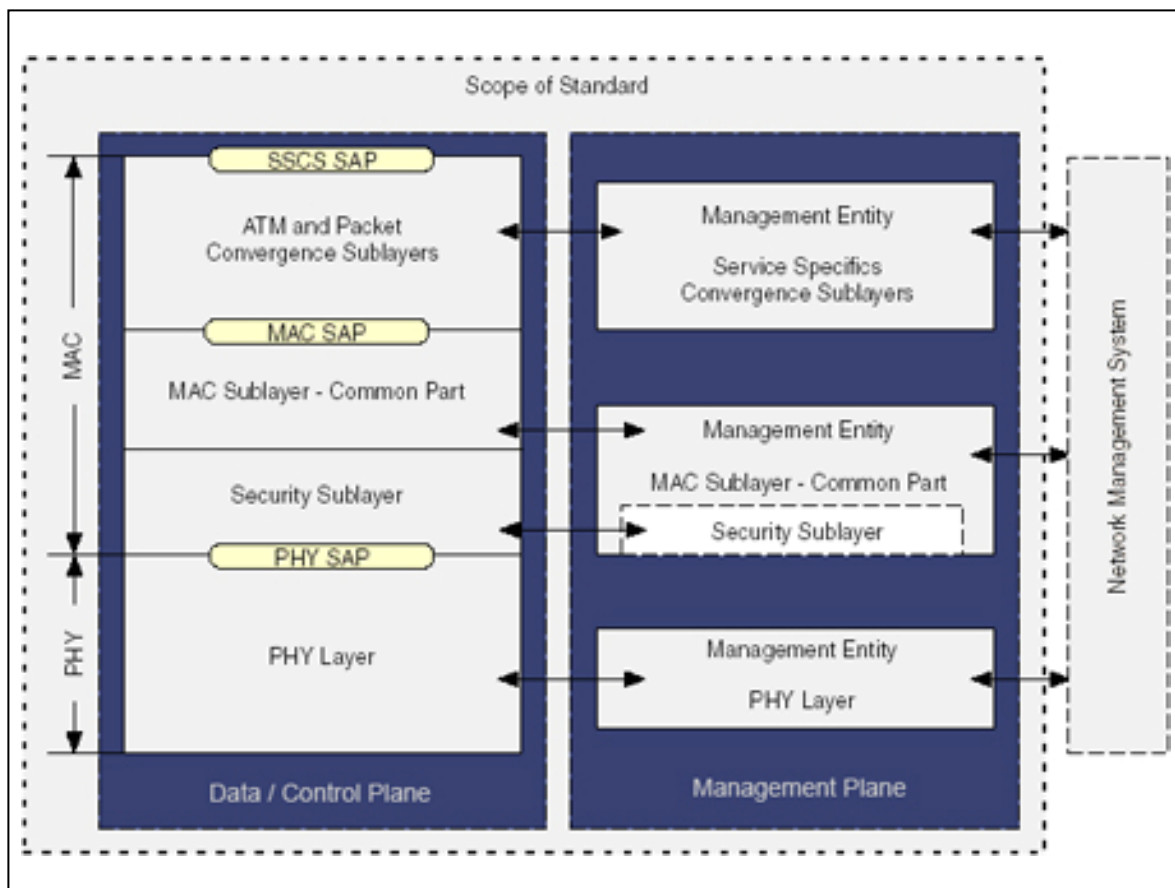


Figure 1.5 Les couches du standard IEEE 802.16.

Tirée du site http://www.supinfo-projects.com/fr/2005/les_reseaux_wimax/2/

Les types de la couche PHY

On distingue trois types de couches physiques :

- **WirelessMAN-SC2** : utilise un format de modulation avec une seule porteuse ;
- **WirelessMAN-OFDM** : utilise un multiplexage orthogonal à division de fréquence avec 256 points de transformation. L'accès à cette couche physique s'effectue en TDMA ;
- **WirelessMAN-OFDMA** : utilise un multiplexage orthogonal à division de fréquence avec 2048 points de transformation. Ce qui permet de supporter de multiples récepteurs.

Les bandes de fréquences utilisées par le WiMAX

Le standard 802.16 utilise les bandes de fréquences suivantes :

- **De 10 à 66 GHZ :** Dans ce cas le LOS (Line of Sight) est exigé, vu que la longueur d'onde est courte et le risque d'atténuation se présente. Cette bande de fréquence offre l'avantage d'avoir des débits élevés. Le type de la couche physique utilisé dans ces fréquences est le WirelessMan-SC. Celle-ci supporte les deux types de duplexage FDD et TDD.
- **De 2 à 11 GHZ :** Le mode NLOS est supporté dans cette bande de fréquence. Ces fréquences utilisent les trois types de couches PHY, à savoir la WirelessMAN-SC2, la WirelessMAN-OFDM et la WirelessMAN-OFDMA.
- **5.7 GHZ :** Contrairement aux autres bandes de fréquences, celle-ci n'exige pas de licence, ce qui permet d'utiliser cette technologie en toute liberté. Cependant, cette liberté d'utilisation de cette catégorie de bande fréquences cause des interférences.

Les modulations

La couche PHY utilise les modulations QPSK (Quadrature Phase Shift Keying), 16QAM (Quadrature Amplitude Modulation à 16 états) et 64QAM (Quadrature Amplitude Modulation à 64 états).

La Figure 1.6 représente l'adaptation du signal, en termes de modulation, aux conditions de l'environnement. Les cercles représentent les couvertures de la station de base BS relatives aux différentes modulations.

Une modulation adéquate est assignée dynamiquement par la station de base. Ceci dépend de la distance séparant les deux stations émettrice et réceptrice, le climat, les interférences du

signal et d'autres facteurs. Dans les hautes fréquences, les modulations 16 et 64 QAM sont automatiquement utilisées par le protocole 802.16.

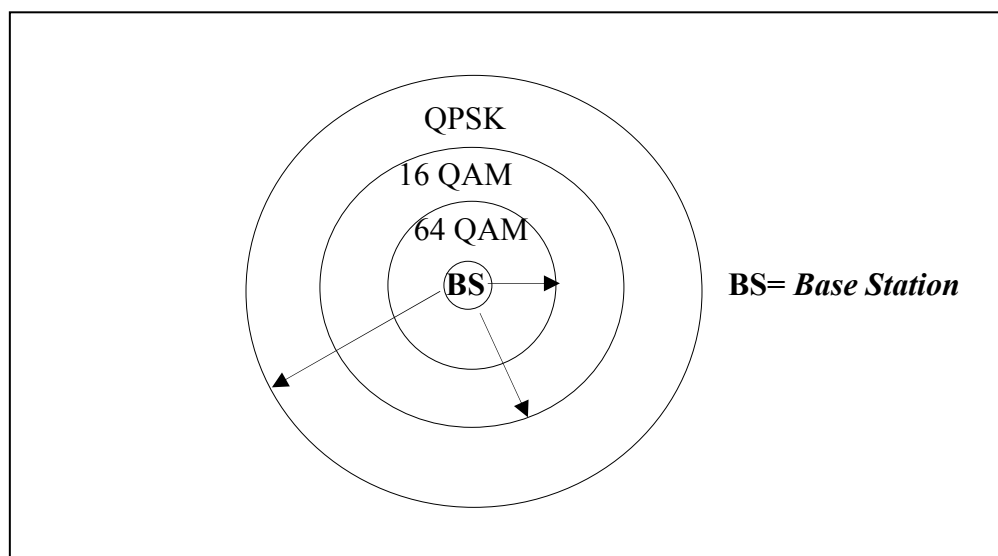


Figure 1.6 Couverture et modulation dans le WIMAX.

Le duplexage

Les deux types de duplexage, TDD (Time Division Duplex) et FDD (Frequency Division Duplex), sont supportés par le WIMAX.

Le FDD requière deux canaux séparés d'une bande de fréquence pour minimiser les interférences, l'un des deux canaux est utilisé pour la transmission et l'autre pour la réception. La majorité des bandes FDD sont allouées pour la transmission de la voix, qui permet de minimiser les délais, par contre le TDD est plus efficace pour la transmission de données et les systèmes IP (Pareek 2006).

Le TDD est utilisé dans les fréquences exemptées de licence. La bande de fréquence est utilisée donc pour les deux liens montant et descendant.

Dans le cas des bandes de fréquence exigeant une licence, l'opérateur aura le choix entre l'utilisation du TDD ou FDD.

La Figure 1.7 montre l'utilisation de la bande de fréquence dans les deux types de duplexage cités ci-dessus. Dans le mode TDD, le lien montant et le lien descendant utilisent la même bande de fréquence et émettent les informations pendant un temps bien défini.

Dns le cas de FDD, le lien montant et le lien descendant exploitent deux bandes de fréquence distinctes, et peuvent émettre en même temps.

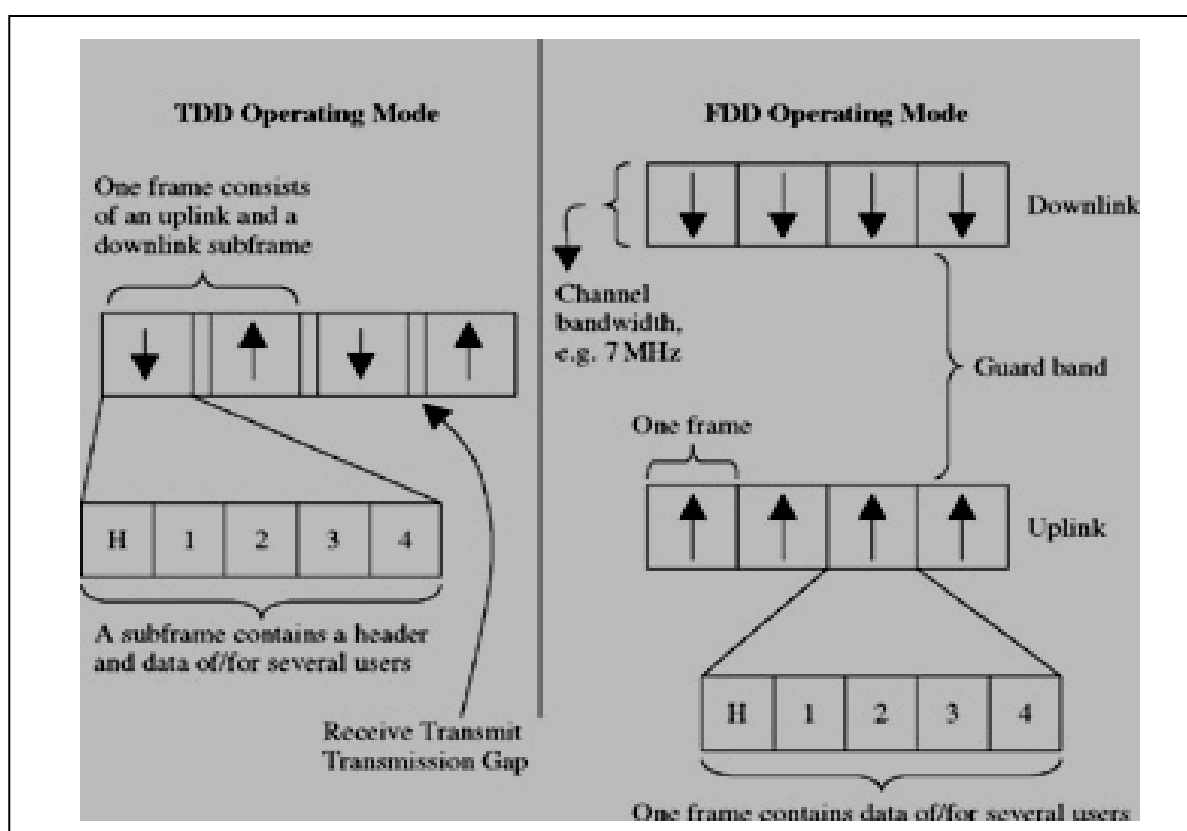


Figure 1.7 Multiplexage TDD et FDD dans les réseaux 802.16.

Tirée de Sauter (2006)

Le code correcteur d'erreur

Le code RS (Reed Solomon) est utilisé comme mécanisme de correction d'erreur FEC dont le rôle est d'augmenter le débit et la résistance aux interférences et au bruit. Le principe de fonctionnement du RS est d'ajouter $2t$ symboles à chaque k symbole de données de telle sorte à former un mot de $n = k + 2t$ symboles. $2t$ correspond au nombre de symboles de parité.

Les trames de la couche physique

Les trames durent 0.5, 1 ou 2 ms, leurs structures diffèrent selon le mode de duplexage utilisé. Dans le cas du FDD, les données transmises sur le uplink vont utiliser une bande de fréquence distincte de celle utilisée par les données transmises sur le downlink, d'où l'utilisation de deux trames. Par contre, dans le cas du TDD, les données transmises sur les deux liens vont utiliser une sous trame de la même trame, pour chaque lien.

La trame de la couche PHY contient un paquet de données (MAP), appelé UL-MAP dans le cas du uplink et DL-MAP dans le cas du downlink. L'UL-MAP contient des informations sur la trame en cours ou bien la suivante. Par contre, le DL-MAP contient des informations sur la trame en cours.

1.2.3 La couche MAC

La couche MAC est la deuxième couche du WiMAX (Figure 1.5). Elle détermine la façon avec laquelle les abonnés accèdent au réseau et comment les ressources leur sont allouées. Trois sous-couches constituent la couche MAC.

Les sous-couches de la couche MAC

Les trois sous-couches MAC sont les suivantes :

SSCS (Service Specific Convergence Sublayer)

Elle inclut un service spécifique de convergence des couches supérieures du modèle OSI avec les sous-couches MAC. Deux sous-couches sont définies : la sous-couche de convergence ATM, relative aux services ATM, et la sous couche de convergence de paquets, qui sont définis pour faire la correspondance des services par paquet, tels qu'IPv4, IPv6 et Ethernet. C'est cette sous-couche qui s'occupe de la transmission des SDU (Service Data Unit) à la connexion MAC désirée et de la préservation ou l'activation de la QoS, ainsi que l'allocation de la bande passante. Un autre rôle assez important assuré par cette sous-couche est la suppression et la génération des entêtes des paquets qui circulent pour améliorer le Payload.

La sous-couche commune (Common Part Sublayer)

Dans la sous-couche commune chaque service est associé à une connexion, du fait que la couche MAC soit orientée connexion. Ceci permet la demande de la bande passante, de la QoS ...etc. Dans le but de faciliter la gestion du trafic et la QoS, le transport de connexions est unidirectionnel.

La couche MAC réserve certaines connexions pour les broadcast et le multicast. Dans le cas du multicast, les stations d'abonnés doivent rejoindre un groupe pour pouvoir bénéficier des informations transmises.

La Security Sublayer

La sous-couche de sécurité contient les fonctions de sécurité relatives aux trames de la couche MAC. Elle comprend deux protocoles, selon (Thomas Hardjono 2005), le protocole d'encapsulation qui définit l'ensemble des algorithmes utilisés pour le cryptage des paquets de données échangés entre la BS et la SS, et le « Key Management Protocol » utilisé pour la gestion du matériel de chiffrement.

Les trames de la couche MAC

Le format général du MAC PDU (MPDU) est représenté sur la Figure 1.8. Chaque trame MAC commence par un en-tête MAC de longueur de 6 octets. Cet en-tête peut être suivi de la charge utile (Payload). Un MPDU peut contenir un CRC (contrôle par redondance cyclique) pour le contrôle d'erreur.

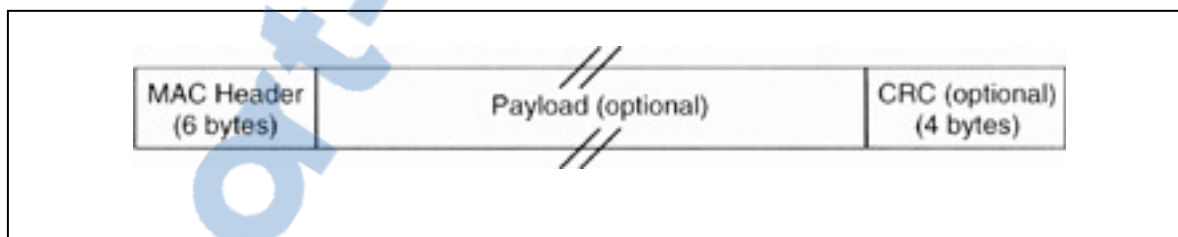


Figure 1.8 Format du MAC PDU.

Tirée de http://www-igm.univ-mlv.fr/~dr/XPOSE2006/aurelie_schell/caracteristiques_tech.php#mac

Il existe deux types d'en-têtes MAC : les en-têtes génériques (Figure 1.9) et les en-têtes de demande de bande passante (Figure 1.10). Le champ commun TYPE permet d'identifier le type de l'entête.

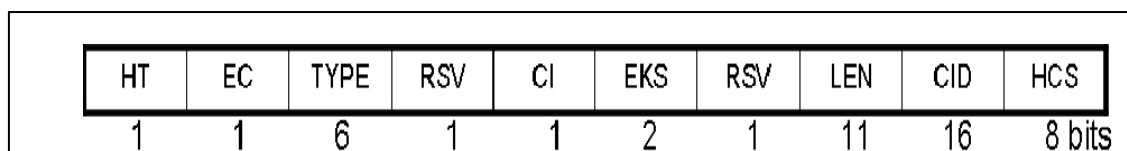


Figure 1.9 En-tête MAC générique.

Tirée de http://www-igm.univ-mlv.fr/~dr/XPOSE2006/aurelie_schell/caracteristiques_tech.php#mac

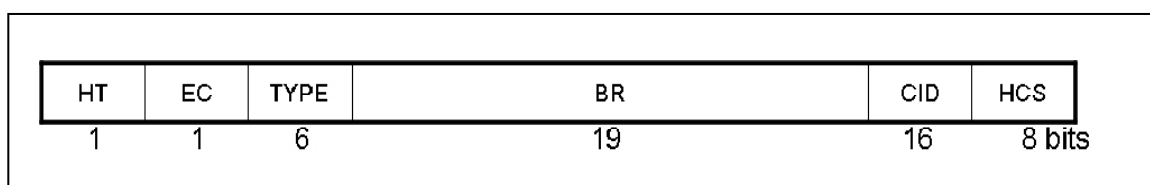


Figure 1.10 En-tête MAC de demande de bande passante.

Tirée de http://www-igm.univ-mlv.fr/~dr/XPOSE2006/aurelie_schell/caracteristiques_tech.php#mac

Les paramètres HT, EC, TYPE, CID et HCS sont communs aux deux types d'entête. Le HT spécifie le type d'entête. Quand il s'agit d'un entête générique, le HT prendra la valeur zéro, par contre, s'il s'agit d'un entête de demande de bande passante, le HT prendra la valeur un. L'EC quant à lui, sert au contrôle du cryptage. Il prendra la valeur zéro quand le Payload est non crypté et la valeur un si celui-ci est encrypté. Le CID est l'identificateur de la connexion de la station réceptrice, il sera utilisé à la place de son adresse MAC. Concernant le HCS, il représente l'entête de contrôle de séquence, pour éviter l'interprétation erronée de l'entête en raison d'erreurs de transmission. Les autres champs non communs sont les suivants : Le paramètre CI est un indicateur de CRC, il indique si le PDU contient un CRC. L'EKS, indique quel trafic, quelle clé de cryptage et quel vecteur d'initialisation vont être utilisés. Le paramètre LEN identifie la longueur totale du MAC PDU incluant l'entête CRC.

La charge utile est la partie de la trame qui contient les données à transmettre. Comme les entêtes MAC, celle-ci est décomposée en différentes parties contenant chacune un sous entête et un message de gestion.

Le CRC (Cyclic Redundancy Check) sert au contrôle d'erreur permettant de vérifier l'intégrité de la trame en utilisant un polynôme basé sur les données de cette trame. Le calcul du polynôme se fait au niveau de la station de base et de l'abonné, ensuite les résultats seront confrontés.

1.2.4 La Qualité de service

La QoS est un élément important pour la mesure de la performance du système. Elle constitue une priorité et un souci pour les constructeurs des réseaux de nouvelle génération.

Le réseau WiMAX est conçu de telle sorte à assurer la QoS pour le trafic montant, ainsi que pour le trafic descendant.

Les éléments de la QoS

Selon (Kalai Kalaichelan 2007) le WiMAX permet de mesurer les quatre éléments de QoS suivants :

- **La disponibilité du service :** Qui représente la proportion du temps durant laquelle un client autorisé est apte à exploiter les ressources du réseau, par rapport au temps total qui lui est alloué. Ce service peut être affecté par le processus de contrôle d'admission. Parfois, un même client peut se voir refuser la transmission d'un fichier FTP, alors qu'il peut utiliser la VoIP en un temps donné. Ceci est dû à l'ajustement de l'utilisation des ressources de la bande passante par le mécanisme de contrôle d'admission afin d'assurer un certain niveau de QoS.

- **Le *Data Throughput* :** C'est la quantité de l'information qui peut être transférée pendant la communication. Le WiMAX est en mesure de fournir différents niveaux de *Data Throughput*, selon les exigences du trafic, tels que le Constant Bit Rates (CBR) et le Best Effort (BE).
- **Le Délai :** Comprend le temps requis pour l'établissement d'une connexion et celui pris pour la transmission de l'information. Le IEEE 802.16 garantit une gestion de trafics par priorité, afin de respecter les exigences de certaines applications en termes de délais maximaux.
- **Le taux de perte de paquets :** C'est le rapport entre le nombre de paquets perdu et le nombre de paquets total transmi, pendant la communication.
- **La gigue :** C'est la variation du délai de transfert de l'information. Pour les applications multimédias, elle doit être la plus faible possible pour pouvoir parler d'un réseau fiable et d'une bonne QoS.

Les classes de la QoS

Afin de garantir la QoS dans les réseaux WiMAX, il faut satisfaire les besoins des différents types d'applications en terme de bande passante. Ces applications ont des exigences et des tolérances distinctes. Un point fort pour les réseaux WiMAX est leur habilité d'appliquer une différenciation de services.

Quatre catégories de qualité de service ont été définies dans la norme 802.16-2004, à savoir UGS (Unsolicited Grant Service), rtPS (real-time Polling Service), nrtPS (non-real-time Polling Service) et BE (Best Effort). Une cinquième a été ajoutée dans la norme 802.16e: la classe ertPS (extended real-time Polling Service).

Cette classification facilite le partage de bande passante entre les différents utilisateurs qui s'effectue selon la classe de QoS à laquelle ils appartiennent, permettant ainsi une

répartition efficace et adaptable des ressources existantes. Par conséquent, une demande de bande passante en temps réel, comme celle effectuée par un utilisateur de VoIP, aura la priorité dans l'allocation de bande passante en comparaison avec celle relative à FTP (File Transfer Protocol) ou à des applications de courrier électronique.(Nuaymi 2007).

Les différentes classes de QoS correspondantes au WiMAX fixe seront présentées dans ce qui suit :

UGS (Unsolicited Grant Service), Service à concession non sollicitée

La classe UGS est désignée pour supporter les flux de données à temps réel, ayant des paquets de données de taille fixe émis à des intervalles de temps réguliers. Ce service garantit une bande passante fixe et un temps de latence constant pour la connexion. Le WiMAX alloue la bande passante selon le besoin maximal des applications de cette classe.

Real-Time Polling Service (rtPS)

Cette classe de QoS est conçue pour transmettre des flux temps réels de taille variable à intervalle régulier comme une vidéo MPEG. La bande passante est allouée aux applications de cette classe selon leur besoin minimal.

Non Real-Time Polling Service (nrtPS):

Elle est conçue pour transmettre des flux tolérants aux délais avec une taille variable, tel que le FTP. La bande passante allouée aux flux de cette classe s'effectue selon leur besoin minimal.

Service Best Effort

Cette classe n'exige aucune QoS. Par conséquent, aucune portion de bande passante ne lui est allouée.

Notons qu'une cinquième classe a été définie pour la version mobile du WiMAX. Il s'agit de la classe extended real-time Polling Service (ertPS). Cette classe a les mêmes exigences en termes de QoS que la classe ertPS. Dans ce qui suit nous n'allons pas la prendre en considération car elle ne concerne pas la version fixe du réseau WiMAX.

1.2.5 Conclusion

Dans ce chapitre, nous avons présenté la norme IEEE802.16 correspondante aux réseaux WiMAX. Nous avons aussi résumé les caractéristiques techniques de ces réseaux. Cette technologie fournit des solutions pour les Backhaul, les réseaux VPN, les réseaux d'entreprise, l'accès Internet pour des clients résidentiels, relier des réseaux WiFi, ainsi que pour l'accès haut débit à Internet avec mobilité. Néanmoins, il existe certains points ouverts dans ce standard qui sont devenus des défis pour les chercheurs. Particulièrement, le support de la QoS en termes de débit, de délai et de pertes de paquet, ce qui nécessite la conception des systèmes de contrôle d'admission d'appels, ainsi que des algorithmes d'ordonnancement pour la gestion des appels admis. Ajouter à cela, la coexistence avec d'autres types de réseaux tels que le Wifi et le réseau filaire et la gestion de mobilité qui forment deux autres points ouverts dans ce standard exigeant une optimisation, tout comme la mise à l'échelle de ce type de réseau.

Bien que le standard donne tous les détails correspondants aux différentes classes de QoS et les méthodes d'allocation de la bande passante, mais l'ordonnancement de la gestion des ressources de la bande passante est laissé non normalisé.

Dans le chapitre 2 nous allons focaliser sur la problématique d'ordonnancement dans les réseaux WiMAX. Dans la présente étude, nous allons travailler sur l'amélioration de la QoS pour les différentes classes de service.

CHAPITRE 2

ÉTAT DE L'ART

2.1 Introduction

Dans le but de garantir la QoS, les réseaux WiMAX ont été dotés d'un ensemble de classes de services, à savoir UGS, rtPS, ertPS, nrtPs et BE. Cette classification permet d'appliquer une différenciation de services et de favoriser certaines applications par rapport à d'autres selon leurs exigences et tolérances. Toutefois, l'ordonnancement au sein de ces réseaux n'a pas été défini. Une multitude d'approches de gestion de trafic existe dans la littérature relativement aux réseaux sans fil, notamment, aux réseaux WiMAX.

Dans ce chapitre, nous allons présenter un survol des travaux de recherches ayant traité la gestion des ressources de la bande passante en présence de différents types d'applications dans les réseaux WiMAX.

2.2 Survol des algorithmes d'ordonnancement existants

Cette section présente les principaux algorithmes d'ordonnancement proposés dans la littérature et appliqués aux réseaux WiMAX. L'ordre de présentation n'interprète pas le degré de robustesse des solutions.

Dans (Xie, Chen et al. 2008) les auteurs présentent un système de gestion de file d'attente pour les différents types de trafic du réseau WiMAX nommé FEQ (Figure 2.1). Les arrivées vers les différentes files d'attente suivent la loi de Poisson de taux d'arrivées moyen λ_i , tel que $i \in [1, k]$, k étant le nombre de classes de trafics.

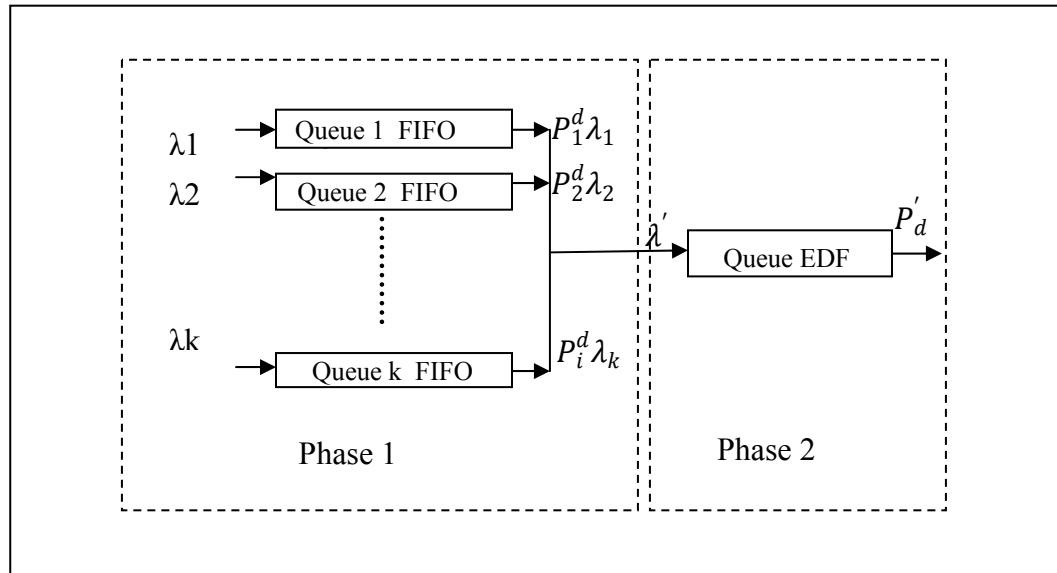


Figure 2.1 Modèle analytique relatif à l'algorithme FEQ.

Tirée de Xie, Chen et al. (2008)

L'allocation de la bande passante se fait en deux parties. Durant la première phase, les files d'attente seront servies selon l'algorithme WRR. En effet, le système allouera pour chaque type de trafic une bande passante égale au MRR (Minimum Reserved Rate). Le MRR relatif à chaque type de trafic représente le poids de la file correspondante durant la phase 1 de l'algorithme FEQ. Plus le trafic est exigeant, plus la valeur MRR qui lui correspond est grande, et plus le poids de la file d'attente est élevé. Cette politique favorisera les trafics ayant moins de tolérances par rapport à ceux ayant moins d'exigences.

Les paquets non servis durant la phase 1, vont être placés dans la file d'attente EDF pour être traités durant la phase 2. Le taux moyen d'arrivée de ces paquets vers la file EDF se donne par la formule suivante :

$$\lambda' = \sum_{i=1}^n \lambda_i P_i^d$$

Avec P_i^d est le taux de perte de paquet dans une file i .

Durant la deuxième phase, le système EDF (Earliest Deadline First) est utilisé. EDF permet de réduire le taux de perte de paquets en servant en premier les paquets dont le temps d'attente est le plus proche de temps maximal toléré.

Les résultats de la simulation montrent une performance de l'équité du système par rapport aux algorithmes Static Priority (SP) et Deficit Fair Priority Queue (DFPQ). Or, les résultats relatifs aux pertes de paquets restent semblables pour les trois solutions.

Les auteurs de (Thaliath, Joy et al. 2008) ont proposé un algorithme d'ordonnancement Service Classe Downlink Scheduling (SCDC), au niveau du lien descendant de la station de base. Leur objectif est de réduire le délai et l'amélioration du throughput pour le trafic à temps réel. L'idée consiste à classer les paquets de toutes les stations d'abonnés selon les classes correspondantes au niveau de la BS (Figure 2.2)

Ce modèle est basé sur le calcul des times slots relatifs à une trame donnée, réservés pour chaque type de trafic. Le time slot relatif au trafic de classe i de la station k se donne par la formule :

$$T_{ik} = \text{floor} (R_{ik} * T_k)$$

Tels que R_{ik} est la priorité d'attente pour le service i dans la station k , et T_k est le nombre de slots correspondant à la station k dans la trame, qui dépend du poids de la SS. Ce poids est proportionnel à la somme des priorités d'attente de tous les services présents dans la station en question. La fonction *floor* est utilisée afin d'obtenir un nombre entier. Les restes des calculs des times slots en utilisant cette fonction seront cumulés et attribués pour la classe BE.

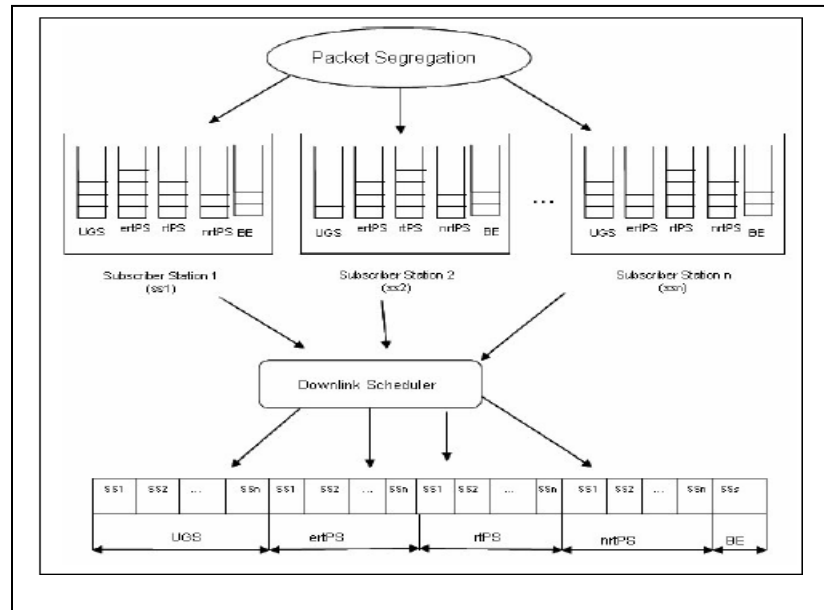


Figure 2.2 Discrimination des paquets au niveau de la BS pour le lien descendant.
Tirée de Thaliath, Joy et al.(2008)

Le modèle de gestion de file d'attente considérée en (Wang, Chan et al. 2008) est composé d'un ensemble de files d'attente regroupées en catégorie où chaque catégorie p contient M_p files d'attente et détient une priorité H . Plus l'ordre de h est élevé plus la priorité est considérée comme haute. Un paquet attend d'abord dans sa file d'attente de groupe p avant d'être transmis à la file d'attente PS qui regroupe tous les paquets en provenance de toutes les files. Le système ne tolère pas la présence de plus d'un message à la fois correspondant à une même file d'attente. Un message se compose de n paquets. Le temps moyen d'attente d'un paquet dans son groupe de file d'attente p est donné par L_p , par contre le temps d'attente d'un paquet dans la file PS est donné par $Sp(n)$. Le temps total d'attente dans le système relatif à un paquet donné est donné comme suit :

$$Dp(n) = Lp + Sp(n)$$

Le serveur traite les paquets présents dans PS en appliquant l'algorithme Round Robin avec considération de la priorité du trafic. L'application de ce modèle s'est effectuée sur les trafics de classe rtPS, nrtPS et BE. Les résultats obtenus par l'analyse numérique coïncident avec ceux fournis par la simulation.

(Prasath, Fu et al. 2008) propose un système d'ordonnancement pour le WiMAX mobile et il considère les deux types de trafic rtPS et nrtPS. Le but de ce mécanisme est d'optimiser le délai de la classe rtPS en permettant à ses applications d'être servies avant que leur temps d'attente n'arrive à échéance, mais sans trop tarder les paquets nrtPS.

La quantité de la bande passante totale requise pour ces deux services est donnée par :

$$B_{req}^{tot} = \sum_{j=1}^l AR_j^{rt} + \sum_{j=1}^m AR_j^{nr}$$

Avec l , m sont les nombres de flux rtPS et nrtPS, et AR_j^{rt} , AR_j^{nr} sont les débits moyens de chaque flux, respectivement.

La quantité de bande passante allouée pour rtPS et nrtPS sont présentées, respectivement, par les formules suivantes :

$$B_{tot}^{rt} = B_{poll} * \frac{N^{rt} * \alpha_i}{(N^{rt} * \alpha_i) + N^{nr}}$$

$$B_{tot}^{nr} = B_{poll} - B_{tot}^{rt}$$

Avec

- B_{poll} est la bande passante allouée pour les services de type polling (real-time Polling Service, non real-time Polling Service);
- N^{rt} est la taille de la file d'attente rtPS;

- N^{nr} est la taille de la file d'attente nrtPS;
- α_i est un facteur d'ajustement permettant de varier la quantité de bande passante allouée pour les deux types de trafic rtPS et nrtPS.

La portion de bande passante allouée pour la classe rtPS est proportionnelle à α_i . De plus, la bande passante additionnelle attribuée à cette classe est empruntée de la quantité de la bande passante réservée pour la classe nrtPS. Cette politique réduit le délai correspondant au trafic à temps réel. Cependant, α_i est choisi de telle sorte à ne pas causer un overflow pour le trafic nrtPS ce qui permet de préserver un bon niveau de QoS pour cette classe de trafic.

La simulation donne des résultats pour trois cas de figure, à savoir, le cas de stationnarité de la station réceptrice, la cas de sa mobilité sans le changement de la valeur de α , et un dernier cas relatif à la mobilité de la station avec l'application du système d'adaptation, donc en changeant la valeur de α . La simulation démontre que pour le trafic nrtPS, bien que le délai moyen de bout en bout pour la solution proposée n'atteigne pas le niveau de performance du cas stationnaire, il arrive quand même à améliorer la valeur obtenue en cas de mobilité sans application de l'adaptation. Pour le trafic rtPS le modèle proposé donne des résultats qui convergent vers ceux obtenus par l'application du modèle stationnaire, qui sont tous les deux meilleurs par rapport au troisième modèle.

Un autre système de gestion de files d'attente a été proposé dans (Lera, Molinaro et al. 2007) (Figure 2.3). Cet algorithme est implémenté au niveau de la BS, il assure l'ordonnancement des trafics rtPS, nrtPS et BE. Une fois que la connexion est acceptée, les paquets seront classifiés selon leur type dans la file d'attente correspondante. La gestion des trois files d'attente concernées par le système d'ordonnancement se fait à travers l'algorithme WF2Q. La sélection d'un paquet pour le servir prend en considération la valeur de son temps virtuel de début de service (Virtual Start Service Time) S_i et son Virtual Finish Service Time F_i . Ces deux valeurs sont calculées comme suit :

$$\begin{aligned}
 F_i &= S_i + \frac{L_i^k}{r_i} \\
 S_i &= \begin{cases} \max(F_i, V_{WF^2Q+}(a_i^k)) & \text{Si la queue } Q_i \text{ est vide} \\ F_i & \text{Sinon} \end{cases}
 \end{aligned}$$

Avec r_i est le taux de bande passante garantie pour le trafic i , L_i^k est la taille d'un paquet de classe i et $V_{WF^2Q+}(a_i^k)$ est le system virtual time. Ce système choisit le paquet dont la valeur F_i est la plus petite. Cependant, avant de procéder à la transmission du paquet choisi, le système vérifie d'abord s'il est marqué par le Carrier to Interference and Noise Ratio(CINR) comme interdit d'être transmis. Si c'est le cas, alors le serveur va choisir un autre paquet non marqué, appartenant à la même classe pour le servir à la place. Notons qu'à chaque début de trame, un flux est marqué relativement à la puissance de réception, dans le cas où elle n'atteint pas une valeur donnée le canal sera marqué comme mauvais et le flux sera considéré comme interdit à transmettre. Si tous les paquets résidents dans une file Q_i sont interdits d'être transmis et ne peuvent pas substituer le paquet choisi pour service, mais qui n'est pas en mesure de l'être, alors le système conserve ses valeurs S_i et F_i . Par conséquent, cet algorithme granit à la file Q_i d'être servi dès qu'elle reçoit un paquet bon à transmettre, étant donné que les valeurs S_j et F_j des autres files vont augmenter avec le temps.

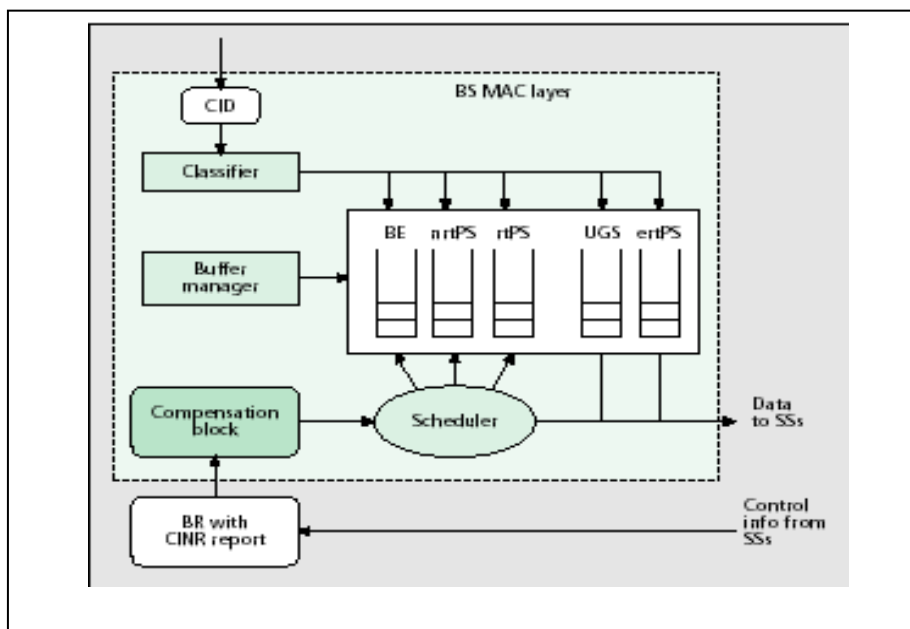


Figure 2.3 Algorithmes d’ordonnancement Channel-aware QoS.

Tirée de Lera, Molinaro et al. (2007)

Dans (Lin, Chou et al. 2008), les auteurs ont comparé quelques algorithmes d’ordonnancement en les déployant au sein d’un réseau WiMAX. Pour ce faire, ils ont considéré deux types de trafics, à savoir CBR et nrtPS, ils ont examiné le throughput, le délai et la perte de paquet au niveau de liens communs par lequel transitent tous les paquets. Deux scénarios ont été implémentés, le premier considère une capacité de lien de 100 Mbps et le second fixe cette valeur à 1 Mbps. Les algorithmes déployés sont : Weighted Fair Queuing (WFQ), Deficit Round Robin (DRR), Random Early Detection (RED), Fair Queuing (FQ), DropTail et RED with In/Out, dans lequel un paquet sera marqué avec « In » si le taux d’arrivée n’excède pas une valeur donnée, dans le cas contraire il sera marqué avec « Out ».

Les résultats de la simulation montrent que l’algorithme DRR est le meilleur à appliquer pour réduire le délai de transition des paquets dans le lien commun (Figure 2.4). Par contre, WFQ représente l’algorithme le plus défavorable pour le délai. Les autres algorithmes donnent des résultats semblables.

Concernant la perte de paquets, elle est considérablement supérieure lors de l'application de WFQ et de DRR. Elle atteint les pourcentages les plus bas lors de l'application de FQ et DropTail (Figure 2.5).

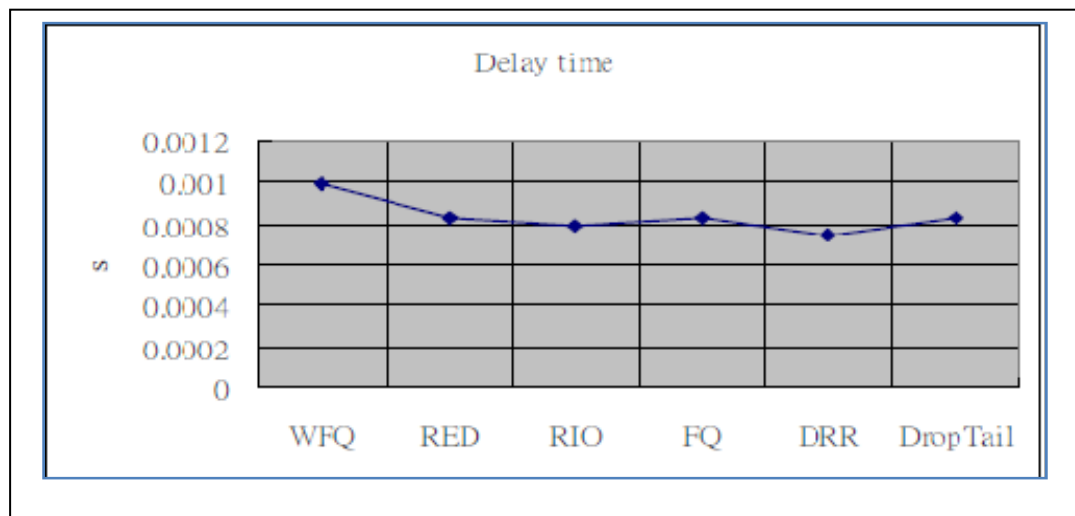


Figure 2.4 Délai d'attente dans le lien commun à UGS et nrtPS.

Tirée de Lin, Chou et al. (2008)

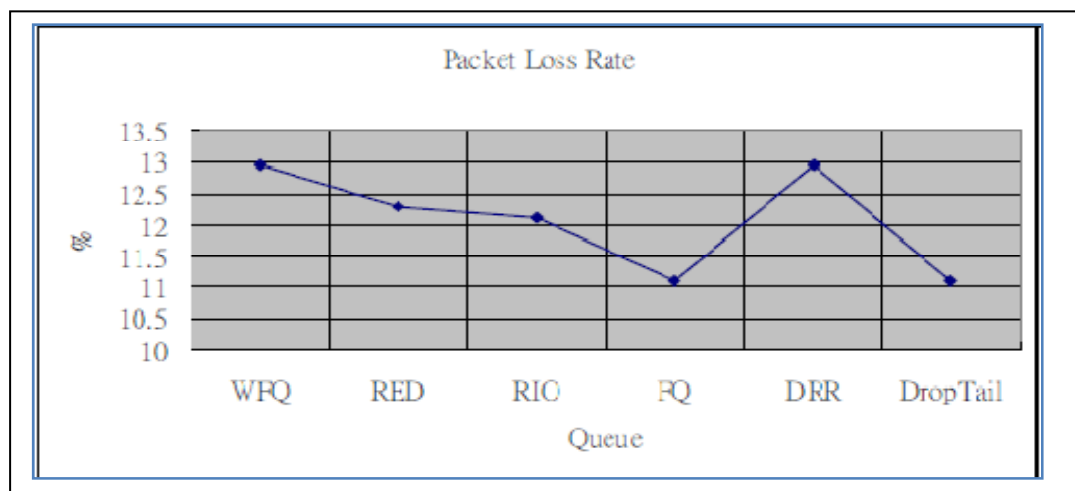


Figure 2.5 Taux de perte de paquets dans le lien commun à UGS et nrtPS.

Tirée de Lin, Chou et al. (2008)

La solution Opportunistic Stable Queue (OSQ) proposée en (Mehrjoo and Shen 2008) assure la gestion de files d'attente des différents trafics, dans les réseaux WiMAX au sein du lien descendant. Elle exploite les informations sur l'état du canal et de la file d'attente des utilisateurs afin d'appliquer l'ordonnancement des trafics avec différentes QoS. Le système calcule pour chaque file d'attente dans le système son degré de stabilité S_i qui représente le rapport entre le nombre moyen des départs estimé et le nombre moyen d'arrivées estimées relatives à la queue i , cette solution consiste à garder la valeur de S la plus proche de 1, autrement dit le système d'ordonnancement tente de garder le taux de départs et le taux des arrivées le plus proche possible pour chaque type de trafic.

L'algorithme d'ordonnancement présenté dans (Lo and Hong 2008) est une combinaison des deux algorithmes d'ordonnancement inter-classe et intra-classe. Les auteurs tiennent à améliorer la QoS en termes de maximisation du throughput, réduction du taux de perte de paquets ainsi que d'équité de service des différents types de trafics. Dans la phase d'ordonnancement inter-classe, chaque classe de trafic va se voir allouée une bande passante égale à la bande passante minimale requise pour ce type de trafic plus une bande passante additionnelle équivalente au minimum entre la valeur de la bande passante maximale relative au type de trafic en question moins la bande passante minimale requise pour cette classe, et la quantité de la bande passante restante divisée par le nombre total de connexions dans le réseau. Finalement, chaque connexion aura la quantité de bande passante $N_i^{min} + N_i^{add}$.

Pour la phase de trafic intra-classe, la valeur de $\sum_i N_i^{min} + N_i^{add}$ de toutes les connexions rtPS est prise en considération. Cette bande passante sera assignée aux différentes connexions rtPS selon la valeur de leur délai dans la file. Les connexions dont le délai est plus important seront servies en premier, pour cela l'algorithme EDF est appliqué. Par la suite, les connexions ne souffrant pas d'excès de délais seront servies en appliquant le WFQ.

Dans (Qu, Fang et al. 2008), Une approche de gestion de files d'attente à multi niveaux a été mise en place dans le réseau WiMAX pour les cinq classes de services, avec une politique de réservation de bande passante aux trafics rtPS et nrtPS selon le minimum qu'ils exigent, ce qui va permettre de libérer une portion de bande passante qui sera utilisée pour compenser

d'autres services qui ont besoin de la bande passante. Le premier niveau concerne l'ordonnancement au sein des files d'attente rtPS, nrtPS ou BE. Chaque flux rtPS est assigné à une file d'attente rtPS qui seront gérées selon l'algorithme Ideal Weighted Fair Queuing (IWFQ), le paquet qui exige moins de temps de traitement sera donc sélectionné. Cette classe de trafic ne sera que peu compensée. Le trafic nrtPS suit le même comportement que rtPS, sauf que cette classe sera plus compensée par le système qui va procéder aussi à garantir la bande passante nécessaire pour la classe BE en exploitant une quantité de ressources destinées à la compensation des services en besoins. IWFQ sera encore une fois adopté pour la gestion des flux de type BE. Le deuxième niveau consiste à la gestion des différentes files d'attente UGS, ertPS, rtPS, nrtPS et BE. Deux mécanismes sont déployés, à savoir une allocation périodique de bande passante pour UGS et ertPS et une gestion des autres queues selon l'algorithme Weighted Round Robin (WRR) en assignant des poids aux différents trafics de telle sorte à privilégier la classe rtPS, ensuite nrtPS et finalement BE. Les résultats de la simulation montrent que cette approche garantit un délai moyen acceptable pour les cinq classes de trafic du réseau WiMAX.

L'étude effectuée en (Belghith and Nuaymi 2008) consiste à comparer les algorithmes Round Robin (RR), maximum Signal-to-Interference Ratio (mSIR), WRR, Temporary Removal Scheduler conjointement avec Round Robin (TRS)+RR et TRS+mSIR. Le principe de TRS est d'exclure les paquets dont la puissance de transmission est jugée comme faible de la liste d'information à envoyer pour un certain temps T_r . La puissance de ce paquet sera analysée de nouveau après l'écoulement de T_r . Il ne va réintégrer la liste d'envoi si et seulement s'il répond à l'exigence du système en terme de puissance de transmission ou s'il reste hors de cette liste pendant un temps T multiple de L fois du temps T_r . Pour mSIR, les ressources de la bande passantes seront allouées à l'utilisateur ayant le SIR le plus élevé. Les résultats de la simulation montrent que RR est le moins performant des algorithmes testés car le taux de paquets servis en appliquant cette approche est le plus faible. Or, TRS-mSIR et mSIR donnent les meilleurs résultats pour le nombre de paquets servis.

Dans (Vinay, Sreenivasulu et al. 2006), les auteurs proposent un système de gestion de file d'attente M/D/1 en appliquant conjointement les deux algorithmes d'ordonnancement Earliest Due Date (EDD) et WFQ. EDD marque chaque paquet arrivé par une étiquette appelée Expected Dead Line Value (ExD^{Qn}), En supposant que :

- AT^{Qn} : Temps d'arrivé du $n^{ième}$ paquet
- ST^Q : Temps maximal du service du $n^{ième}$ paquet
- $Xmin$: Temps d'inter arrivée minimal

ExD^{Qn} sera donné par le système d'équation suivant :

$$ExD^{Q1} = AT^{Q1} + ST^Q \text{ first paquet}$$

$$ExD^{Qn} = \max\{ExD^{Qn-1} + Xmin, \quad AT^{Qn-1} + ST^Q \}$$

Pour cet algorithme, plus la valeur de ExD^{Qn} sera élevée plus le paquet sera moins prioritaire. Les auteurs ont comparé les délais de bout en bout relatifs aux différentes classes de service en appliquant le modèle EDD seul, ensuite en l'appliquant conjointement avec WFQ (modèle hybride). La simulation montre la performance du modèle EDD-WFQ pour les services à temps réel concernant la réduction du délai de bout en bout.

Les auteurs dans (Ruangchaijatupon, Wang et al. 2006) ont réalisé une simulation pour la comparaison de l'utilisation des deux algorithmes EDF et WFQ dans un réseau WiMAX en présence de plusieurs types de trafic. Les résultats montrent que EDF offre les meilleurs délais de bout en bout pour les applications real-time cependant il défavorise le trafic non-real-time. Or, WFQ donne un délai considérablement haut pour la vidéo par rapport à FTP et la voix, à cause du taux élevé de la perte de paquets. Les auteurs concluent que ni EDF ni WFQ ne puissent garantir à la fois un bonne QoS pour toutes les classes de service.

Contrairement aux approches précédentes, la solution proposée dans (Wenfeng Du, Zhen Ji et al. 2008) tient à favoriser les trafics de basses priorités. En effet, la BS va récupérer la bande passante allouée et non utilisée des SSs qui sont dans sa couverture et la redistribue selon les besoins des deux classes de trafic nrtPS et BE. Pour ce faire, toutes les requêtes de demande de bande passante relatives à ces deux dernières classes seront mises en attente dans un buffer et seront donc servi selon la disponibilité de bande passante restituée des SSs. Le modèle mathématique développé dans cette étude a permis de calculer L , la longueur de la file d'attente des trafics nrtPS et BE, ainsi que D , le délai moyen d'attente selon les deux formules suivantes :

$$L = \sum_{k=1}^{M_{temp}} k * p_k$$

$$D = \sum_{k=1}^{M_{temp}} p_k * d_k$$

Notons que :

- k : le nombre de requêtes mises en attente dans le buffer nrtPS_BE;
- p_k : la probabilité que k requêtes de type nrtPS ou BE soient mises en attente;
- d_k : le délai d'attente de la requête k dans la file d'attente;
- M_{temp} : La taille maximale de la file d'attente nrtPS_BE.

Les résultats de cette étude ont été validés par la simulation.

2.3 Synthèse et limites des solutions existantes

Les approches exposées dans la littérature permettent d'apporter des améliorations pour la gestion des différents types de trafics dans les réseaux WiMAX. L'objectif principal des différentes approches est de concevoir des modèles permettant de réduire les délais d'attente et de pertes de paquets, ainsi que de maximiser le débit. Plusieurs algorithmes d'ordonnancement ont été utilisés pour cette fin. Entre autres, l'application de l'algorithme WFQ seul ou conjointement avec un autre ordonnanceur, était très fréquente dans les solutions que nous avons examinées. La bande passante allouée aux différentes classes de trafic s'effectue de telle sorte à assurer l'équité de service pour tous les trafics. Chaque file d'attente correspondant à une classe donnée va se voir attribuée un poids. Ce poids va correspondre à la portion de bande passante qui lui est réservée. Cependant, il arrive parfois que les ressources allouées ne seront pas entièrement utilisées. Nous parlerons donc d'une sous utilisation de la bande passante. Parfois, d'autres formes de WFQ sont utilisées telles que IWFQ ou WRR. Les limites de ces algorithmes sont semblables à ceux de WFQ. En outre, EDF, à son tour est l'un des principaux algorithmes de gestion de files d'attente auquel on fait recours. Le principe de cet algorithme est de traiter les paquets ayant l'échéance la plus proche en priorité. Cet algorithme n'est pas optimal pour le cas non préemptif (Ekelin 2006). Sans oublier les deux systèmes RR et FQ qui malgré leur souci de garantir l'équité de service, ils négligent le fait que certains trafics doivent être traités en priorité car ils ont plus d'exigences en terme de besoin en bande passante. D'autres approches ont été proposées afin de contourner les faiblesses des solutions citées dessus. En revanche, la quasi-totalité favorise les trafics de haute priorité.

Il est évident que les trafics non temps réel ont moins d'exigences par rapport aux services real-time, mais cela n'empêche pas à songer à une performance de la QoS de ces services d'autant plus si cette amélioration n'affecte point la performance des trafics prioritaires. De ce fait, notre approche proposée, dite algorithme de courtoisie, permet de réduire les délais d'attente et le taux de pertes de trafic pour les classes de basse priorité en bénéficiant de la courtoisie des trafics des classes prioritaires, grâce au mécanisme de calcul de priorité de

transmission des paquets et du temps d'attente supplémentaire toléré pour les paquets de la classe courtoise durant lequel les paquets moins exigeants seront servis à la place. Notre approche sera détaillée dans le chapitre 3.

CHAPITRE 3

ALGORITHME DE COURTOISIE

3.1 Introduction

Le nouveau millénaire voit naître l'évolution fulgurante des réseaux sans fil large bande, dont notamment le WiMAX qui a été conçu principalement pour la desserte d'internet à haut débit dans les zones dépourvues d'accès filaire.

Ces réseaux peuvent être interconnectés selon deux topologies, à savoir la topologie Point-Multipoint (PMP) ou la topologie maillée. Contrairement à la première, le réseau maillé permet aux stations d'abonnés d'échanger l'information sans passer par une station de base. Dans le mode PMP, les clients WiMAX « SSs » se communiquent à travers la station de base BS, chaque SS doit donc établir une connexion avec la BS, chaque connexion est unique et caractérisée par un identificateur, appelé connexion identifier «CID ». Les paquets à transmettre à travers les uplink sont placés dans des flux de trafic, chaque flux regroupe les paquets de même classe de service. La connexion établie entre la SS et la BS va être exploitée pour la transmission des flux qui lui sont relatifs.

Bien que le standard 802.16d définisse les classes de services assurant la différenciation des types trafic selon la priorité et les exigences de chacun, à savoir UGS, rtPS, nrtPS et BE, il ne décrit pas un système d'ordonnancement pour les liens uplink et downlink. Comme nous l'avons vu dans le chapitre précédent, maintes solutions ont été développées pour la gestion des trafics avec priorités dans les réseaux WiMAX. Seulement, la quasi-totalité de ces solutions favorise les trafics de hautes priorités, telles que, la VoIP et la vidéo-conférence. Ce qui implique une optimisation qui n'atteint pas l'entière satisfaction des clients. La performance de l'allocation de la bande passante pour les classes de basses priorités se montre donc capitale.

Dans la suite de ce chapitre nous allons présenter notre approche de performance d'ordonnancement dans les réseaux WiMAX fixes, que nous proposons pour, non seulement diminuer le délai d'attente dans les queues de basses priorités, mais aussi pour réduire le taux de pertes de paquets au sein des classes correspondantes. Cela dit, la performance de notre solution n'affecte point la QoS des classes de hautes priorités.

La section suivante va servir d'une description de notre solution, à savoir l'algorithme de courtoisie. La section 3.2 illustre les conditions d'utilisation de cet algorithme. Celle qui succède présente le détail du développement du modèle mathématique correspondant au cas d'un système de deux files d'attente M/G/1 avec priorité non préemptive. Une extension de notre étude pour couvrir le cas de n files d'attente sera inclus dans cette section. Ultérieurement, nous allons mettre en évidence l'importance de l'analyse mathématique entamée dans ce travail. Par la suite, la structure de notre algorithme sera présentée. Finalement, nous allons achever ce chapitre par une conclusion

3.2 Description de l'algorithme de courtoisie

Quatre classes de QoS ont été définies dans les réseaux WiMAX fixe, connues sous les noms d'UGS, rtPS, nrtPS et BE. Cette classification a permis aux trafics exigeants en terme de bande passante d'être mieux servis, grâce aux différents niveaux de priorités assignées aux maintes classes de services. Cependant, le standard ne définit pas des solutions d'ordonnancement correspondant à ces trafics, par conséquent, plusieurs propositions ont été faites dans le but de compléter ce manque. Or, des solutions telles que le WFQ, PQ, WRR se montrent défavorables pour la classe nrtPS. À défaut ou par manque de ressources de bande passante allouées au trafic moins prioritaire, les paquets correspondants doivent attendre davantage dans leur file. En effet, une attente excessive d'un paquet peut entraîner son rejet. Parfois ce manque est insensé du moment qu'il existe des quantités de bande passante non utilisées, mais détenues dans le réseau, un cas de figure se présente quand une classe de haute priorité, comme la VoIP, se fait allouée des taux de ressources qui ne reflètent pas son

besoin réel, ce qui entraîne une sous utilisation de la bande passante, alors que les autres classes, moins prioritaires, telle que nrtPS, subissent des pertes de paquets considérables.

L'objectif de notre algorithme est d'optimiser l'utilisation de la bande passante, en termes de délai et de perte de paquets correspondants aux classes défavorisées. Pour ce faire, un système de gestion de files d'attente sera implémenté. Dans un premier moment, nous allons définir deux classes de QoS à savoir la classe rtPS, relative aux applications de la VoIP, et la classe nrtPS relative aux applications FTP. Notre solution tient à ne pas affecter la QoS des classes prioritaires.

Chaque paquet qui arrive à une file d'attente, va se doter d'une priorité de service. Nous allons expliquer, ultérieurement dans ce chapitre, comment sera calculée cette priorité et comment procéder pour inclure la différenciation de service dans notre modèle. Les paquets seront servis selon la valeur de leur priorité, commençant par celui qui en détient la valeur la plus grande. Il est important de mentionner que le paquet d'entête de chaque file d'attente est considéré comme le plus prioritaire de sa classe, étant donné qu'il s'agit de files d'attente de type FIFO.

Ce mécanisme va assurer une équité de service au sein de notre système de file d'attente. Cependant, pour améliorer notre solution, nous l'avons orné d'un mécanisme de courtoisie, permettant aux paquets de la classe nrtPS d'être servis, au lieu de ceux qui appartiennent à la classe rtPS. Cette solution n'est possible que si les paquets courtois ont suffisamment du temps pour être servis sans que leur QoS ne soit affectée.

Tel que nous l'avons décrit, notre solution porte sur la gestion des deux types de trafic rtPS et nrtPS, relatifs à la VoIP et FTP, respectivement. Toutefois, il est possible de l'étendre pour un nombre plus important de classes de trafics en explorant le cas de la présence de k appels dans le réseau et l'arrivée de m nouveaux appels de types multiples, UGS, rtPS, nrtPS et BE. Nous pouvons même l'appliquer pour un réseau WiMAX mobile en considérant les classes de services correspondantes aux différents types de handoffs.

Notons que nous implémentant cette solution au niveau des liens montants, au sein de la station de base, mais cela ne nous empêche pas de la déployer au sein des relais ou même au niveau d'une station d'abonné, qui jouera le rôle d'un point d'accès pour un réseau WLAN. Cette approche ne s'applique pas à toutes les transmissions de paquets. Elle est plutôt recommandée dans le cas de manque de bande passante pour le trafic FTP, alors que la VoIP détient une quantité en excès. Dans la section suivante, nous allons détailler les conditions d'utilisation de cette solution.

3.3 Conditions d'application

Les conditions d'application de l'algorithme de courtoisie vont garantir aux paquets courtois d'attendre davantage dans leur file d'attente, sans affecter leur QoS.

Afin de simplifier notre étude, nous allons considérer deux classes de QoS, à savoir C_1 et C_2 , relatives respectivement aux trafics rtPS et nrtPS. Les paquets de la classe C_1 auront la priorité $Pr1$, tandis que ceux de la classe C_2 auront la priorité $Pr2$.

Les quatre conditions suivantes doivent être satisfaites pour que les paquets de la classe C_1 cèdent leur tour de service à ceux de la classe C_2 :

3.3.1 Condition 1

Comme nous l'avons spécifié ci-dessus, l'objectif de notre solution est de permettre aux paquets de moindres priorités d'être servi à la place de ceux qui détiennent une haute priorité. En effet, dans notre cas, la première condition d'application de cet algorithme est la suivante :

$$Pr1 > Pr2 \quad (3.1)$$

3.3.2 Condition 2

Nous avons déjà mentionné que l'application de cette solution ne doit pas affecter la QoS de la classe courtoise. Autrement dit, le taux de perte de paquets relatif à celle-ci ne doit pas dépasser une certaine valeur ω_1 qui représente le seuil toléré de perte de paquets pour le trafic de classe C_1 . Ceci doit être vérifié même après la transmission des paquets bénéficiaires de ce mécanisme. Par conséquent, la probabilité de perte de paquets η_1 en temps t' , relative au trafic de la classe C_1 , ne doit pas atteindre la valeur ω_1 .

t' représente la valeur du temps lors de l'achèvement du service des paquets bénéficiaires de notre solution. Il représente la somme de t et τ_2 , tels que t est le temps initial de l'exécution de notre solution, et τ_2 est le temps de service des paquets bénéficiant de la courtoisie.

$$t' = t + \tau_2 \quad (3.2)$$

Soit μ le temps moyen de service d'un paquet. Soit γ_2 le nombre de paquets de classe C_2 traités en appliquant l'algorithme de courtoisie, alors τ_2 sera calculé comme suit :

$$\tau_2 = \gamma_2 * \mu \quad (3.3)$$

De (3.2) et (3.3), on peut calculer t' comme suit :

$$t' = t + \gamma_2 * \mu \quad (3.4)$$

La condition 2 se présente donc comme suit :

$$\eta_1(t') < \omega_1 \quad (3.5)$$

En prenant en considération (3.4), (3.5) va donc s'écrire de la sorte :

$$\eta_1 (t + \gamma_2 * \mu) < \omega_1 \quad (3.6)$$

3.3.3 Condition 3

Cette condition concerne la probabilité de perte de paquets de la classe C_2 , notée η_2 , juste avant l'application de l'algorithme de courtoisie, c'est-à-dire en temps t .

La valeur de η_2 permet de déterminer si la classe C_2 a besoin de bande passante supplémentaire ou non. En fait, si η_2 est supérieure au seuil de perte de paquets tolérés pour la classe C_2 , nommé ω_2 , on considérera les paquets appartenant à cette classe en besoin d'être servi. Par conséquent, l'application de l'algorithme de courtoisie sera avantageuse.

La troisième condition sera donc donnée comme suit :

$$\eta_2 (t) > \omega_2 \quad (3.7)$$

3.3.4 Condition 4

Cette condition permet de vérifier si τ_2 , le temps de service des paquets bénéficiaires de la classe C_2 , ne va pas excéder ξ_1 , le temps d'attente supplémentaire toléré pour les paquets de la classe courtoise C_1 . À compter du temps de début d'exécution de notre solution.

On peut donc formuler la quatrième condition comme suit :

$$\tau_2 < \xi_1 \quad (3.8)$$

La non-satisfaction de cette condition va augmenter le taux de perte de paquets de classe C_1 . En conséquence, la valeur de η_1 va excéder le seuil de perte de paquets tolérée ω_1 . Ce qui affecte la QoS de la classe C_1 . Cette dernière condition va permettre donc de maintenir la

valeur η_1 en dessous de ω_1 , avant et après l'application de la solution en question. Ce qui mène à reformuler la condition 2 comme suit :

$$\eta_1(t) < \omega_1 \quad (3.9)$$

Pour conclure, les quatre conditions d'application de l'algorithme de courtoisie sont les suivantes :

- $Pr1 > Pr2$
- $\eta_1(t) < \omega_1$
- $\eta_2(t) > \omega_2$
- $\tau_2 < \xi_1$

3.4 Analyse mathématique

Dans cette section nous allons développer le modèle mathématique relatif à notre solution. Nous allons commencer par l'étude de cas de deux files d'attente, pour finir par une généralisation sur n files d'attente.

3.4.1 Cas de deux files d'attente

Système de file d'attente

La Figure 3.1 représente le système de file d'attente considéré dans notre étude. Il s'agit d'un modèle de file d'attente M/G/1 avec priorité non préemptive. Il se compose de deux files d'attente rtPS_Queue et nrtPS_Queue, ainsi que d'un seul serveur.

La file d'attente $rtPS_Queue$ est de taille maximale $K1$. Elle contient les paquets de la classe prioritaire C_1 . La file d'attente $nrtPS_Queue$ est de taille maximale $K2$. Elle regroupe les paquets de la classe de basse priorité C_2 .

La classe C_1 correspond au trafic $rtPS$ et la classe C_2 est relative au trafic $nrtPS$.

Les arrivées au système de files d'attente suivent une loi de poisson de taux moyens d'arrivées λ_1 paquet/s pour les paquets de classe C_1 et λ_2 paquets/s pour ceux de la classe C_2 . Notons λ le taux moyen total d'arrivées par seconde dans le système. On obtient alors :

$$\lambda = \lambda_1 + \lambda_2$$

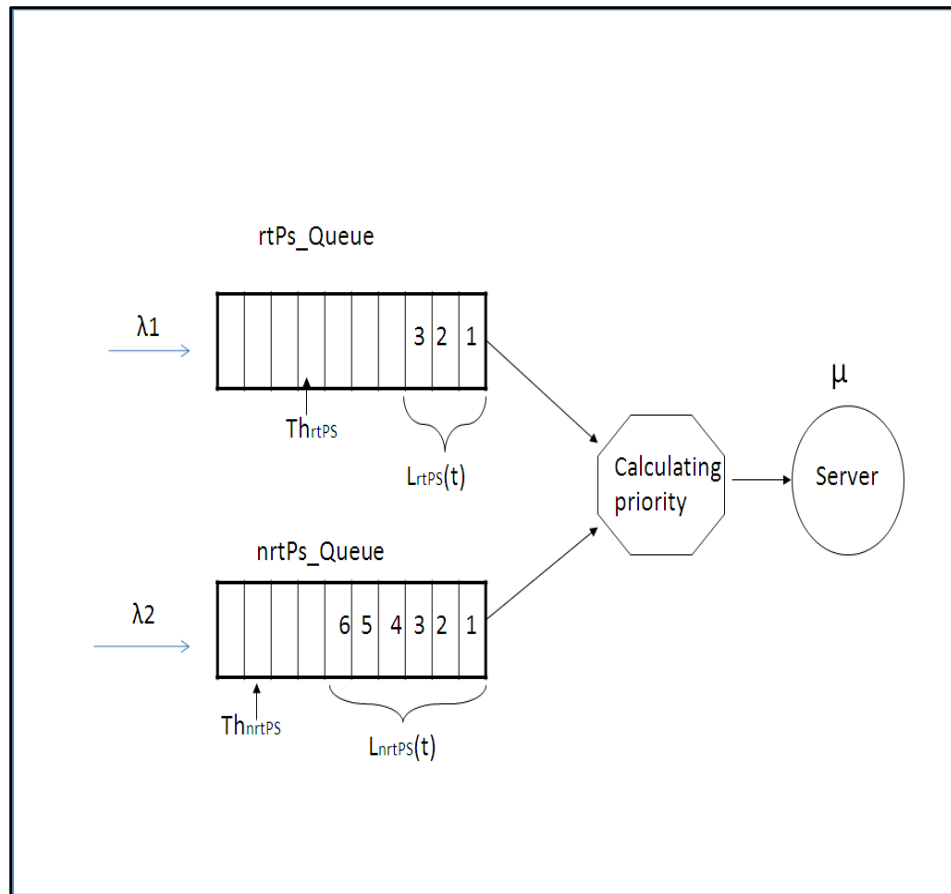


Figure 3.1 Système de file d'attente M/G/1 courtoisie : deux queues.

Le service correspondant aux paquets de la classe C_1 et ceux de la classe C_2 , suivent une loi exponentielle d'un temps de service moyen $1/\mu$ paquet/sec.

Les temps d'interarrivées des paquets de la classe C_1 et ceux de la classe C_2 suivent tous les deux une loi exponentielle. Les temps inter arrivés moyens relatifs aux paquets de la classe C_1 et ceux de la classe C_2 sont égaux à $1/\lambda_1$ et $1/\lambda_2$, respectivement.

Les longueurs moyennes des files d'attente $rtPS_Queue$ et $nrtPs_Queue$, nommées $Lq1$ et $Lq2$, respectivement, suivent une loi géométrique. Il est de même pour $L1$ et $L2$ les nombres de paquets moyens dans le système correspondant aux classes C_1 et C_2 , respectivement.

On note $Wq1$ et $Wq2$, les temps d'attente des paquets de la classe C_1 et ceux de la classe C_2 , respectivement, dans les files d'attente correspondantes.

Nous supposons que $P_{m,n,r}$ est la probabilité qu'il y ait m paquets de classe C_1 et n paquets de classes C_2 dans le système, et un paquet de classe C_1 ou de classe C_2 dans le service. Si $r=1$, alors on considérerait que le serveur traite un paquet de classe C_1 , sinon si $r=2$ alors on dirait que le paquet en service appartient à la classe C_2 .

La Figure 3.2 représente le diagramme d'état correspondant au modèle M/G/1 avec Priority Queuing. Il sera le modèle de base que nous considérons dans notre solution. Les arcs en bleu représentent des transitions suite à des arrivées λ_1 . Ceux en rouge représentent des arrivées λ_2 . Par contre ceux en vert représentent des services.

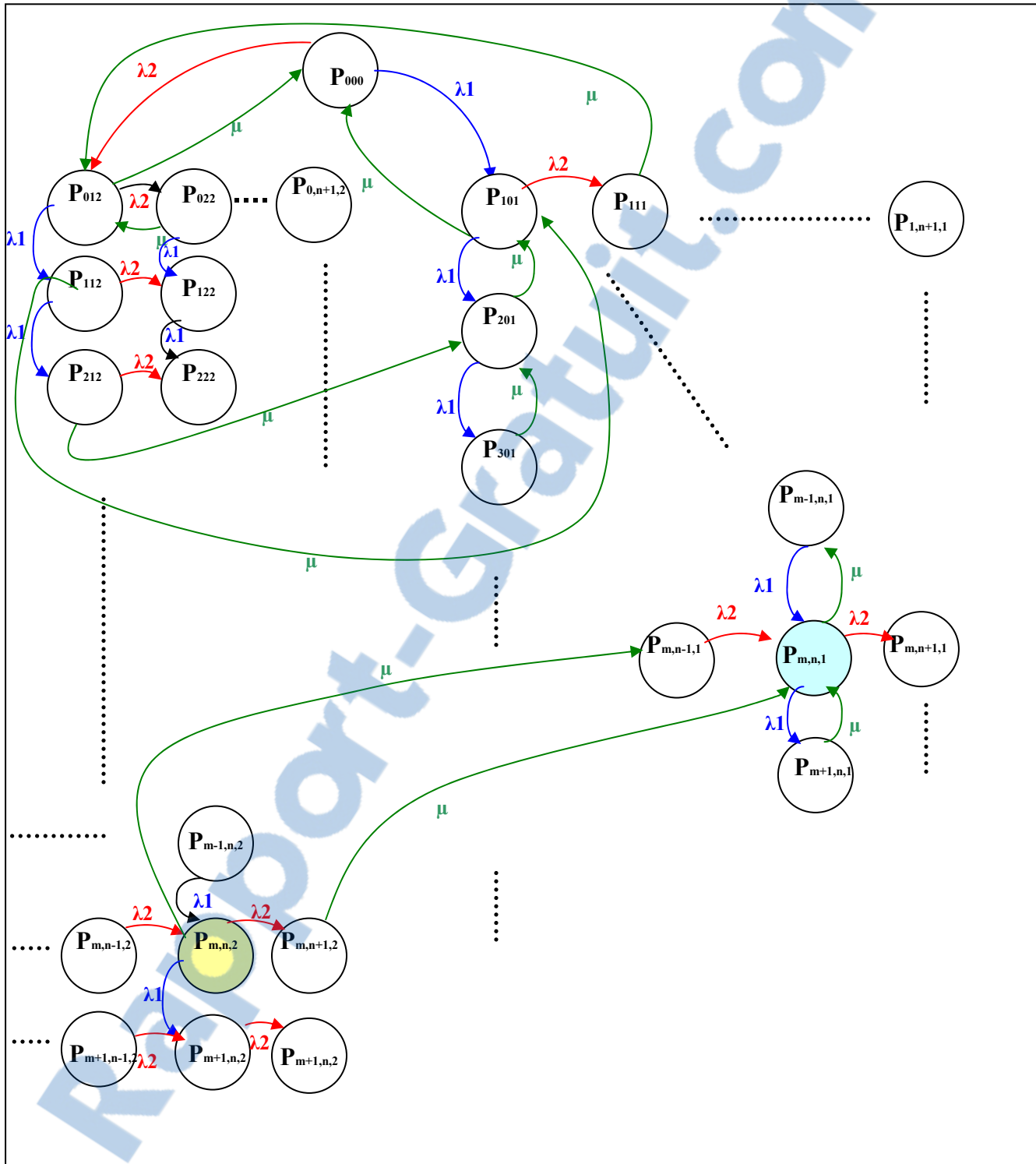


Figure 3.2 Diagramme d'état du système de files d'attente M/G/1 avec PQ.

Suivant la Figure 3.2, dans l'état d'équilibre, nous aurons les équations différentielles suivantes, selon les valeurs de m et de n :

- **m>0, n>1**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,1} = \lambda_1 P_{m-1,n,1} + \lambda_2 P_{m,n-1,1} + \mu(P_{m+1,n,1} + P_{m,n+1,2})$$

- **m>1, n=0**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,2} = \lambda_1 P_{m-1,n,2} + \lambda_2 P_{m,n-1,2}$$

- **m=0, n=0**

$$(\lambda_1 + \lambda_2)P_{0,0,0} = \mu(P_{1,0,1} + P_{0,1,2})$$

- **m=0, n>1**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,2} = \lambda_2 P_{m,n-1,2} + \mu P_{1,1} + \mu P_{m,n+1,2}$$

- **m=1, n>0**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,1} = \lambda_2 P_{m,n-1,1} + \mu P_{m+1,n,1} + \mu P_{m,n+1,2}$$

- **m>1, n=0**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,1} = \lambda_1 P_{m-1,n,1} + \mu P_{m+1,n,1} + \mu P_{m,n,2}$$

- **m>0, n=1**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,1,2} = \lambda_1 P_{m-1,n,2}$$

- **m=0, n=1**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,2} = \lambda_2 P_{0,0,0} + \mu P_{m+1,n,1} + \mu P_{m,n+1,2}$$

- **m=1, n=0**

$$(\lambda_1 + \lambda_2 + \mu)P_{m,n,1} = \lambda_1 P_{0,0,0} + \mu P_{m+1,n,1} + \mu P_{m,n+1,2}$$

Avec: selon (Gross 2008)

$$P_{0,0,0} = 1 - \rho \quad \text{avec} \quad \rho = \lambda/\mu \quad (3.10)$$

$$P_n = \sum_{m=0}^{n-1} (P_{n-m,m,1} + P_{m,n-m,2}) \quad (3.11)$$

En prenant en compte l'étude faite en (Gross 2008) on pourra considérer le système de formule (3.12), pour le diagramme d'état illustré dans la Figure 3.2.

$$\begin{aligned} L1 &= \frac{(\lambda 1/\mu) (1 + \rho - (\lambda 1/\mu))}{1 - (\lambda 1/\mu)} \\ Lq1 &= \frac{\lambda 1/\mu}{1 - \lambda 1/\mu} \\ Wq1 &= \frac{\rho}{\lambda 1 - \mu} \\ L2 &= \frac{(\lambda 2/\mu) (1 + \rho (\lambda 1/\mu) - (\lambda 1/\mu))}{(1 - \rho) (1 - (\lambda 1/\mu))} \\ Lq2 &= \frac{\rho \lambda 2/\mu}{(1 - \rho) (1 - \lambda 1/\mu)} \\ Wq2 &= \frac{\rho}{(1 - \rho) (\lambda 1 - \mu)} \end{aligned} \quad (3.12)$$

Les équations (3.12) vont permettre de déduire celles qui correspondent à notre modèle, après avoir calculé le nombre de paquets courtois, qu'on note $coef f_{courtois}$ et qui représente le coefficient de courtoisie, et ξ_1 , le temps de tolérance d'attente supplémentaire de ces paquets dans leur queue.

Il est clair que les temps d'attente moyens des paquets de la classe rtPS dans leur queue et dans le système vont augmenter de ξ_1 . En revanche, les paquets de la classe nrtPS auront profité d'une diminution d'attente moyenne dans le système ainsi que dans leur queue de cette même valeur ξ_1 .

L'attente supplémentaire dans la file rtPS va engendrer un accroissement de sa longueur moyenne, ainsi que du nombre moyen de paquets rtPS qui attendent dans le système. Cette valeur additionnelle correspond aux paquets courtois, donc à $Coef f_{courtois}$.

Comme le temps moyen de service est identique pour les deux files d'attente, nous constatons que le nombre de paquets nrtPS, qui tirent profit de notre solution, est similaire aux nombres de paquets courtois. Par conséquent, la longueur moyenne de la queue relative à la classe nrtPS et le nombre de paquets de cette même classe vont se réduire de $coef f_{courtois}$.

Finalement, pour les deux classes de service, nous exprimons les délais moyens d'attente des paquets dans leur file d'attente et dans le système, ainsi que les longueurs moyennes des deux buffers et le nombre des paquets dans le système, par l'ensemble d'équations (3.13).

$$\begin{aligned}
L_{rtPS} &= L_1 + coeff_{courtois} \\
Lq_{rtPS} &= Lq_1 + coeff_{courtois} \\
Wq_{rtPS} &= Wq_1 + \xi_1 \\
L_{nrtPS} &= L_2 - coeff_{courtois} \\
Lq_{nrtPS} &= Lq_2 - coeff_{courtois} \\
Wq_{nrtPS} &= Wq_2 - \xi_1
\end{aligned} \tag{3.13}$$

Calcul de ξ_1 et de $Coef_{courtois}$:

Le temps de tolérance d'attente supplémentaire des paquets prioritaires dans leur file d'attente rtPS_Queue, défini par la valeur ξ_1 , est le temps durant lequel les paquets courtois cèdent leur tour de service aux paquets de basse priorité sans détériorer leur QoS.

Nous avons déjà mentionné que le nombre de paquets rtPS qui cèdent leur tour de service et temporisent un temps supplémentaire ξ_1 dans la queue de haute priorité, est similaire au nombre de paquets nrtPS servis durant ce temps là. Il s'agit de « $coeff_{courtois}$ ». Ainsi, sachant que le temps de service moyen d'un paquet nrtPS est égal à μ , ξ_1 se calculera comme suit :

$$\xi_1 = coeff_{courtois} * \mu \tag{3.14}$$

Une autre interprétation du coefficient de courtoisie peut être le nombre de paquets attendus pour atteindre Th_1 , le seuil de remplissage de la file rtPS_Queue. Ce seuil correspond à une probabilité de perte de paquets égale à ω_1 pour le trafic de classe v. La formule (3.15) sert à calculer ce coefficient :

$$coef_{courtois} = Th_1 - L_1 \quad (3.15)$$

Th_1 se calcule comme suit :

$$Th_1 = K1 - R1 \quad (3.16)$$

R1 est le nombre maximal que puisse contenir une rafale (burst) de classe C1. Sa valeur est donnée par le produit du taux d'arrivées maximal des paquets de la classe C1 et le temps d'attente maximal dans la file rtPS_Queue.

Le taux d'arrivées maximal est la somme du taux d'arrivée moyen et sa variance, correspondants au trafic de la classe C1.

D'autre part, la somme du temps d'attente moyenne des paquets de classe C1 dans le système et la variance de ce temps-ci, permet de donner le temps d'attente maximal dans la file correspondante. On aura donc :

$$R1 = (\lambda_1 + \sigma_{\lambda_1}) * (Wq_1 + \sigma_{Wq_1}) \quad (3.17)$$

De (3.16) et (3.17), l'équation (3.15) va se donner comme suit :

$$coef_{courtois} = (K1 - ((\lambda_1 + \sigma_{\lambda_1}) * (Wq_1 + \sigma_{Wq_1}))) - L_1 \quad (3.18)$$

En remplaçant (3.18) dans (3.14) nous aurons la formule finale qui sert à calculer le temps de tolérance d'attente supplémentaire des paquets de haute priorité dans leur file. Ce temps sera celui que nous allouons aux paquets de la classe C2 pour être servis.

$$\xi_1 = ((K1 - ((\lambda_1 + \sigma_{\lambda_1}) * (Wq_1 + \sigma_{Wq_1}))) - L_1) * \mu \quad (3.19)$$

Calcul de la priorité

Comme nous l'avons déjà mentionné, notre modèle de files d'attente traite les paquets arrivés par ordre de priorités. Avant d'en servir un. Le système doit calculer la priorité des paquets en tête de chacune des deux files d'attente. Celui qui détient la valeur la plus haute sera traité en premier.

Notons qu'un paquet du trafic de la classe $C1$ détient la priorité $Pr1$, et celui de la classe $C2$ la priorité $Pr2$.

Ce mécanisme de gestion de file d'attente, à savoir un système M/G/1 avec priorité non préemptive incluant deux files d'attente, permet de constater qu'en tout instant t , il y aura i paquets de priorité $Pr1$ dans la file $rtPS_Queue$, j paquets de priorité $Pr2$ dans la file $nrtPS_Queue$ et 1 paquet de priorité $PrMax$ dans le serveur, telle que :

$$PrMax = \text{Max} (Pr1, Pr2) \quad (3.20)$$

Afin de pouvoir calculer $PrMax$, il faut d'abord déterminer les valeurs $Pr1$ et $Pr2$. Nous supposons que :

$$0 \leq Prk \leq 1, \text{ tel que } k = \{1,2\} \quad (3.21)$$

Soient β_1 et β_2 les poids relatifs aux files d'attente $rtPS_Queue$ et $nrtPS_Queue$. Ils sont proportionnels aux pourcentages de bande passante allouée à la classe C_1 et à la classe C_2 , respectivement. Notons que :

$$\sum \beta_i \leq 1 \quad | i = \{1,2\} \quad (3.22)$$

$$0 \leq \beta_2 < \beta_1 \leq 1 \quad (3.23)$$

La formule (3.23) permet de favoriser le trafic rtPS par rapport au trafic nrtPS. Ceci implique l'application du Diffserv dans notre solution.

D'autres parts soient W_{s1} et W_{s2} les temps d'attente dans le système, correspondants aux paquets de la classe $C1$, et ceux de la classe $C2$, respectivement, tel que :

$$W_{sk} = W_{qk} + \mu \quad | \quad k = \{rtPS, nrtPS\} \quad (3.24)$$

Rappelons que Th_1 est le seuil de remplissage de la file d'attente rtPS_Queue, correspondant à une probabilité de perte de paquets de classe $C1$ égale à ω_1 .

Soit ψ le taux de congestion dans le canal de transmission.

La priorité d'un paquet pk dépend de plusieurs facteurs, à savoir :

- La classe à laquelle appartient le paquet. Notre objectif est de donner une priorité aux paquets du trafic rtPS. Ceci est faisable en considérant la formule (3.23).
- Le temps d'attente du paquet en question dans le système, à savoir W_{s1} si le paquet appartient à $C1$, et W_{s2} s'il appartient à la classe $C2$. Dans cette solution, les paquets qui attendent plus dans le système, se voient leur priorité augmentée. Ceci permet de réduire la perte de paquets.
- La congestion ψ dans le canal de transmission, de telle manière que plus la valeur de ψ augmente, plus la priorité du paquet augmente. Ce qui multiplie sa chance être servie plutôt que d'être rejeté à cause d'une attente excessive dans le système.

Ces trois facteurs vont permettre de conclure les formules de calcul de priorité que nous présentons comme suit :

$$Pr1 = \beta_1 + f(W_{s1}, \psi) \quad si \quad f(W_{s1}, \psi) \leq 1 - \beta_1 \quad (3.25)$$

$$Pr2 = \begin{cases} \beta_2 + f(W_{s2}, \psi) & \text{si } f(W_{s2}, \psi) \leq 1 - (\beta_1 + \beta_2) \quad L_{rtPS} \geq Th_1 \\ Pr1 & \text{si } L_{nrtPS} \geq Th_2 \quad L_{rtPS} < Th_1 \end{cases} \quad (3.26)$$

Les formules (3.25) et (3.26), permettent de calculer la priorité d'un paquet rtPS ou nrtPS. Comme nous l'avons déjà mentionné, la priorité d'un paquet dépend du poids que détient la classe à laquelle il appartient. Cette valeur est reflétée par l'un des deux facteurs β_1 ou β_2 , selon sa classe de trafic. Un autre élément qui contribue à la détermination de cette priorité est la fonction $f(W_{si}, \psi)$ avec $i = \{1,2\}$. Celle-ci sert à mesurer la congestion et le temps d'attente d'un paquet dans le système.

- **Calcul de $f(W_{si}, \psi)$ | $i = \{1,2\}$:**

La congestion dans le canal de transmission peut s'interpréter par la fraction du nombre de paquets des deux types de trafic dans le système, et la somme des tailles des buffers et du serveur. On peut la calculer comme suit :

$$\psi = \frac{L_{rtPS} + L_{nrtPS}}{K1 + K2 + 1} \quad (3.27)$$

Rappelons que L_{rtPS} et L_{nrtPS} sont le nombre de paquets dans le système relatif, respectivement, au trafic de classes $C1$ et $C2$, et $K1$ et $K2$ sont les tailles respectives de rtPS_Queue et nrtPS_Queue. Tandis que le 1 représente le nombre maximal de paquets qui peuvent se trouver dans le serveur.

Reste à déterminer les formules relatives au temps d'attente dans le système W_{s1} et W_{s2} . Les résultats doivent être sans unités. Ceci nous mène à considérer les deux fractions suivantes :

$$Fract_W_{s1} = W_{s1} / (W_{s1} + W_{s2}) \quad (3.28)$$

$$Fract_{W_{s2}} = W_{s2} / (W_{s1} + W_{s2}) \quad (3.29)$$

$f(W_{s1}, \psi)$ et $f(W_{s2}, \psi)$ se calculent comme suit:

$$f(W_{s1}, \psi) = Fract_{W_{s1}} + \psi \quad (3.30)$$

$$f(W_{s2}, \psi) = Fract_{W_{s2}} + \psi \quad (3.31)$$

En remplaçant (3.27) et (3.28) dans (3.30), et (3.27) et (3.29) dans (3.31), on déduira les formules suivantes :

$$f(W_{s1}, \psi) = W_{s1} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1) \quad (3.32)$$

$$f(W_{s2}, \psi) = W_{s2} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1) \quad (3.33)$$

Or, nous avons la supposition suivante :

$$Pr1 \leq 1 \quad (3.34)$$

$$Pr2 \leq (1 - \beta_1) \quad (3.35)$$

Tenant compte des équations données en (3.25) et (3.26), les équations (3.34) et (3.35) seront reformulées comme suit :

$$f(W_{s1}, \psi) \leq 1 - \beta_1 \quad (3.36)$$

$$f(W_{s2}, \psi) \leq 1 - \beta_1 - \beta_2 \quad (3.37)$$

Par conséquent, on aura le système d'équations suivant :

$$f(W_{s1}, \psi) = \begin{cases} W_{s1} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1) & \text{si } f(W_{s1}, \psi) \leq 1 - \beta_1 \\ 0 & \text{sinon} \end{cases} \quad (3.38)$$

$$f(W_{s2}, \psi) = \begin{cases} W_{s2} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1) & \text{si } f(W_{s2}, \psi) \leq 1 - \beta_1 - \beta_2 \\ 0 & \text{sinon} \end{cases} \quad (3.39)$$

Finalement, nous aurons le système de calcul de priorité suivant :

$$Pr1 = \begin{cases} \beta_1 + (W_{s1} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1)) & \text{si } f(W_{s1}, \psi) \leq 1 - \beta_1 \\ \beta_1 & \text{sinon} \end{cases} \quad (3.40)$$

$$Pr2 = \begin{cases} \beta_2 + W_{s2} / (W_{s1} + W_{s2}) + (L_{rtPS} + L_{nrtPS}) / (K1 + K2 + 1) & \text{si } f(W_{s2}, \psi) \leq 1 - (\beta_1 + \beta_2) \text{ et } L_{rtPS} \geq Th_1 \\ Pr1 & \text{si } L_{nrtPS} \geq Th_2 \text{ et } L_{rtPS} < Th_1 \\ \beta_2 & \text{si } f(W_{s2}, \psi) > 1 - (\beta_1 + \beta_2) \end{cases} \quad (3.41)$$

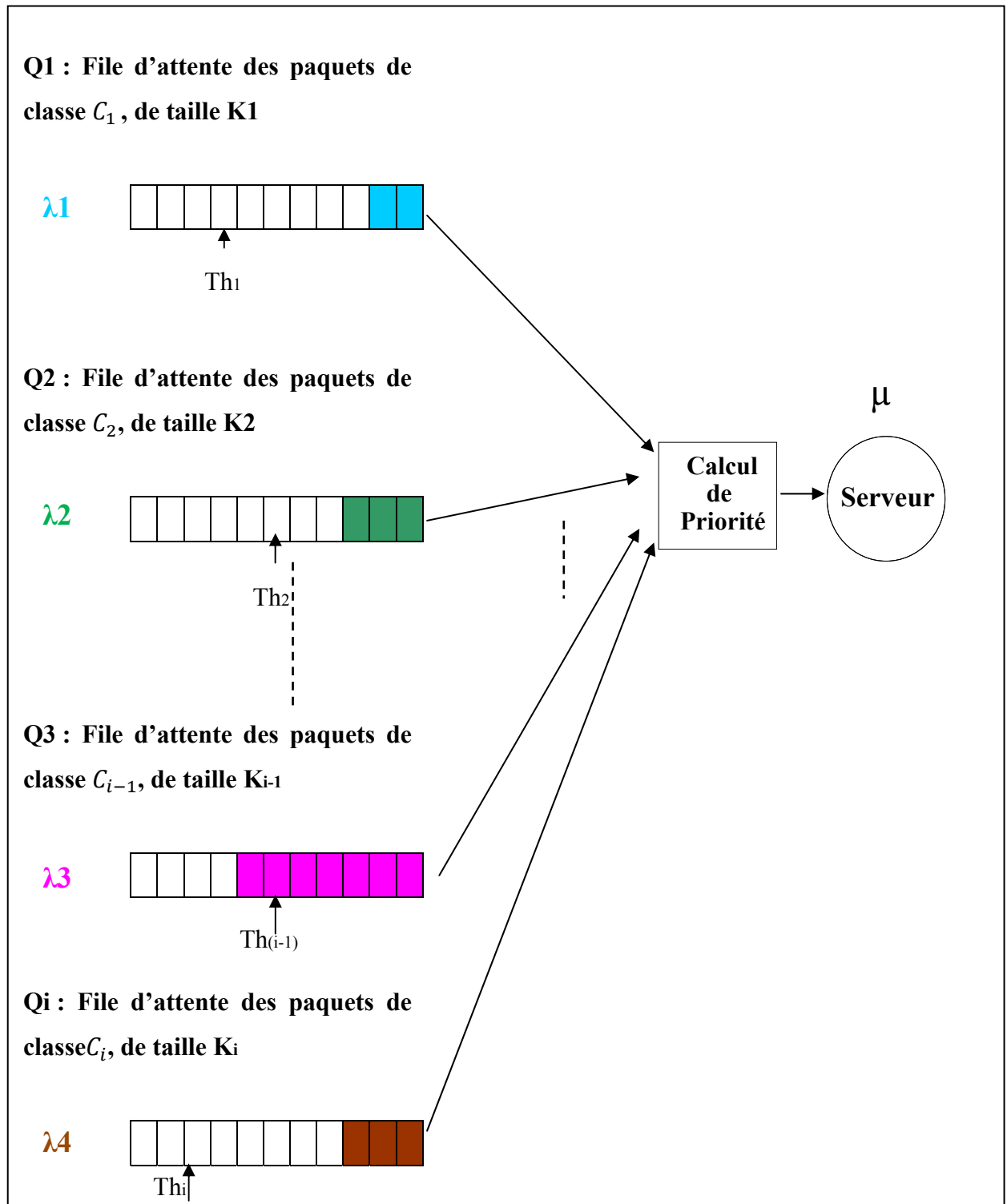


Figure 3.3 Système de file d'attente M/G/1 avec courtoisie (n files d'attente).

La condition $f(W_{s1}, \psi) \leq 1 - \beta_1$ est une conséquence de $Pr1 \leq 1$. La condition $f(W_{s2}, \psi) \leq 1 - (\beta_1 + \beta_2)$ est dû au fait que $Pr2 \leq 1$ et qu'un pourcentage β_1 de bande passante a été alloué pour le trafic de classe $C1$.

Voir les formules (3.13) et (3.24) pour le calcul de W_{qi} et W_{sj} , $i = rtPS, nrtPS$, $j = 1, 2$.

3.4.2 Cas de n files d'attente

Système de file d'attente :

La Figure 3.3 est une généralisation du cas d'un système de file d'attente M/G/1 avec priorité non préemptive contenant 2 files d'attente.

Dans cette section, nous considérons le système de files d'attente représenté dans la Figure 3.3 et qui est composé de n files d'attente Q_i , v , correspondantes à i classes de service C_i , tel que $i=1, 2, \dots, n$.

Les arrivées du trafic de classe C_i suivent une loi de poisson de taux moyen d'arrivées λ_i paquet/s. Notons λ le taux moyen total d'arrivées par seconde dans le système. Nous aurons

$$\text{alors } \lambda = \sum_{i=1}^k \lambda_i$$

Le service des paquets de la classe C_i , suit une loi exponentielle de paramètres $1/\mu$ paquet/s.

Le temps d'interarrivées des paquets de la classe C_i , suit une loi exponentielle de paramètre $1/\lambda_i$ sec.

Les longueurs moyennes des files d'attente Q_i , tel que $i \in \{1, \dots, n\}$, noté L_{qi} , suivent une loi géométrique. Il est de même pour L_i , les nombres de paquets moyens de classe C_i dans le système.

Calcul de ξ_{i-1} :

Dans cette section, nous supposons que le réseau contient i classes de trafic où $i = \{1, 2, \dots, k\}$. Nous allons calculer ξ_{i-1} , le temps de tolérance d'attente supplémentaire des paquets de la classe de trafic $C_{(i-1)}$. Nous supposons que $Pr_1 > Pr_2 > \dots > Pr_k$, tels que $k = \{1, \dots, n\}$ et Pr_k est la priorité de la classe k . Pour cela nous considérons le système de files d'attente représenté dans la figure 3.3.

L'équation suivante va nous permettre de calculer la valeur de ξ_{i-1} , tel que $i \in \{1, \dots, k\}$.

$$\xi_{i-1} = coeff_{courtoisie(i-1)} * \mu \quad (3.42)$$

D'un raisonnement similaire à celui effectué en section (3.3.1.2), nous pourrions déduire les résultats suivants :

$$R_{i-1} = (\lambda_{i-1} + \sigma_{\lambda_{i-1}}) * (Wq_{i-1} + \sigma_{Wq_{i-1}}) \quad (3.43)$$

$$coeff_{courtoisie(i-1)} = K_{i-1} - \left((\lambda_{i-1} + \sigma_{\lambda_{i-1}}) * (Wq_{i-1} + \sigma_{Wq_{i-1}}) \right) - L_{i-1} \quad (3.44)$$

Similairement au cas de deux files d'attente, R_{i-1} est le nombre maximal que peut contenir une rafale (burst) de classe C_{i-1} . Sa valeur est donnée par le produit du taux d'arrivées maximal des paquets de cette classe et le temps d'attente maximal dans la file correspondante.

$coeff_{courtoisie(i-1)}$ est le nombre de paquets qui cèdent leur tour de service aux paquets de la classe C_i . Notons qu'étant donné que le temps de service moyen est identique pour tous les paquets, quel que soit leur priorité, alors nous pouvons considérer que $coeff_{courtoisie(i-1)}$ est aussi le nombre de paquets de classe C_i Bénéficiant de la courtoisie.

Finalement, le temps de tolérance d'attente supplémentaire dans la file d'attente de classe C_{i-1} sera donné par la formule suivante :

$$\xi_{i-1} = [K_{i-1} - ((\lambda_{i-1} + \sigma_{\lambda_{i-1}}) * (Wq_{i-1} + \sigma_{Wq_{i-1}})) - L_{i-1}] * \mu \quad (3.45)$$

Notons que :

- K_{i-1} : est la taille de la file d'attente des paquets courtois.
- λ_{i-1} : est le taux d'arrivée moyen des paquets de classe C_{i-1}
- $\sigma_{\lambda_{i-1}}$: c'est la variance de λ_{i-1}
- Wq_{i-1} : est le temps d'attente moyen des paquets de classe C_{i-1} dans leur file
- $\sigma_{Wq_{i-1}}$: est la variance de Wq_{i-1}
- μ est le temps de service moyen d'un paquet. Tous les paquets ont le même temps de service moyen.

En tenant compte de l'étude effectuée dans (Gross 2008), concernant les systèmes de files d'attente M/G/1 avec plusieurs priorités, non préemptifs, nous pouvons considérer les formules suivantes pour le calcul du Wq_{i-1} , le temps d'attente dans la file Q_{i-1} , et sa longueur, $L_{q_{i-1}}$:

$$Wq_{i-1} = \frac{\sum_{k=1}^{i-1} \frac{\rho_{k-1}}{\mu}}{(1-\sigma_{i-1})*(1-\sigma_i)} \quad (3.46)$$

$$\sigma_r = \sum_{k=1}^r \rho_k < 1 \quad (3.47)$$

$$L_{q_{i-1}} = \frac{\lambda_{i-1} \sum_{k=1}^n \rho_k / \mu}{(1-\sigma_{i-1})(1-\sigma_i)} \quad (3.48)$$

« i-1 » étant le rond de la classe C_{i-1} .

Calcul de la priorité - Généralisation à i classes de trafic

Tenant compte de l'analyse effectué dans la section (3.3.1.3) pour le calcul de la priorité des paquets dans un système de gestion de files d'attente avec priorité et courtoisie, nous pouvons déduire les formules qui déterminent les valeurs de priorité pour n classes de trafic.

En plus des deux formules (3.40) et (3.41) qui sert à donner les valeurs des deux premières priorités relatives aux deux classes les plus prioritaires, nous présentons ci-dessous les équations qui permettent de calculer les restantes.

$$Pr3 = \left\{ \begin{array}{ll} \beta_3 + W_{s3} / (W_{s1} + W_{s2} + W_{s3}) + (L_1 + L_2 + L_3) / (K1 + K2 + K3 + 1) & \\ \quad si \ f(W_{s3}, \psi) \leq 1 - (\beta_1 + \beta_2 + \beta_3) \text{ et } L_2 \geq Th_2 & \\ Pr2 & si \ L_3 \geq Th_3 \text{ et } L_2 < Th_2 \\ \beta_3 & si \ f(W_{s3}, \psi) > 1 - (\beta_1 + \beta_2 + \beta_3) \end{array} \right. \quad (3.49)$$

⋮

$$Pri = \left\{ \begin{array}{ll} \beta_i + W_{si} / (\sum_{j=1}^i W_{sj}) + (\sum_{j=1}^i L_j) / (\sum_{j=1}^i K_j + 1) & \\ \quad si \ f(W_{si}, \psi) \leq 1 - ((\sum_{j=1}^i \beta_j)) \text{ et } L_{i-1} \geq Th_{i-1} & \\ Pr_{i-1} & si \ L_i \geq Th_i \ L_{i-1} < Th_{i-1} \\ \beta_i & si \ f(W_{si}, \psi) > 1 - ((\sum_{j=1}^i \beta_j)) \end{array} \right. \quad (3.50)$$

3.5 Importance de l'analyse mathématique

L'analyse mathématique est primordiale pour la conception de notre algorithme. Plusieurs facteurs influent sur le comportement de notre système de file d'attente. Par conséquent, une bonne gestion des différents types de trafic nécessite la mesure de ces facteurs, à savoir le temps d'attente supplémentaire tolérée pour les paquets courtois, noté ξ_1 dans le modèle analytique à deux files d'attente, le nombre de paquets courtois exprimé par la valeur de $coef_{courtois}$, qui détermine aussi le nombre de paquets qui bénéficient de notre approche, la taille maximale d'une rafale, notée par R_1 dans le cas de deux files d'attente et R_i dans le cas généralisé. Les seuils de remplissage des files d'attente ont aussi un impact sur notre modèle du moment qu'ils soient relatifs au seuil de pertes de paquets tolérées pour les trafics en question.

Vu la complexité de l'analyse mathématique, nous allons nous contenter des résultats obtenus jusque-là. Ces résultats vont être exploités pour la conception de l'algorithme de courtoisie. La validation de notre algorithme va se faire à travers la simulation dans le chapitre 4.

3.6 Structure de l'algorithme de courtoisie

La Figure 3.4 représente l'algorithme de courtoisie. Dans cette section nous allons appliquer l'analyse mathématique réalisée dans la section précédente, sur un réseau WiMAX à deux classes de service.

Les deux types de trafic pris en considération sont rtPS pour le trafic de classe C_1 et nrtPS pour le trafic de la classe C_2 . Les files d'attente correspondantes sont respectivement rtPS_Queue et nrtPS_Queue.

Comme nous l'avons déjà mentionné, ω_1 et ω_2 représentent, respectivement, les seuils des taux de perte de paquets rtPS et nrtPS tolérées dans ce système. Ces deux valeurs sont

proportionnelles à Th_1 et Th_2 , respectivement. Ces derniers sont les seuils de remplissage des files d'attente $rtPS_Queue$ et $nrtPS_Queue$.

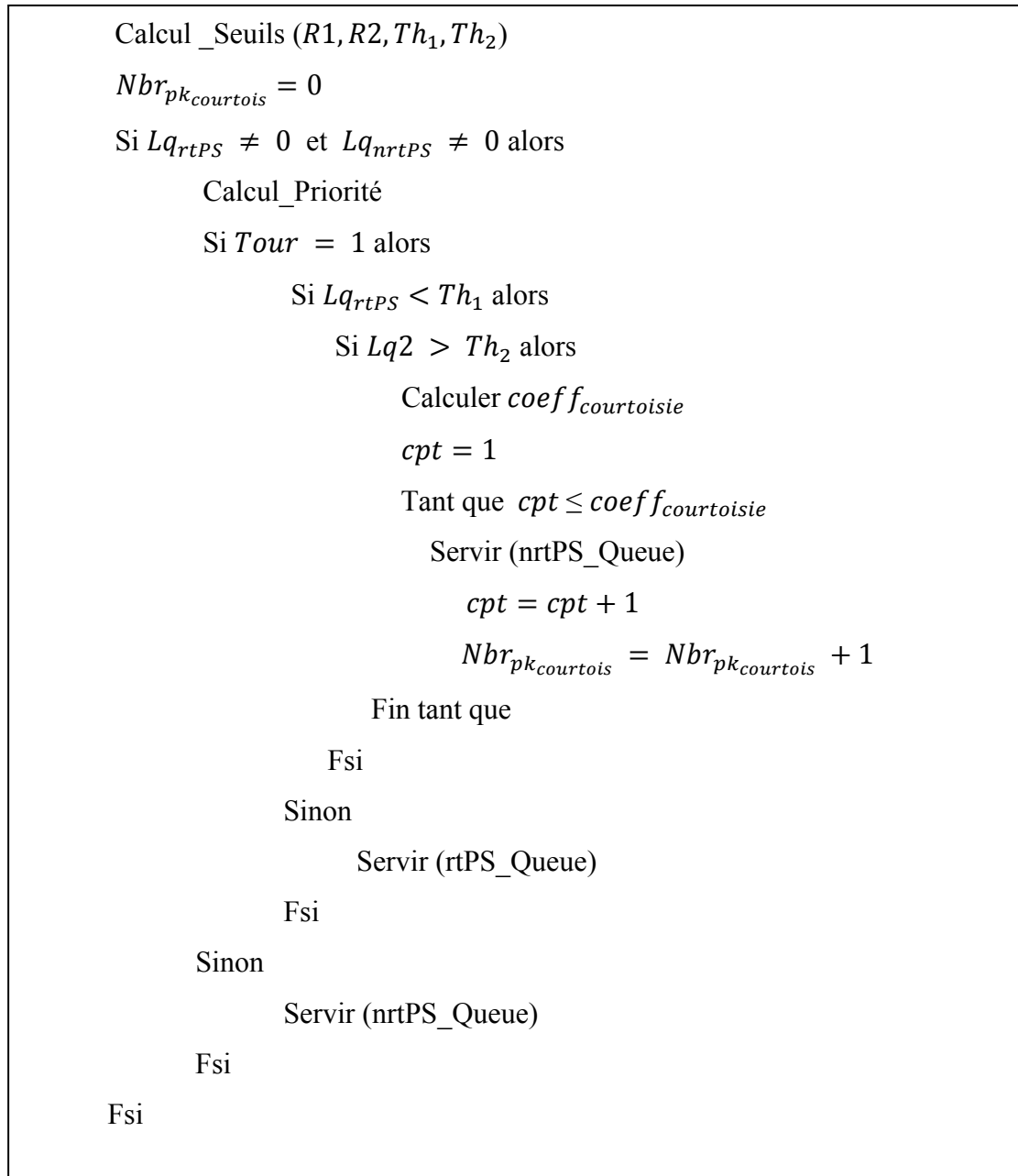


Figure 3.4 Algorithme de courtoisie pour le réseau WiMAX : Deux classes de service.

Cette solution consiste à calculer la priorité des paquets rtPS et nrtPS, résidants dans les files d'attente, notre système est illustré par la Figure 3.1. C'est le résultat que donne la fonction *Calcul_Priorité*, présentée dans la figure 3.5. Cette procédure renvoie la valeur de *Tour*. Dans le cas où *Tour* est égale à 1 alors c'est la file rtPS_Queue qui sera servie, sinon c'est l'autre file qui sera servie.

Si le tour de service est pour la file rtPS, alors les paquets qui y résident peuvent se montrer courtois dans le cas où le seuil Th_1 n'est pas encore atteint. Ceci s'interprète par une valeur de $coeff_{courtoisie}$ supérieure ou égale à 1.

Pour chaque itération de courtoisie, le nombre de paquets nrtPS à servir est égale à une valeur que nous représentons par le symbole *cpt*, de telle sorte que cette valeur est inférieure ou égale à $coeff_{courtoisie}$. Rappelons que $coeff_{courtoisie}$ est le nombre de paquets attendus pour atteindre le seuil Th_{rtps} relatif à la probabilité de perte de paquets tolérés ω_1 correspondante à la classe rtPS du réseau WiMAX en question. Dans notre solution on initialise *cpt* à 1.

La valeur de $Nbr_{pk_{courtois}}$ va permettre d'identifier le nombre total des paquets rtPS courtois durant un intervalle de temps prédéfini. Comme le temps de service moyen est le même pour tous les paquets, quelle que soit leur classe de service $Nbr_{pk_{courtois}}$ sera aussi égale au nombre de paquets nrtPS ayant bénéficié de notre solution, ce qui nous permet de mesurer la performance apportée par notre algorithme.

Le temps de tolérance d'attente supplémentaire relative aux paquets de la classe rtPS est relatif au $coeff_{courtois}$, représentant le nombre de paquets qui cèdent leur tour de service. Ces deux facteurs étant proportionnel, nous avons opté pour l'utilisation de $coeff_{courtois}$ pour une raison de simplification de notre algorithme. Les seuils $R1, R2, Th_1, Th_2$, permettent de calculer la valeur de $coeff_{courtois}$, et de décider si l'application de notre algorithme répond aux conditions de son utilisation (voir les sections 3.2 et 3.3).

Il est clair que les paquets rtPS ne vont pas céder leur place dans les deux cas suivants :

- La probabilité de perte de paquets nrtPS n'est pas importante. Ce qui s'explique par la non-atteinte de Th_2
- Le temps de service des paquets nrtPS excède la valeur du temps d'attente supplémentaire tolérée pour un paquet rtPS. Ce qui donne $\tau_2 > \xi_1$.

$$Fpr1 = f(W_{s1}, \psi) \quad (\text{voir l'équation 4.38})$$

$$Fpr2 = f(W_{s2}, \psi) \quad (\text{voir l'équation 4.39})$$

Si $Fpr1 > 1 - \beta_1$ alors

$$Pr1 = \beta_1$$

Sinon

$$Pr1 = \beta_1 + Fpr1$$

Fin si

Si $Fpr2 > 1 - \beta_1 - \beta_2$ alors

$$Pr2 = \beta_2$$

Sinon

$$Pr2 = \beta_2 + Fpr2$$

Fin si

$$MaxPr = \text{Max}(Pr1, Pr2)$$

Si $MaxPr = Pr1$ alors

$$Tour = 1$$

sinon

$$Tour = 2$$

Fsi (15)

Figure 3.5 Fonction Calcul_Priorité.

La Figure 3.5 représente la fonction Calcul_Priorité. Celle-ci permet de calculer les priorités des paquets rtPS et nrtPS, pour déterminer lequel des deux paquets est plus prioritaire. Pour ce faire, elle renvoie la valeur de Tour qu'utilise l'algorithme de courtoisie. Le système pourra donc décider à qui est le rôle de traitement parmi les paquets des deux classes $C1$ et $C2$.

$$\begin{aligned} R1 &= (\lambda1 + \sigma_{\lambda1}) * (Wq1 + \sigma_{Wq1}) \\ R2 &= (\lambda2 + \sigma_{\lambda2}) * (Wq2 + \sigma_{Wq2}) \\ Th_1 &= K1 - R1 \\ Th_2 &= K2 - R2 \end{aligned}$$

Figure 3.6 Fonction Calcul _Seuils.

La Figure 3.6 représente la fonction Calcul_Seuils ($R1, R2, Th_1, Th_2$) utilisée dans l'algorithme de courtoisie. Celle-ci consiste à calculer les seuils Th_1 et Th_2 correspondant aux seuils probabilités de pertes de paquets tolérés ω_1 et ω_2 relatifs aux files d'attente rtPS_Queue et nrtPS_Queue, respectivement. En effet, Th_1 et Th_2 vont assurer la satisfaction des conditions d'application de l'algorithme de courtoisie.

Cette procédure permet aussi de calculer les valeurs de $R1$ qui sont nécessaires pour déterminer $coeff_{courtoisie}$, qui est le nombre de paquets qui peuvent céder leur tour de service au profil de la classe nrtPS, attendant ainsi une durée égale au temps de tolérance d'attente supplémentaire ξ_1 .

3.7 Conclusion

Tel que nous l'avons mentionné au début de ce chapitre, le WiMAX a défini quatre classes de services afin d'assurer une différenciation de service et garantir la transmission des paquets en respectant leurs exigences et en considérant leurs tolérances. Néanmoins, ce standard a omis de déterminer le mécanisme adéquat et pertinent pour l'ordonnancement de ces trafics. Ce point ouvert a suscité le développement de plusieurs solutions tentant ainsi de proposer la meilleure approche apte à optimiser la performance de la QoS dans les réseaux IEEE802.16. Or, la majorité des solutions existantes dans la littérature défavorise les trafics de basses priorités. Ce qui ne va point apporter la satisfaction entière aux utilisateurs de ces réseaux.

Conformément aux besoins de combler la totalité des abonnés. Notre approche vient compléter les solutions favorisant les trafics de hautes exigences. Elle consiste à optimiser la QoS pour les trafics appartenant aux classes moins prioritaires en termes de délais et de pertes de paquets, en leurs permettant d'être servi à la place des paquets plus prioritaires pendant un temps ξ_1 . Celui là est considéré comme le temps d'attente supplémentaire tolérée, correspondant au trafic courtois. Cette attente additionnelle n'affecte point la QoS de la classe prioritaire.

L'application de l'algorithme de courtoisie ne s'applique pas à tous les coups. Quatre conditions ont été définies pour que le système puisse avoir recours à cette solution. L'objectif de ces contraintes est d'empêcher les classes prioritaires de venir en aide aux autres classes dans le cas où elles n'ont pas besoin, autrement dit, leur taux de perte de paquets n'est pas considérable. En outre, ces conditions vont garantir aux classes courtoises la préservation de leur QoS en refusant ou cessant l'exécution de cette solution dès que le taux de pertes de paquets qui leur sont correspondant atteint un seuil donné.

Un modèle mathématique relatif à notre système de files d'attente a été entamé pour le cas de deux files d'attente rtPS et nrtPS du réseau WiMAX et aussi pour le cas de plusieurs files

d'attente. L'analyse accomplie dans ce chapitre est primordiale pour la conception de l'algorithme de courtoisie, car elle a permis de déterminer les formules relatives aux calculs des facteurs influant sur le comportement de notre modèle de gestion de files d'attente, tels que ξ_1 qui représente le temps d'attente supplémentaire toléré pour les paquets de la classe de haute priorité, et $coeff_{courtoisie}$ qui détermine le nombre de paquet de la classe de basse priorité qui bénéficient de la courtoisie, ainsi que la taille maximale d'une rafale que nous avons appelé R1 dans le cas d'un système de gestion de files d'attente à deux classes de service. Ces formules ont été exploitées pour la conception de l'algorithme de courtoisie.

La solution mathématique se montre assez complexe à achever. Cette complexité demeure dans la difficulté de présenter une formule fermée correspondante à notre système M/G/1. Un tel résultat demande un développement et un travail plus poussés et nécessite beaucoup plus de temps. Par conséquent nous n'avons pas eu des résultats numériques pour les délais et les longueurs des files d'attente. Pour cet effet, nous avons opté pour la simulation afin de valider notre approche. Notons que la simulation nous a permis de tester plusieurs cas de figures, ce qui a rendu notre étude plus complète. Cette partie sera traitée dans le prochain chapitre.

CHAPITRE 4

SIMULATIONS ET RÉSULTATS

4.1 Introduction

Comme nous l'avons mentionné dans le chapitre précédent, nous avons opté pour la simulation afin de tester et valider notre solution, à savoir l'algorithme de courtoisie nommé Courtesy Priority Queuing (CPQ).

Dans l'étude mathématique, nous avons montré que certains facteurs, tels que λ_1 , λ_2 , R_1 et ξ_1 ont un impact direct sur le comportement de notre modèle. Dans le but de mesurer l'effet de chaque paramètre, nous avons réalisé un ensemble de scénarios en faisant varier les valeurs de ces facteurs. L'étude du comportement de notre solution (CPQ) et sa comparaison avec celui de PQ et WFQ, nous ont permis de déterminer les cas recommandés pour l'application de notre approche.

Dans le reste de ce chapitre, nous allons exposer notre approche adaptée dans les simulations, expliquer les modèles PQ et WFQ considérés dans nos comparaisons, présenter les paramètres de tests pour chaque scénario, ainsi que ses résultats numériques et graphiques. Nous allons conclure par une synthèse et conclusion.

4.2 Approche adaptée dans les simulations

Nous avons réalisé un ensemble de scénarios où nous avons implémenté un système de gestion de files d'attente M/G/1 avec priorité non préemptive au sein d'une station de base d'un réseau WiMAX pour le canal montant (uplink). Pour chaque scénario, nous avons appliqué les algorithmes d'ordonnancement CPQ, PQ et WFQ séparément, en considérant les mêmes paramètres de bases, tels que les taux d'arrivées des paquets de types rtPs et nrtPS, la

taille des files d'attente et leur poids, la taille maximale d'une rafale du trafic voix, le délai d'attente maximal d'un paquet rtPS, etc.

Le choix de PQ comme premier algorithme de comparaison est dû au fait qu'il permet de déterminer le comportement d'un système de files d'attente favorisant exclusivement les paquets prioritaires, ce qui permet de déterminer les plus petits délais et taux de pertes de paquets que puisse avoir la classe de trafic rtPS. Quant au choix de WFQ comme une deuxième solution de confrontation est dû au fait que cette approche est très répondue dans les solutions proposées dans la littérature. De plus, cet algorithme garantit l'équité de service des différents types de trafic dans les réseaux sans fil, particulièrement dans les réseaux WiMAX, surtout que notre algorithme « CPQ » tient à assurer un service équitable pour les différentes classes en calculant et considérant la priorité de chaque paquet avant sa sélection par le serveur de la file d'attente pour transmission.

Pour commencer, nous avons effectué un scénario de référence, nous avons choisi nos paramètres de telle façon à éliminer les pertes de paquets. Ce scénario représente le cas stable où le nombre d'arrivées des paquets des deux types de trafic n'excède pas la capacité du canal. D'autres scénarios ont été accomplis en changeant à chaque fois les valeurs des inter-arrivées rtPS et nrtPS, ils reflètent l'impact de l'intensité du trafic de la voix et celui des données sur le comportement de notre algorithme et les deux algorithmes PQ et WFQ. Nous avons aussi augmenté la taille de notre échantillon, à savoir le nombre de paquets arrivés, afin d'obtenir des résultats plus significatifs. Un dernier scénario a été réalisé pour tester l'impact du paramètre $R1$, la taille maximale d'une rafale, sur le comportement de notre système, bien que la formule (3.17) permet le calcul de $R1$, mais pour des raisons de simplifications nous avons opté à le fixer dans nos simulations. Notons que plus la valeur de $R1$ est importante plus le temps de tolérance d'attente supplémentaire ξ_1 relative aux paquets de la voix sera réduit.

Nous avons considéré un délai d'attente maximal pour le trafic voix dans chaque scénario. Étant donné que le délai de bout en bout acceptable pour la voix ne doit pas dépasser 200 ms

et que l'acheminement d'un paquet de la SS vers la BS passe par plusieurs sauts, nous devons donc choisir une valeur raisonnable pour le délai d'attente maximal dans la file rtPS. Dans nos scénarios nous avons opté pour des délais égaux à 10 ms et à 20 ms. Un paquet qui réside dans la file d'attente de la voix plus que ce délai va être rejeté par le système. Ceci à été appliqué dans les trois modèles, à savoir PQ, WFQ et CPQ.

Pour mieux analyser nos résultats nous avons procédé, dans certain cas, à l'agrandissement de quelques portions des graphes que nous ayons jugé être les plus significatifs et dont l'allure et la plus fréquente. L'analyse de ces portions de graphes va refléter le comportement général de la solution en question. Nous pouvons donc effectuer nos synthèses en se basant sur de tels résultats.

Dans quelques scénarios, la comparaison des résultats de notre solution n'apporte pas d'amélioration, nous allons quand même les inclure dans ce chapitre étant donné que nous examinons aussi les cas extrêmes et les limites de notre approche. D'autant plus que le fait que notre algorithme n'obtient pas des résultats moins bons que ceux obtenus avec PQ et WFQ, est considéré comme un avantage pour CPQ.

Finalement, notons que les graphes correspondants à la longueur des files d'attente ainsi que ceux relatifs aux délais d'attente dans les files sont représentés en bleu pour WFQ, en rouge pour CPQ et en vert pour PQ.

4.2.1 Description du modèle PQ considéré

L'algorithme PQ simulé dans ce travail, consiste à servir la file d'attente prioritaire (VoIP) jusqu'à ce qu'elle soit vide, après quoi la file d'attente de moindre priorité sera servie. La Figure 4.1 illustre un exemple de système de gestion de file d'attente avec PQ.

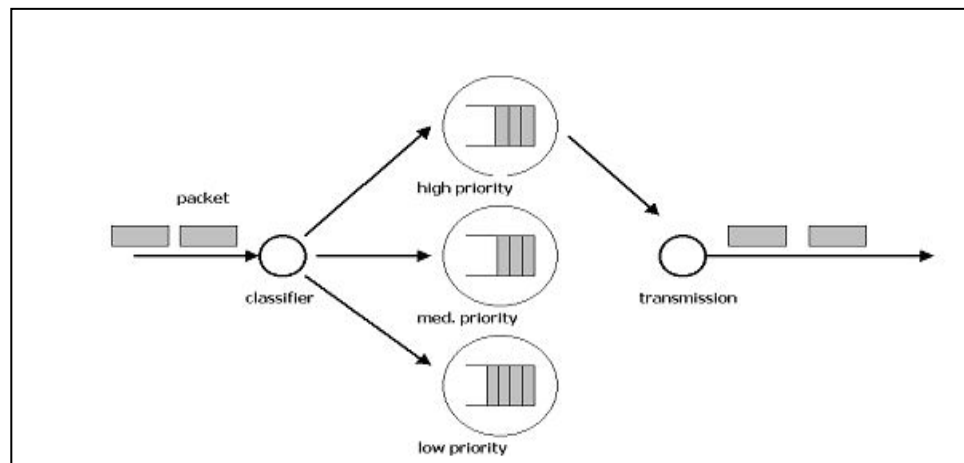


Figure 4.1 Exemple d'un modèle de Priority Queuing.

Tiré de <http://www.coredump.fr.to/altq-fifoq-priq-wfq/>

4.2.2 Description du modèle WFQ considéré

L'algorithme WFQ pris en considération dans ce travail consiste à sélectionner le paquet ayant la valeur Finish Time la plus petite pour le servir. Les deux files d'attente ont été pondérées afin de prioriser la voix, qui représente le trafic rtPS, par rapport au data, qui représente le trafic nrtPS.

Pour le calcul de Finish Time, nous considérons la formule :

$$Ft(i) = Ft_{prev}(i) + Taille_{Paquet(i)} * Wi$$

Avec:

- $Ft(i)$: Finish Time du paquet de trafic i , il correspond au temps virtuel du départ du paquet du système
- $Ft_{prev}(i)$: Correspond au Finish Time du dernier paquet servi du trafic i
- $Taille_Paquet(i)$: Taille du paquet de type i
- Wi : Inverse du poids de la file d'attente du trafic i
- $i = 1$ pour Voix ou 2 pour FTP

4.3 Scénarios de tests Matlab

4.3.1 Scénario1 : scénario de référence

Ce scénario sert d'un test de base pour notre solution, il représente l'état stable du système de files d'attente. Pour cela, nous avons considéré un réseau WiMAX avec 40 connexions FTP de 1Mbits/s et 50 connexions PCM. Un exemple de réseau WiMAX avec système de gestion de file d'attente M/G/1 est présenté dans la Figure 4.2.

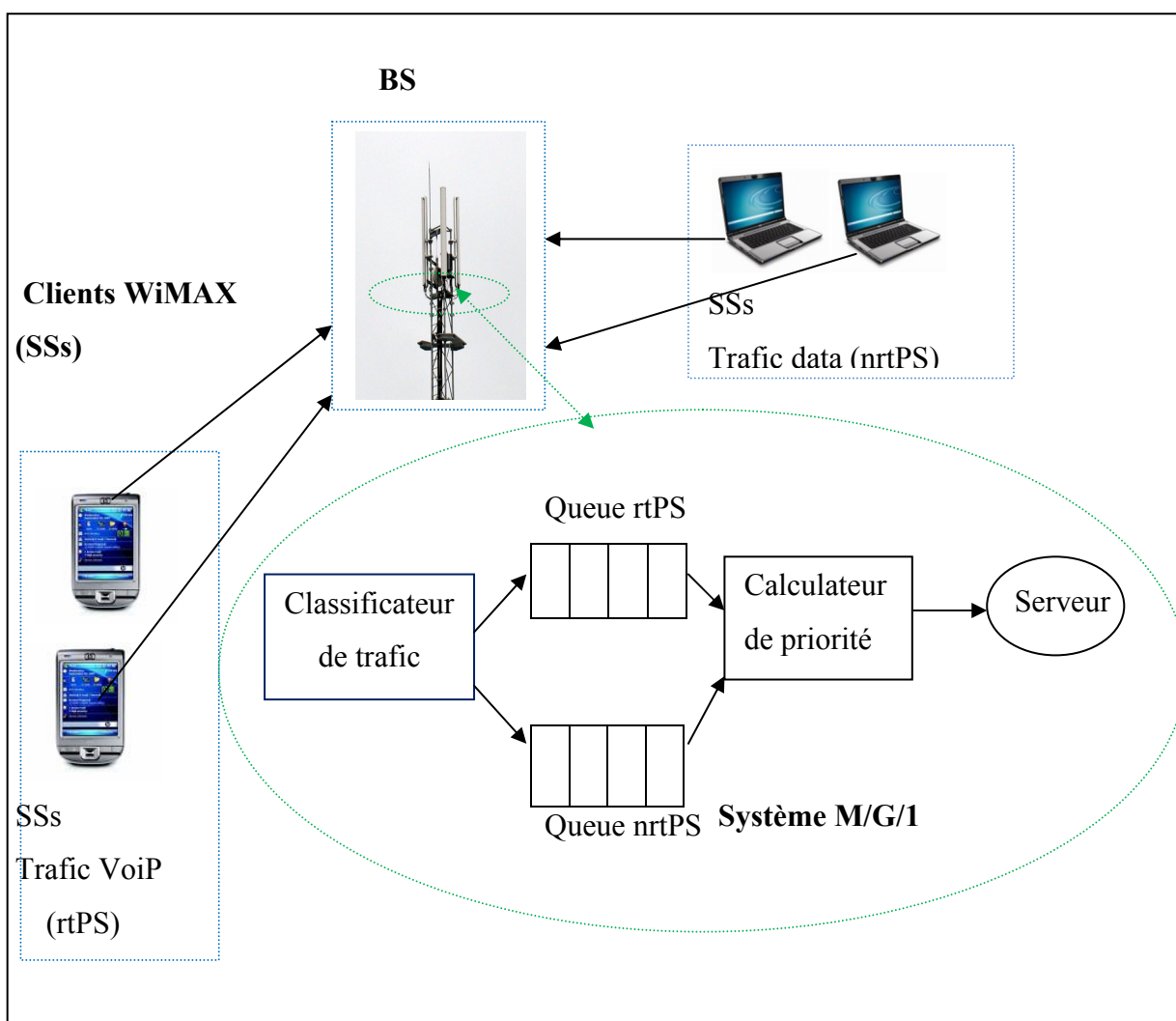


Figure 4.2 Exemple d'un système M/G/1 pour un réseau WiMAX fixe : deux queues.

Le Tableau 4.1 regroupe les paramètres de la simulation utilisés dans le scénario de référence pour les trois approches PQ, CPQ et WFQ, leurs valeurs numériques, ainsi que leurs descriptions.

Tableau 4.1 Paramètres de la simulation, relatifs au scénario de référence

Paramètre	Description	Valeur
Paquets simulés	Taille de l'échantillon	1000 paquets
Inter_arrival1	Temps d'inter-arrivée moyen des paquets VoIP	0.0004 s
Inter_arrival2	Temps d'inter-arrivée moyen des paquets Data	0.0001 s
μ	Temps de service moyen d'un paquet	10 μ s
Délai Max	Délai maximum d'attente dans la file d'attente de la voix	10 ms
Taille Q1	Taille de la file d'attente de la voix	10 paquets
Taille Q2	Taille de la file d'attente du data	12 paquets
Taille Pk1	Taille d'un paquet de type rtPS (voix)	160 bytes
Taille Pk2	Taille d'un paquet de type nrtPS (data)	512 bytes
w1	Poids de la file d'attente VoIP (WFQ)	1.0
w2	Poids de la file d'attente Data (WFQ)	0.2
β_1	Pourcentage de la bande passante allouée pour la VoIP (CPQ)	0.50
β_2	Pourcentage de la bande passante allouée pour le Data (CPQ)	0.008
R1	Taille maximale d'une rafale	2 paquets

Calcul de λ_1 , λ_2 , β_1 , β_2 , w_1 , w_2

Rappelons que λ_1 est le taux moyen d'arrivées de type VoIP, λ_2 est le nombre d'arrivées moyen des paquets de type data. β_1 , β_2 , sont les proportions de la bande passante allouée aux trafics VoIP et data, respectivement dans la solution CPQ, w_1 , w_2 sont les rapports inverses des poids respectifs aux deux files d'attente VoIP et data dans la solution WFQ. Nous disposons de 40 connexions FTP à 1Mb/s chacune, avec des paquets de taille de 512 Octets, et 50 connexions voix PCM avec des paquets de taille de 160 Octets chacun.

Calcul de λ_1

Nous avons 50 connexions à 64 kb/s chacune, donc nous avons (50×160) Octets chaque 20ms, ce qui implique l'arrivée de $(50 \times 160) \times 50$ Octets chaque 1s. Ce qui donne 2500 paquets de 160 Octets de voix par seconde. Alors :

$$\lambda_1 = 2500 \text{ pk/s} \text{ et } 1/\lambda_1 = 0.0004\text{s}$$

Calcul de λ_2

Nous avons 40 connexions FTP à 1Mb/s chaque 1s chacune, ce qui donne $(40 \times 1 \times 1000000)$ bit /s, donc 40/8 MB/s. En terme de paquets nous aurons $40/(8 \times 512)$ Mpk/s. Alors :

$$\lambda_2 = 9765.6 \text{ pk/s} \text{ et } 1/\lambda_2 = 0.0001\text{s}$$

Calcul de w_1 , w_2

Nous avons 40 Mbit/s de FTP ($40 \times 1\text{Mb/s}$) contre 3.2 Mb/s de voix ($50 \times 64 \text{ Kb/s}$), ce qui fait que :

$$\text{Poids de FTP} = (40/3.2) \text{ poids de la VoIP}$$

Ceci implique que :

$$Poids\ de\ FTP = 12.5\ poids\ de\ la\ VoIP$$

Si on pose $Weight_FTP$ comme étant le poids de FTP et $Weight_Voice$ comme étant le poids de la VoIP , alors nous aurons :

$$Weight_Voice = 1$$

$$Weight_FTP = 12.5$$

Afin d'assurer une priorité pour le trafic voix nous allons diviser le poids du trafic FTP par 2.
Ainsi :

$$Weight_Voice = 1$$

$$Weight_FTP = 12.5/2 = 6.25$$

Donc :

$$W1 = 1/Weight_Voice = 1.0$$

$$W2 = 1/Weight_{FTP} = 0.16 \cong 0.2$$

Calcul de β_1 et β_2

Afin d'assurer une priorité du trafic voix par rapport au trafic FTP, nous allons prendre :
 $\beta_1 = 0.50$ et $\beta_2 = 0.008$

Résultats numériques

Le Tableau 4.2 résume les résultats les plus importants pour les trois algorithmes CPQ, WFQ et PQ.

Tableau 4.2 Résultats numériques pour le scénario 1

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	0.052552	0.047557	0.041649
Nombre de paquets moyen dans la queue nrtPS (data)	0.19138	0.19638	0.20229
Délais moyens dans la queue rtPS	4.1163e-006	3.725e-006	3.2623e-006
Délais moyens dans la queue nrtPS	1.4991e-005	1.5382e-005	1.5845e-005
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	0	0	0
Nombre de paquets nrtPS bénéficiant de la courtoisie	60		
% de paquets ayant bénéficiés de la courtoisie	7.5758%		

En observant les résultats numériques, nous constatons que PQ donne de meilleurs résultats pour la voix concernant la taille moyenne de la file rtPS et le délai moyen d'attente dans la file rtPS, par contre, CPQ offre les meilleurs résultats pour le délai d'attente moyen dans la file nrtPS, ainsi que pour sa longueur moyenne. Ceci implique que notre solution apporte une amélioration de la QoS pour le trafic moins prioritaire en termes d'équité de service et de délai d'attente dans la file. Reste à dire que la QoS est garantie dans les trois solutions : PQ,

CPQ et WFQ, car le taux de perte de paquets est nul dans les trois cas et les délais d'attente sont acceptables

Plus de 7% des paquets nrtPS servis par notre solution ont bénéficié de la courtoisie. Ce gain s'interprète par la diminution du délai d'attente dans la file nrtPS.

Résultats graphiques

Étude de la longueur de la file d'attente de la VoIP :

Les résultats graphiques montrent une légère différence dans la longueur de la file d'attente prioritaire dans le cas de CPQ par rapport à PQ et WFQ. La Figure 4.3 et la Figure 4.4 montrent cet écart.

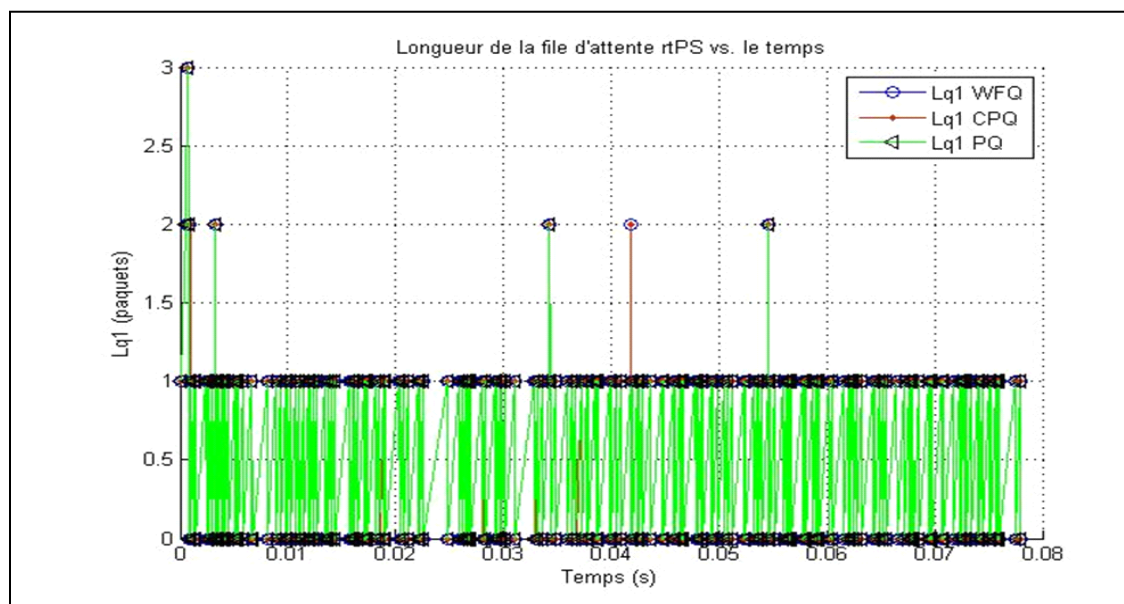


Figure 4.3 Longueur de la file d'attente de la VoIP vs. le temps.



La Figure 4.4 représente le comportement général des trois systèmes WFQ, PQ et CPQ, relatif à la variation de la longueur de la file d'attente de la voix en fonction du temps. Il est clair que le nombre de paquets dans la file d'attente rtPS atteint ses valeurs minimales dans le cas d'application de l'algorithme PQ.

PQ à une gestion des files d'attente qui consiste à servir les paquets prioritaires tant qu'ils sont présents dans leur file. Or, WFQ et CPQ appliquent des méthodes de calcul de priorité afin d'assurer l'Équité de service au sein des différents types de trafic. Par conséquent, certains paquets de classe rtPS peuvent attendre davantage dans leur buffer jusqu'à ce que d'autres paquets nrtPS plus anciens, soient servis. Ceci explique l'augmentation de la taille de la file rtPS dans les cas de WFQ et CPQ.

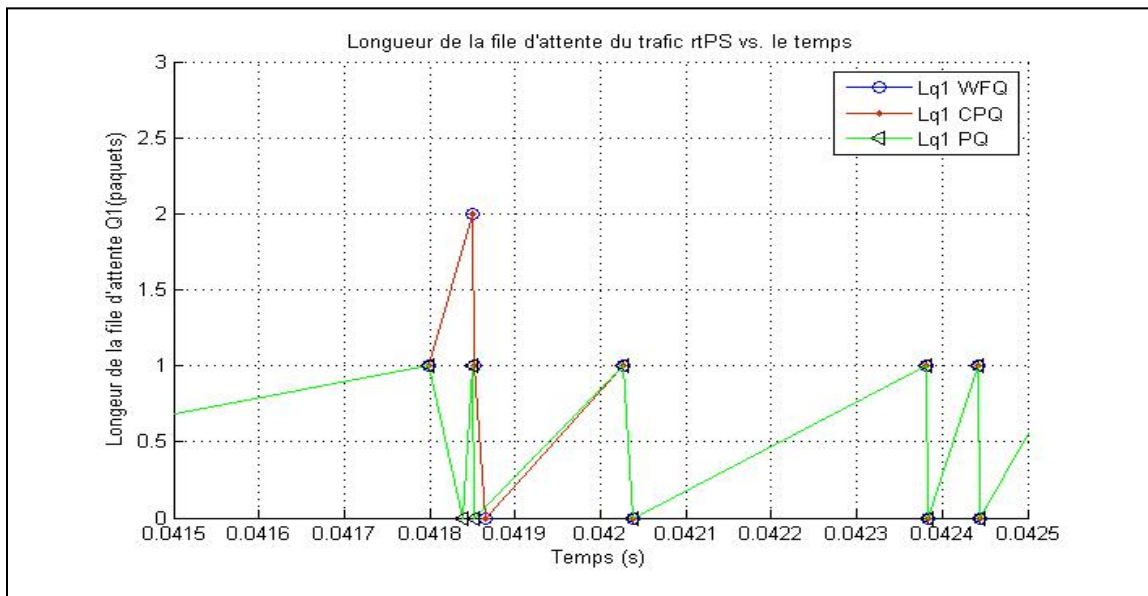


Figure 4.4 Longueur de la file d'attente de la VoIP vs. le temps.

Étude de la longueur de la file d'attente des données :

La longueur de la file d'attente relative aux trois solutions PQ, WFQ et CPQ est pratiquement la même (Figure 1.1). Cependant, si on fait référence aux résultats numériques relatifs à la longueur moyenne de la file nrtPS correspondante à l'application des trois approches, on constatera qu'il y a une légère amélioration apportée par CPQ par rapport aux deux autres algorithmes (cf. Tableau 4.2), représentée par une valeur moyenne de 0.19138 paquet dans la file nrtPS contre 0.19638 pour WFQ et 0.20229 pour PQ relativement à la même file d'attente.

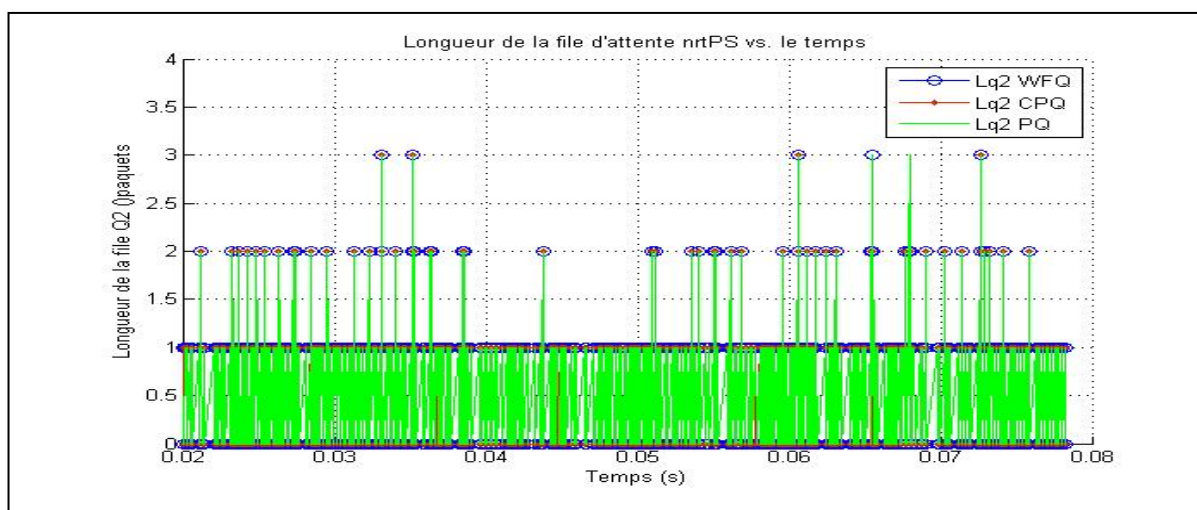


Figure 4.5 Longueur de la file d'attente de données vs. le temps.

Étude du délai d'attente dans la file de la VoIP :

La Figure 4.6 représente les délais d'attente dans la file rtPS pour chaque algorithme. Nous constatons une légère augmentation du délai dans CPQ par rapport aux deux autres solutions. Ce résultat s'explique par l'attente supplémentaire des paquets courtois dans leur file d'attente. PQ donne de meilleurs résultats, car cet algorithme fonctionne de telles manières à discriminer les paquets rtPS, la file d'attente nrtPS ne sera servie seulement si

l'autre file soit vide. WFQ calcule le finish time de chaque paquet d'entête avant de sélectionner la file à desservir. Elle sert la file ayant la valeur de finish time la plus petite. Il est donc évident que le trafic voix sera servi en premier. Or, plus le nombre de paquets d'un trafic donné sera servi plus la valeur de leur finish time augmente, ce qui implique que la priorité de service bascule d'une file à une autre selon la taille et le nombre de paquets déjà servis. Ceci explique la raison pour laquelle WFQ se comporte parfois comme PQ, en favorisant les paquets de la voix, et d'autres fois non.

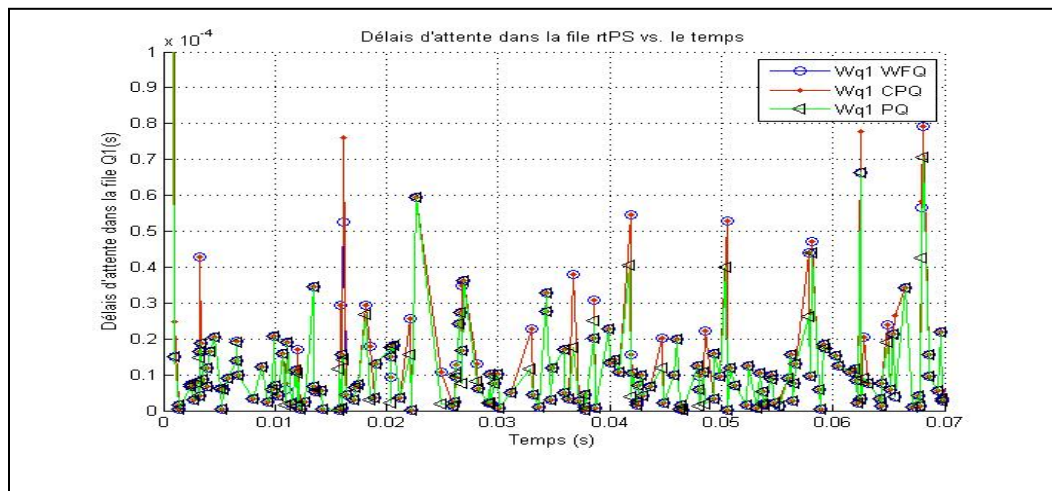


Figure 4.6 Délai d'attente dans la Queue rtPS vs. le temps.

Étude du délai d'attente dans la file des données:

La Figure 4.7 illustre le délai d'attente des paquets nrtPS dans leur file, bien que les écarts soient minces, mais nous voyons que les paquets attendent plus dans le cas de l'approche PQ, et ils attendent moins dans le cas de la solution CPQ. Ce qui est logique vu que dans cette dernière approche, un ensemble de paquets data a tiré profil de la courtoisie pour être servi à la place des paquets de la voix, avant l'arrivée de leur tour, ceci diminue leur temps d'attente dans la file.

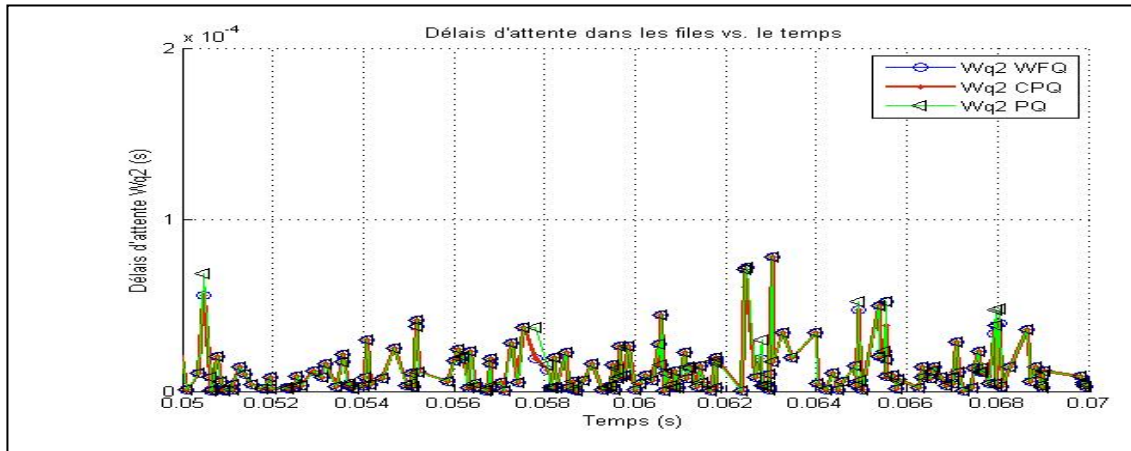


Figure 4.7 Délai d'attente dans la file d'attente nrtPS vs. le temps.

4.3.2 Scénario 2: l'effet de la diminution de λ_1

Ce scénario consiste à tester le cas de la diminution du taux du trafic prioritaire afin d'analyser le comportement des trois approches adaptées, notamment CPQ. La question que nous nous sommes posée est la suivante : quel serait le pourcentage de paquets nrtPS qui seront privilégiés par CPQ dans le cas où le nombre de paquets de la voix n'est pas important? Pour ce faire, nous avons augmenté la valeur d'inter arrivées du trafic rtPS pour atteindre 0.004s. Les autres paramètres demeurent identiques à ceux du scénario de référence (voir Tableau 4.1).

Résultats numériques :

Le Tableau 4.3 résume les résultats les plus importants pour les trois algorithmes CPQ, WFQ et PQ, relatifs au scénario 2.

Nous constatons un faible apport positif pour CPQ, en termes de délais moyens et de nombre de paquets moyens résidents dans la file d'attente de données. Effectivement, nous remarquons que seuls 2.15% des paquets de la file nrtPS ont bénéficié de la courtoisie. Ce résultat s'interprète par le fait que le nombre de paquets moyens arrivant de classe prioritaire

est très petit par rapport à celui des paquets de basse priorité. Le premier est de l'ordre de 250 paquets/sec, donc 2,5 % du trafic global, et le second est de l'ordre de 10000 paquets/sec, qui représente 97.5% de tous les paquets. Il est clair donc que nombreux sont les paquets de données qui seront servis par CPQ comme étant prioritaire.

Tableau 4.3 Résultats numériques pour le scénario 2

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	0.12106	0.11999	0.11907
Nombre de paquets moyen dans la queue nrtPS (data)	1.058	1.059	1.0599
Délais moyens dans la queue rtPS	1.3061e-005	1.2946e-005	1.2847e-005
Délais moyens dans la queue nrtPS	0.00011414	0.00011426	0.00011436
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	82	82	82
Nombre de paquets nrtPS bénéficiant de la courtoisie	21		
% de paquets ayant bénéficiés de la courtoisie	2.1561%		

Résultats graphiques

Étude de la longueur de la file d'attente de la voix :

Les résultats graphiques montrent une convergence des résultats obtenus des trois algorithmes PQ, WFQ et CPQ (Figure 4.8)

Comme nous l'avons mentionné ci-dessus, un faible taux de trafic rtPS ne peut pas offrir l'opportunité au trafic nrtPS d'être favorisé par l'algorithme de courtoisie, pour la simple raison : Beaucoup de paquets nrtPS auront une priorité élevée. Même si les résultats obtenus pour CPQ étant semblables à ceux correspondants aux autres approches, nous les considérons comme encourageants, car cela nous permet de dire que CPQ offre une équité de service pour les trafics de type prioritaire et de type moins prioritaire.

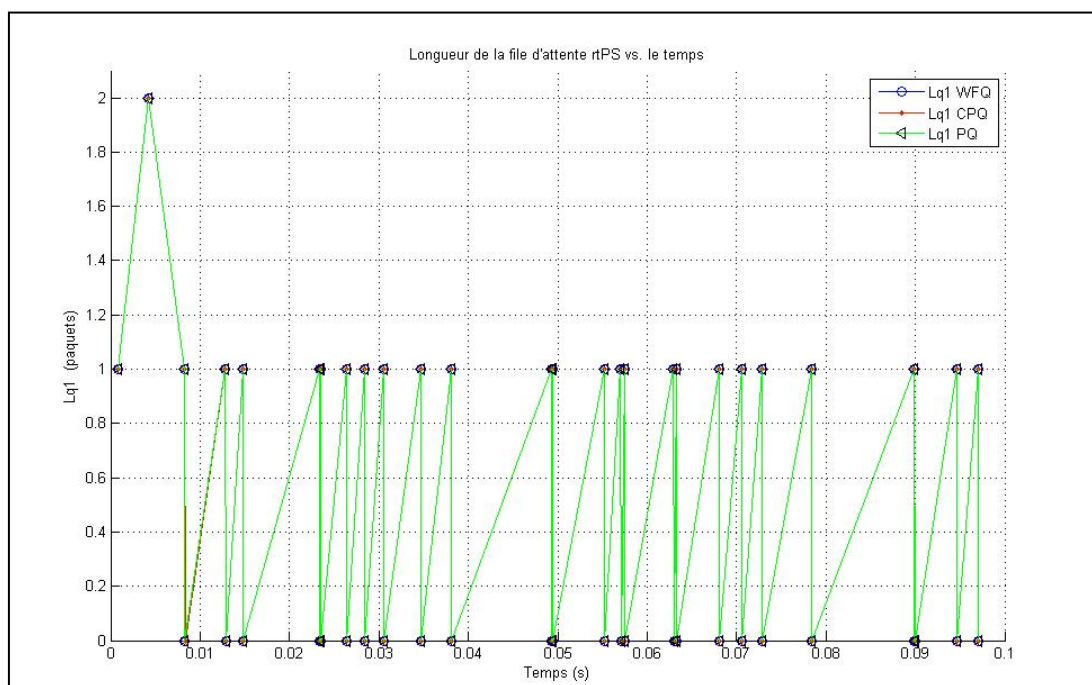


Figure 4.8 Longueur de la file d'attente de la voix pour le scénario 2.

Étude de la longueur de la file d'attente des données :

Comme pour le cas de la file d'attente de la voix, l'utilisation d'une solution ou d'une autre n'améliore quasiment pas l'état du buffer. Étant donné que le réseau WiMAX dispose d'un trafic prioritaire très faible, ce qui mène à dire que le nombre de paquets rtPS courtois est très petit par rapport au nombre de paquets nrtPS qui ont besoin d'être servis.

D'autre part, comme le taux d'arrivée des paquets de la voix est très faible en le comparant avec celui des données, la file d'attente de classe prioritaire va se trouver vide à plusieurs occasions dans les trois solutions. Par conséquent, WFQ et CPQ vont se comporter comme PQ. Ceci est illustré dans la Figure 4.9.

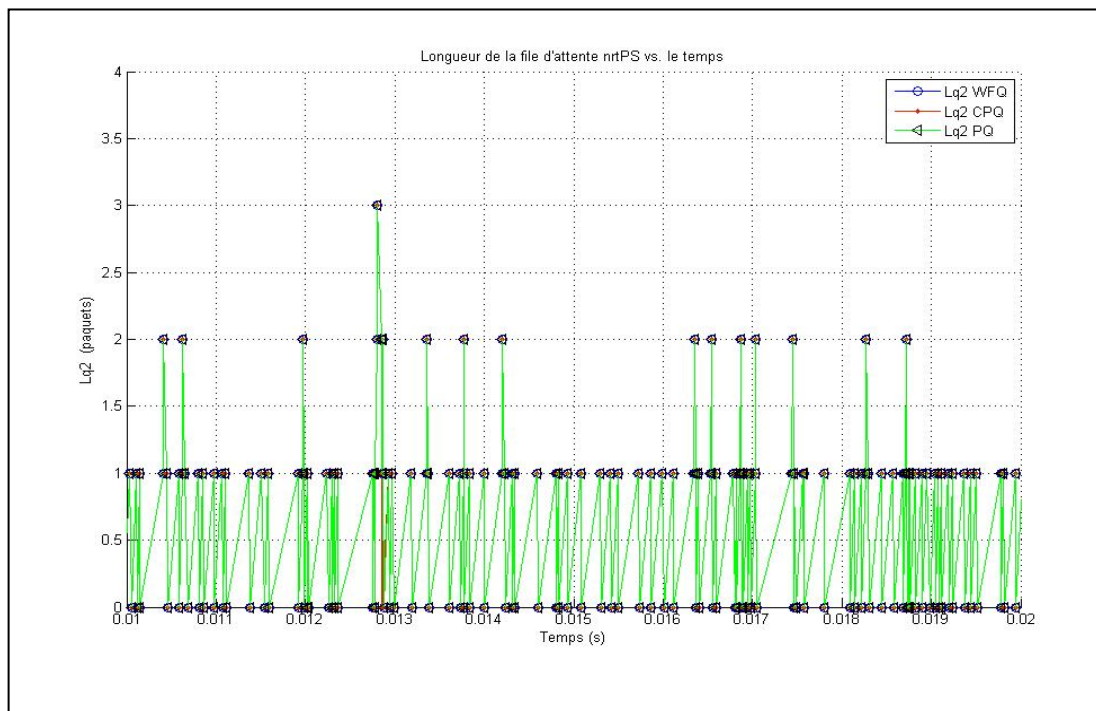


Figure 4.9 Longueur de la file d'attente nrtPS pour le scénario 2.

Étude du délai d'attente dans la file de la voix :

La Figure 4.10 est une portion de la figure globale représentant le délai d'attente dans la file rtPS. Cette figure illustre l'allure assez fréquente correspondante à la figure originale elle montre que, sur certains intervalles de temps, le délai d'attente dans la file rtPS est légèrement plus important dans le cas de CPQ (en rouge) et le moins important dans le cas de PQ (en vert). Ceci s'explique par le temps d'attente supplémentaire ξ_1 des paquets courtois relatif à l'application de l'algorithme de courtoisie.

Le résultat que donne WFQ (en bleu) est plus proche de celui de PQ. Même si nous constatons une différence entre les résultats des trois algorithmes, cette dissimilitude n'est pas considérable, nous concluons donc que les trois solutions convergentes pour donner des résultats pratiquement très proches.

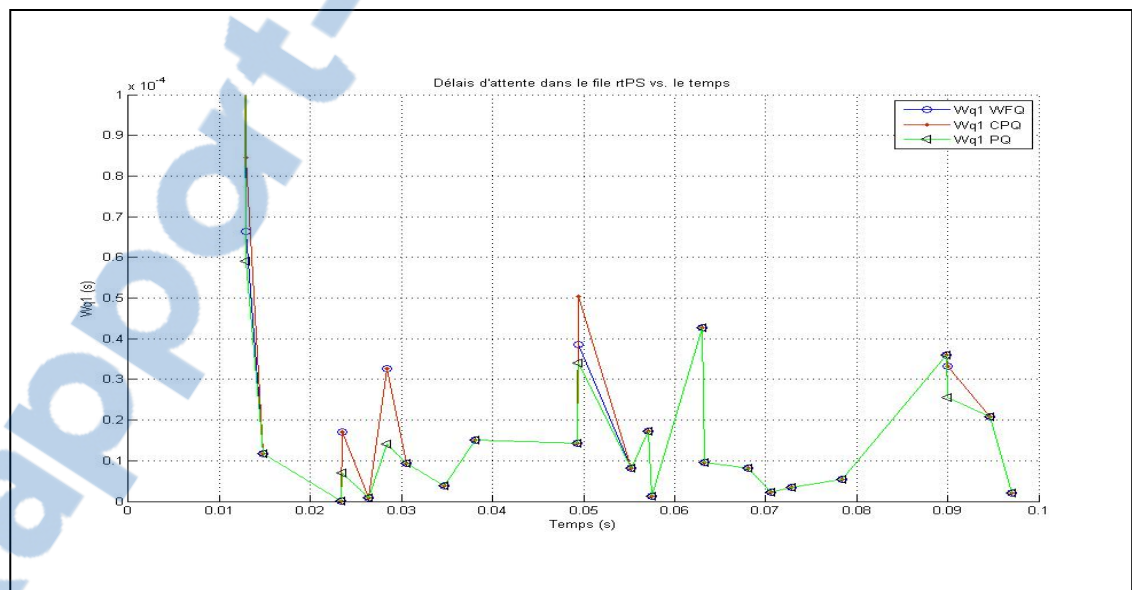


Figure 4.10 Délai d'attente dans la file rtPS pour le scénario 2.

Étude du délai d'attente moyen dans la file de données :

La Figure 4.12 est une partie de la Figure 4.11. Toutes les deux représentent le délai d'attente dans la file nrtPS pour les trois approches PQ, en vert, WFQ, en bleu et CPQ en rouge.

Sauf pour des cas isolés, nous constatons que le comportement de WFQ et CPQ converge vers celui de PQ car la file de la voix est souvent vide, étant donné que le temps d'inter arrivé relatif à la classe rtPS est assez grand par rapport à celui qui correspond à la classe de données. La raison d'avoir ce type de résultat pour notre approche est identique à celle relative aux résultats de Wq1, Lq1 et Lq2. Il s'agit du faible nombre de paquets rtPS existants dans le réseau, ce qui rend le taux de paquets courtois très faible.

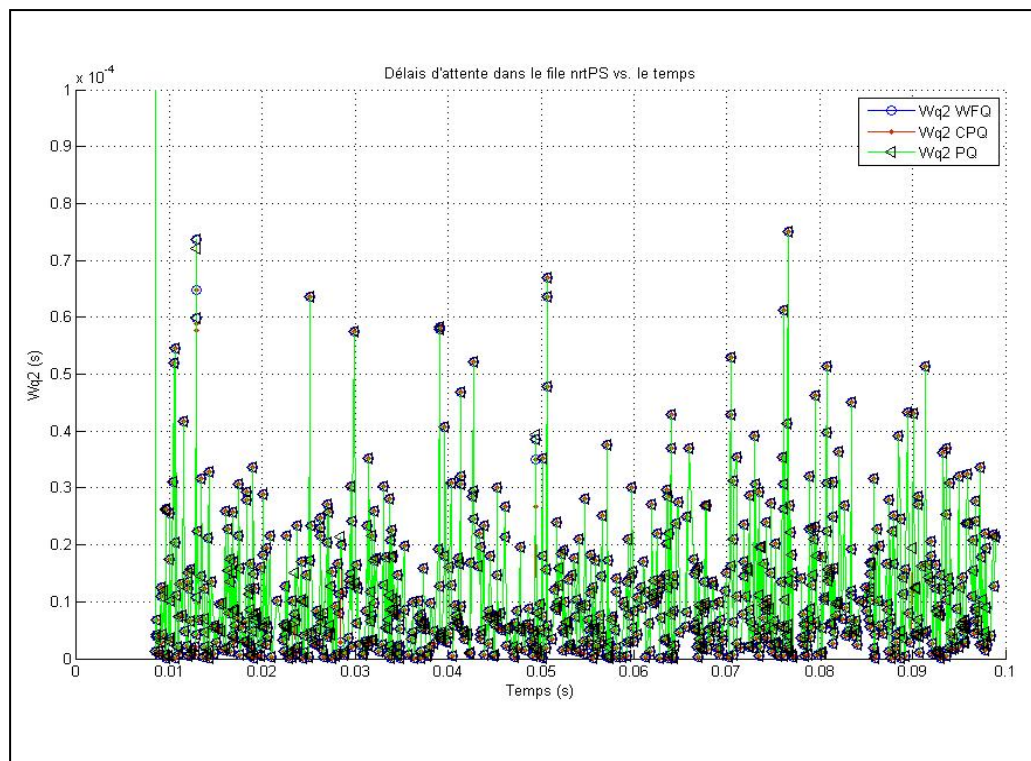
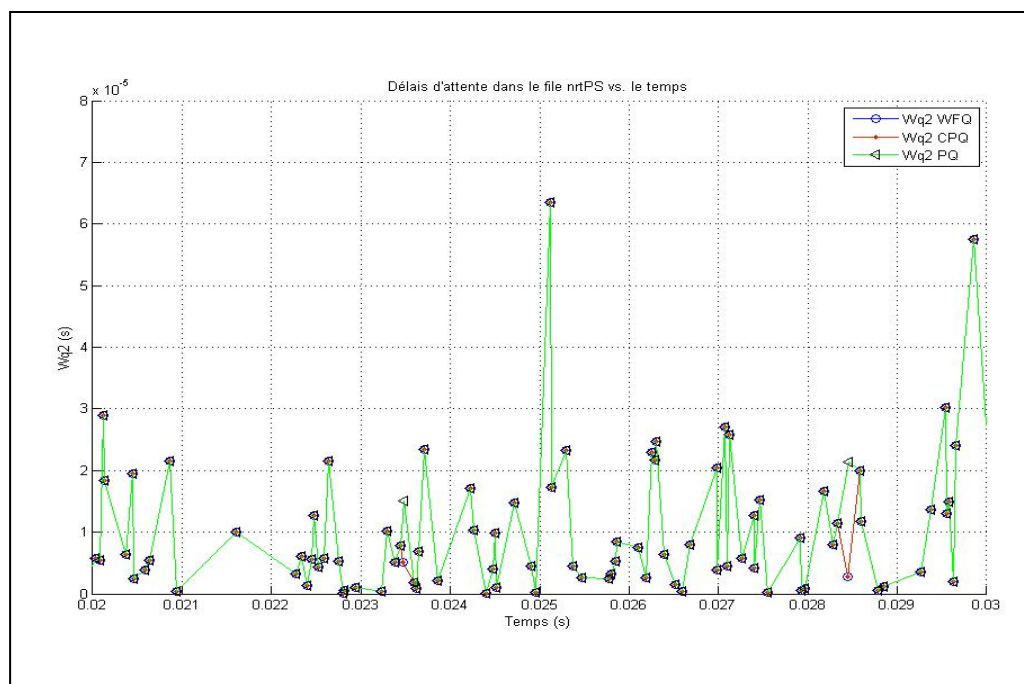


Figure 4.11 Délai d'attente dans la file nrtPS (scénario 2).



**Figure 4.12 Délai d'attente dans la file nrtPS :
Temps = [0.02, 0.03] sec.**

Étude de la perte de paquets nrtPS :

La perte de paquets durant la simulation, que nous avons représentée dans le Tableau 4.3 par la valeur de 82 paquets pour chaque solution, et qui est encadrée dans les graphes de la Figure 4.13, survient juste au début de la simulation, nous n'allons donc pas considérer ce résultat comme fiable étant donné que dans cet intervalle de temps les données sont biaisées.



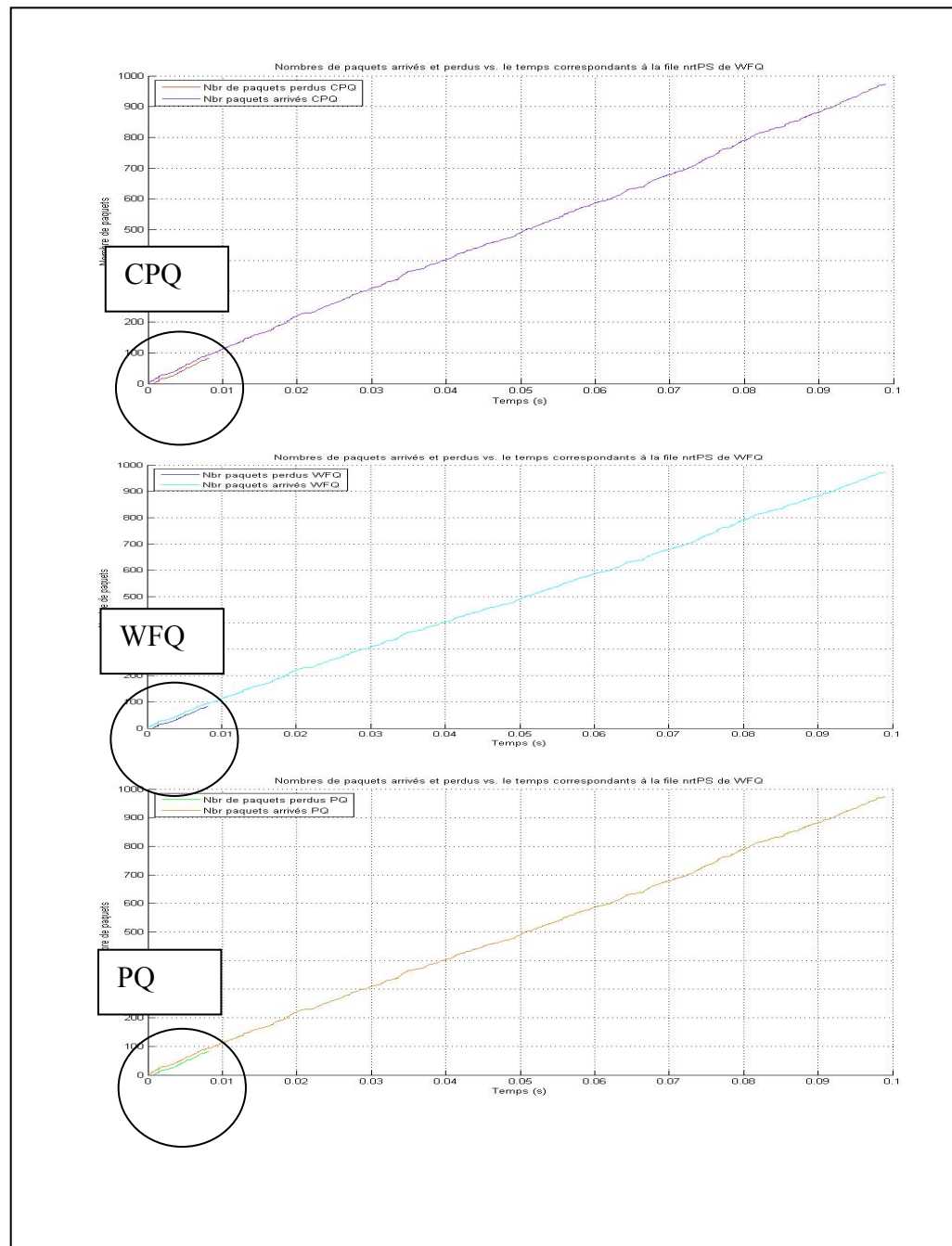


Figure 4.13 Pertes de paquets nrtPS pour le scénario 2.

4.3.3 Scénario 3: Étude de l'effet de l'augmentation de λ_1 , λ_2 et $D_{\max_rtPS}$

Dans ce scénario nous avons élevé le taux d'arrivées des deux types de trafic par rapport au scénario de référence, nous avons donné les valeurs respectives suivantes 0.0002s et 0.00005s pour les inter-arrivées voix et data. Nous voulons donc connaître les comportements des trois algorithmes dans le cas de l'existence d'un trafic important relatif aux deux classes de service. Ajoutant à cela, nous avons augmenté le temps d'attente maximal toléré ($D_{\max_rtPS}$) pour les paquets de la voix. Celui là a été fixé à 20 ms. Nous avons déjà mentionné que le délai acceptable bout en bout pour le trafic voix est de l'ordre de 200 ms. Sachant que ce délai comprend tous les temps d'attente relatifs aux différents sauts séparant la source de la destination, on constatera ainsi que plus le nombre de sauts est élevé, plus le délai maximal d'attente dans une file d'attente doit être réduit. En conséquence, le fait d'augmenter cette valeur revient à considérer un nombre de sauts plus restreint.

Résultats numériques :

Le Tableau 4.4 résume les résultats numériques les plus importants pour les trois algorithmes CPQ, WFQ et PQ. Nous observons que le délai moyen dans la file d'attente $rtPS$ est le plus bas dans le cas de PQ et le plus haut dans le cas de CPQ, il est de même pour la longueur moyenne de la file $rtPS$. Par contre, le CPQ donne un meilleur délai moyen dans la file d'attente $nrtPS$ et PQ offre le délai moyen le plus long dans cette file, pareillement à la longueur moyenne de la file $nrtPS$. WFQ donne des résultats qui se situent entre ceux de l'algorithme PC et ceux de l'algorithme CPQ

Notons que le nombre de paquets moyens du trafic $nrtPS$ ayant bénéficié de notre solution, à savoir CPQ, est de l'ordre de 106. Ce nombre représente 13,26 % du trafic global des données.

Tableau 4.4 Résultats numériques pour le scénario 3

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	0.10141	0.082676	0.063051
Nombre de paquets moyen dans la queue nrtPS (data)	0.35169	0.37042	0.37991
Délais moyens dans la queue rtPS	3.9888e-006	3.5218e-006	2.4849e-006
Délais moyens dans la queue nrtPS	1.3833e-005	1.457e-005	1.4973e-005
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	0	0	2
Nombre de paquets nrtPS bénéficiant de la courtoisie	106		
% de paquets ayant bénéficiés de la courtoisie	13.2666%		

Résultats graphiques

Étude de la longueur de la file d'attente de la voix :

En observant la Figure 4.14 nous constatons que le délai d'attente des paquets de la voix est plus grand dans le cas des algorithmes WFQ et CPQ que dans le cas de PQ. Cela est dû au fait que ce dernier sert tous les paquets de la file d'attente rtPS jusqu'à ce qu'elle devient vide.

Les différences dans la longueur de la file d'attente de la voix ne sont pas très importantes pour les résultats numériques. Ceci interprète la raison pour laquelle les graphes de la Figure 4.14 n'illustrent pas assez visiblement ces différences. Mais ces résultats ne nous empêchent pas de considérer notre solution comme assez bonne que WFQ pour un tel cas de figure.

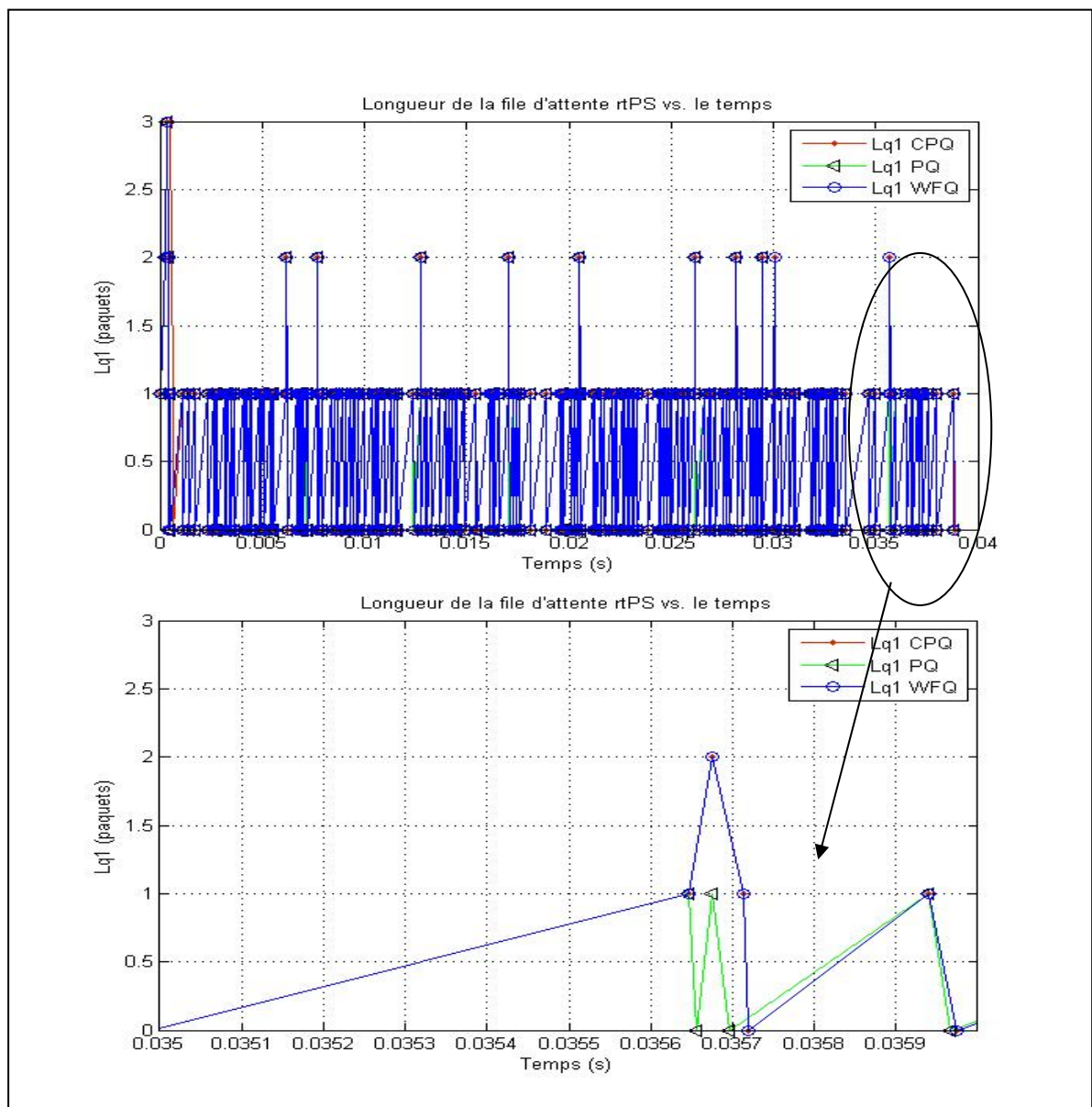


Figure 4.14 Longueur de la file d'attente rtPS pour le scénario 3.

Étude de la longueur de la file d'attente des données :

La Figure 4.15 représente la longueur de la file d'attente de données, elle illustre des graphes tracés pour toute la durée de la simulation, ainsi qu'une portion agrandie qui sert à faciliter l'interprétation des résultats étant donné que cette allure est la plus répétitive.

En regardant la Figure 4.15 nous concluons que notre solution apporte une amélioration en terme de nombre de paquets résidents dans la file d'attente des données. Alors que les deux autres solutions ont un comportement semblable entre elles. Le temps ξ_1 relatif à l'attente supplémentaire tolérée pour les paquets rtPS dans leur file d'attente, a permis de servir quelques paquets additionnels de type nrtPS. Par conséquent, le nombre de paquets de données a diminué.

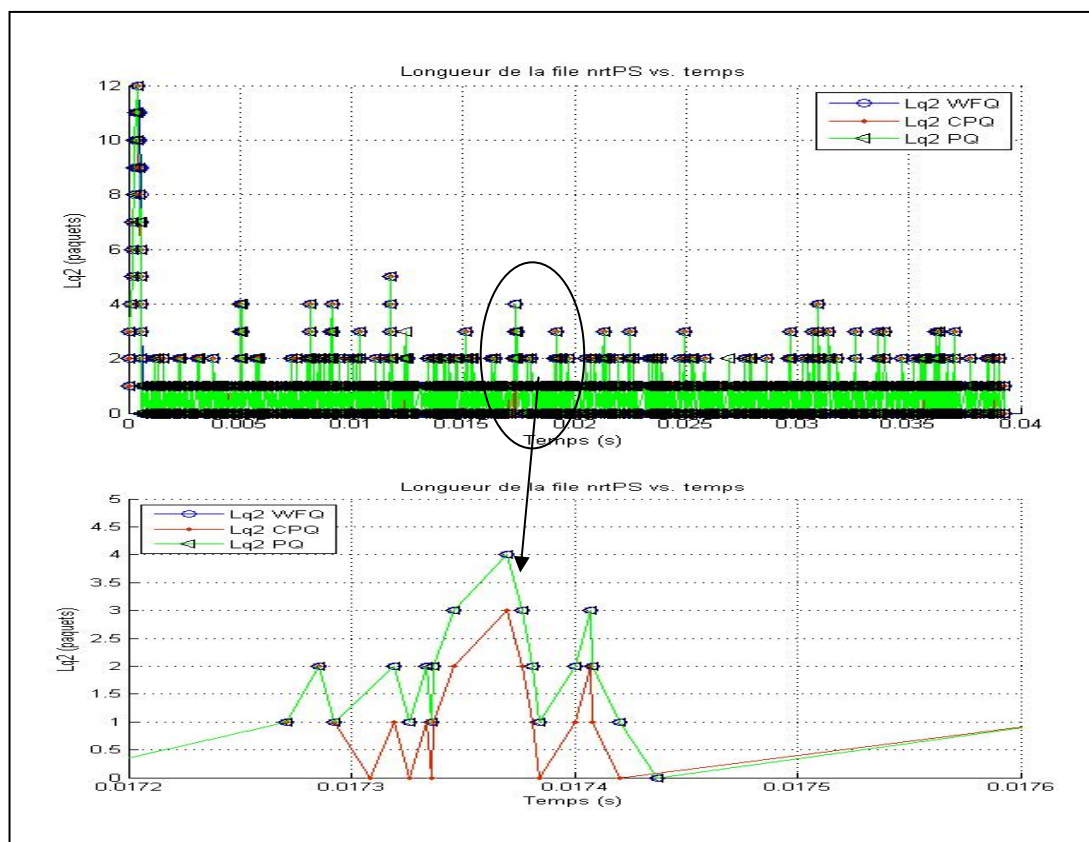


Figure 4.15 Longueurs de la file de données pour le scénario 3.

Étude du délai d'attente dans la file prioritaire (rtPS) :

La Figure 4.16 illustre les résultats du scénario 3 correspondants au délai d'attente dans la file rtPS.

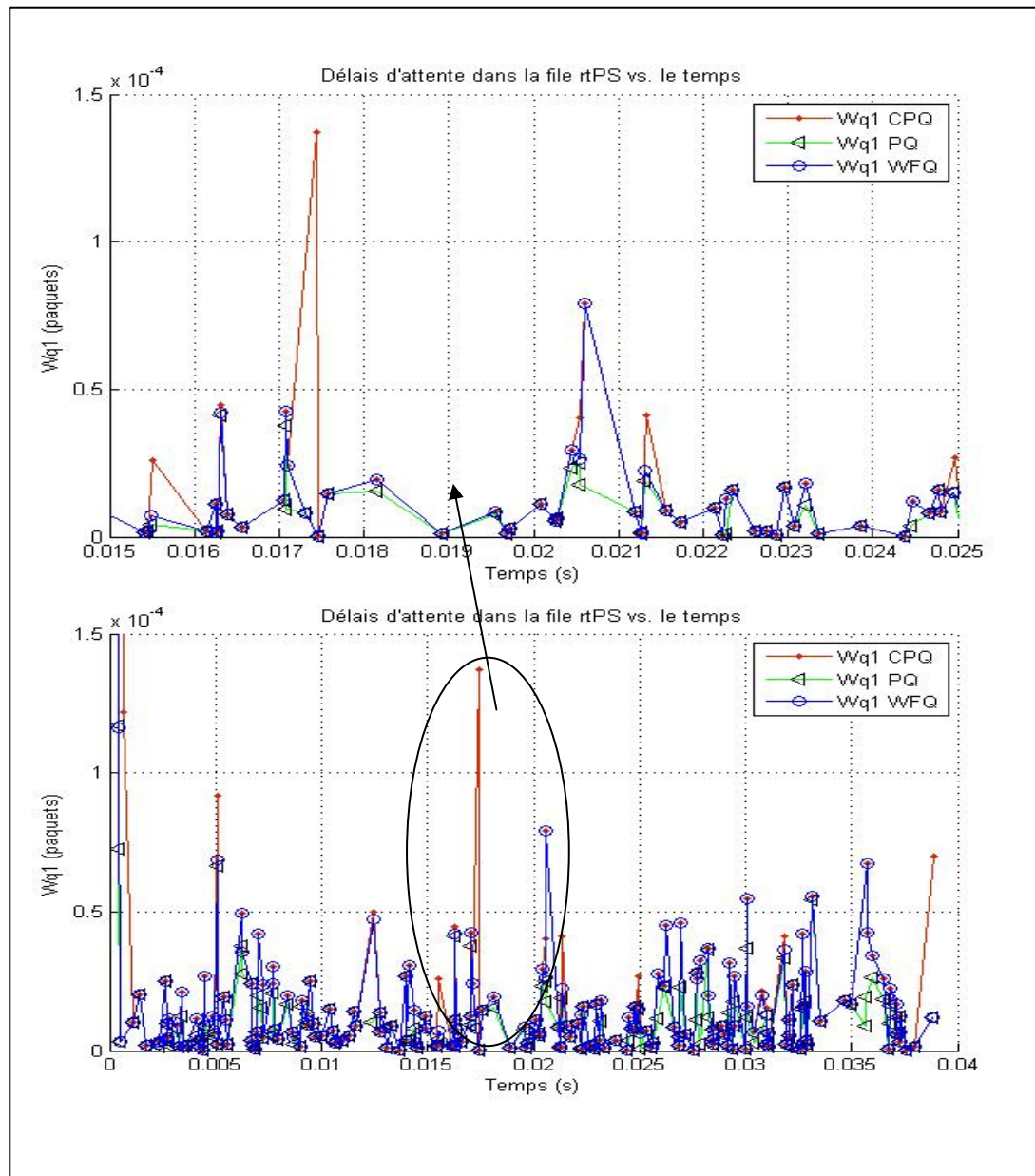


Figure 4.16 Délais d'attente dans la file rtPS pour le scénario 3.

Dans cette figure, nous avons représenté en grand une partie de nos résultats afin de mieux les analyser. En examinant les trois graphes rouge, vert et bleu relatifs au délai d'attente dans la file rtPS obtenus en appliquant CPQ, PQ et WFQ, respectivement, nous constatons que les délais les plus longs sont donnés par CPQ, les meilleurs sont obtenus par PQ, alors que WFQ se comporte parfois comme PQ et parfois comme CPQ. Les meilleurs délais sont obtenus en favorisant la classe rtPS tout le long de la simulation. Il s'agit de la servir jusqu'au dernier paquet qu'y réside avant de passer à la seconde file d'attente. C'est le mécanisme suivi par PQ. Alors que CPQ et WFQ ne suivent pas ce principe, les deux files d'attente sont servies selon la priorité des paquets. Ces deux algorithmes tiennent à assurer une certaine équité de service au sein du système de file d'attente. De plus, CPQ en plus de calculer la priorité des paquets avant de les traiter, applique une politique de courtoisie en cédant la place des paquets rtPS aux paquets nrtPS, ce qui explique la raison pour laquelle WFQ donne de meilleurs résultats par rapport à CPQ relativement au délai d'attente dans la file rtPS.

Il est primordial de noter que malgré que CPQ cause un temps supplémentaire d'attente, mais le délai maximal, que nous observons sur le graphe reste très bon, il est inférieur à 0.15 ms.

Étude du délai d'attente dans la file de données :

Les graphes qui représentent le délai d'attente dans la file nrtPS sont représentés dans la Figure 4.17. L'application de l'algorithme de courtoisie a donné ces fruits, nous voyons bien que le temps d'attente est réduit dans le cas de CPQ, alors qu'il est plus haut pour les deux autres solutions. PQ se montre parfois la plus mauvaise solution pour un tel scénario, car elle donne plus d'importance aux trafics à temps réel. WFQ a un comportement qui oscille entre celui de CPQ et celui de PQ.

Il arrive dans certains cas que les trois algorithmes donnent les mêmes délais, donc suivent la politique de PQ, nous voyons cela sur les graphes qui se superposent. Il s'agit des intervalles de temps où les paquets de la voix deviennent plus prioritaires donc doit être servi en premiers, ou alors dans le cas où la queue rtPS soit vide. Notons que dans une telle situation,

CPQ se trouve incapable d'appliquer la courtoisie pour les deux raisons : l'absence de paquets rtPS ou le strict besoin de ces derniers d'être servi à cause du taux de pertes de ces paquets qui dépasse le seuil toléré.

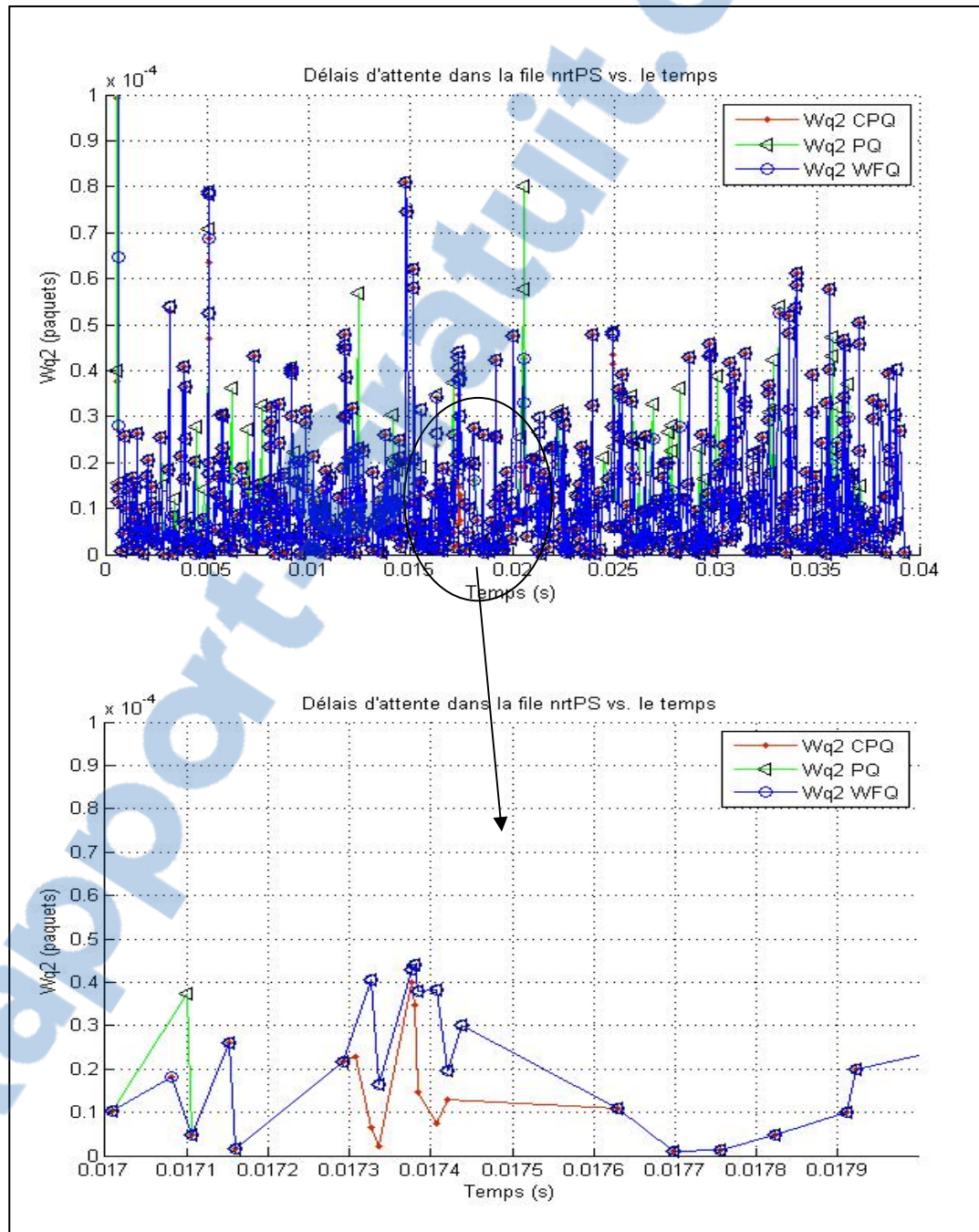


Figure 4.17 Délai d'attente dans la file nrtPS pour le scénario 3.

4.3.4 Scénario 4: Étude de l'effet de l'augmentation de λ_2

Dans ce scénario tous les paramètres demeurent identiques à ceux du scénario 3, nous avons juste augmenté le taux d'arrivée de type nrtPS du double, nous avons donc divisé le taux moyen d'inter arrivés des paquets de données sur deux, ce qui mène à le fixer à 0.000025s.

Notre objectif est d'étudier l'effet de la quantité d'information de type nrtPS transmise, sur le comportement des trois algorithmes. Autrement dit, nous souhaitons connaître les délais d'attente dans les files ainsi que les taux de pertes de paquets relatifs à chaque queue. Nous désirons aussi savoir si l'augmentation de la quantité des données par rapport à la voix dans un réseau WiMAX recommande l'utilisation de notre approche ou non en comparant ses résultats avec ceux de PQ et WFQ.

Résultats numériques :

Les résultats numériques relatifs au scénario 4 (voir Tableau 4.5) sont encourageants pour notre solution. Nous allons résumer les résultats les plus importants dans ce qui suit.

L'application de CPQ permet de réduire la longueur moyenne de la file d'attente nrtPS par rapport aux deux autres algorithmes. Bien que le nombre de paquets rtPS est plus important dans notre solution par rapport au scénario de référence, mais cela ne va pas affecter la QoS de cette classe étant donné que le délai d'attente est bon, même s'il est un peu élevé par rapport aux délais moyens donnés par PQ et WFQ.

L'apport de notre solution consiste à la diminution du délai d'attente dans la file nrtPS sans causer des pertes de paquets dans la file rtPS. De plus, nous avons réduit le taux de perte de paquets par rapport à PQ de deux paquets. Notons que le nombre de paquets nrtPS qui est servi à la place des paquets courtois représente 22.24%, un résultat qui est loin d'être négligeable.

Tableau 4.5 Résultats numériques correspondants au scénario 4

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	0.24336	0.14911	0.10157
Nombre de paquets moyen dans la queue nrtPS (data)	0.86018	0.95589	0.98725
Délais moyens dans la queue rtPS	5.4396e-006	3.3329e-006	2.275e-006
Délais moyens dans la queue nrtPS	1.9219e-005	2.1358e-005	2.2104e-005
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	12	12	14
Nombre de paquets nrtPS bénéficiant de la courtoisie	194		
% de paquets ayant bénéficiés de la courtoisie	22.24%		

Résultats graphiques

Étude de la longueur de la file d'attente rtPS

En examinant la Figure 4.18, nous constatons que la longueur de la file d'attente rtPS augmente dans le cas d'application de CPQ. Alors que WFQ montre des résultats qui convergent. Le comportement de courtoisie des paquets de la voix invoqué par notre solution entraîne une attente supplémentaire des paquets, ce qui va accroître la taille de la queue. Ce phénomène étant absent dans les deux autres solutions, nous voyons bien que le nombre de paquets résidants dans la file rtPS est souvent inférieur à celui obtenu par CPQ.

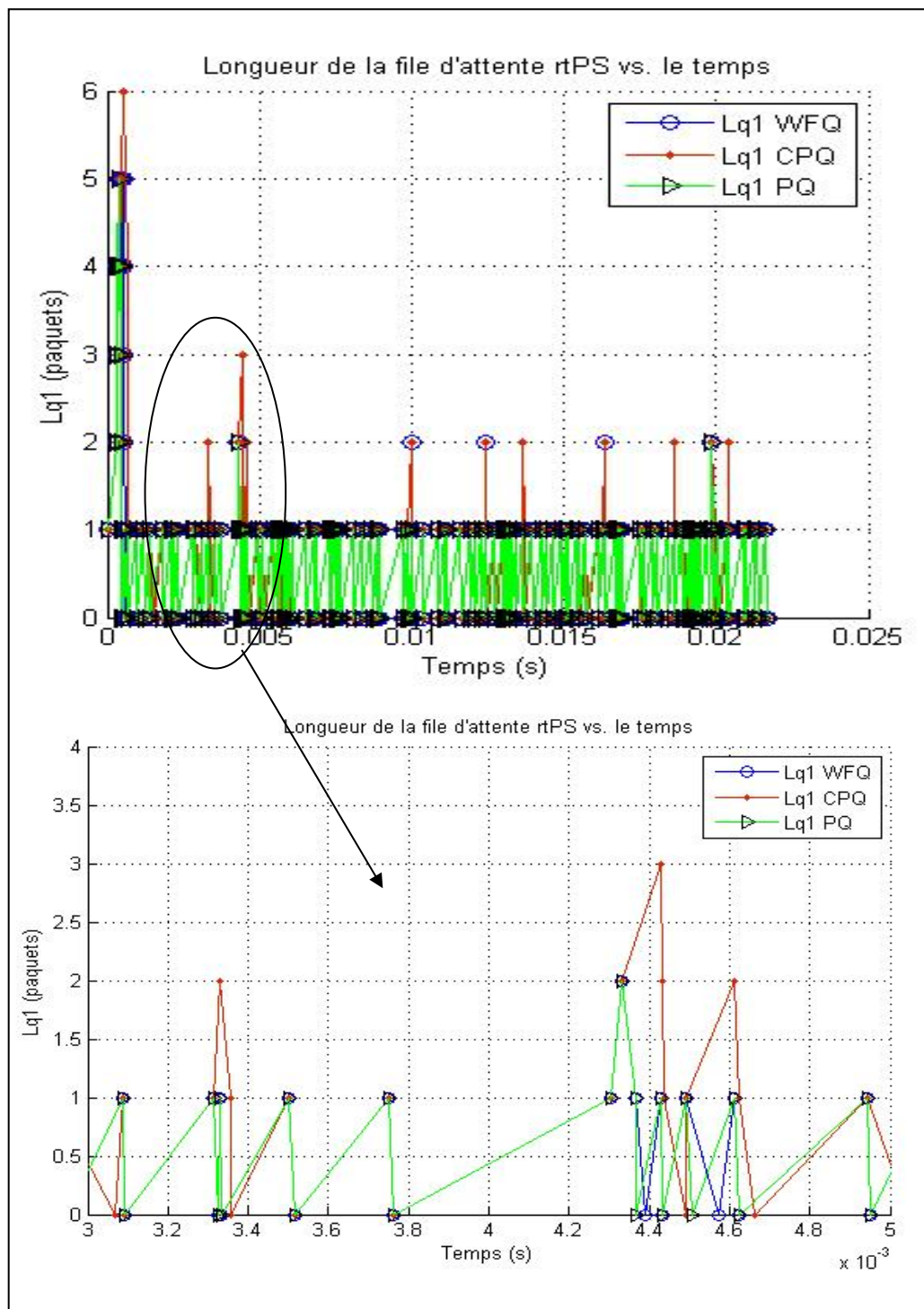


Figure 4.18 Longueur de la file d'attente rtPS (scénario 4).

Étude de la longueur de la file d'attente nrtPS

Contrairement aux résultats de la Figure 4.18 relatifs à la longueur de la file d'attente de la voix, ceux qui correspondent au nombre de paquets dans la queue des données représentent une amélioration de l'état de la file dans le cas de l'application de CPQ (Figure 4.19).

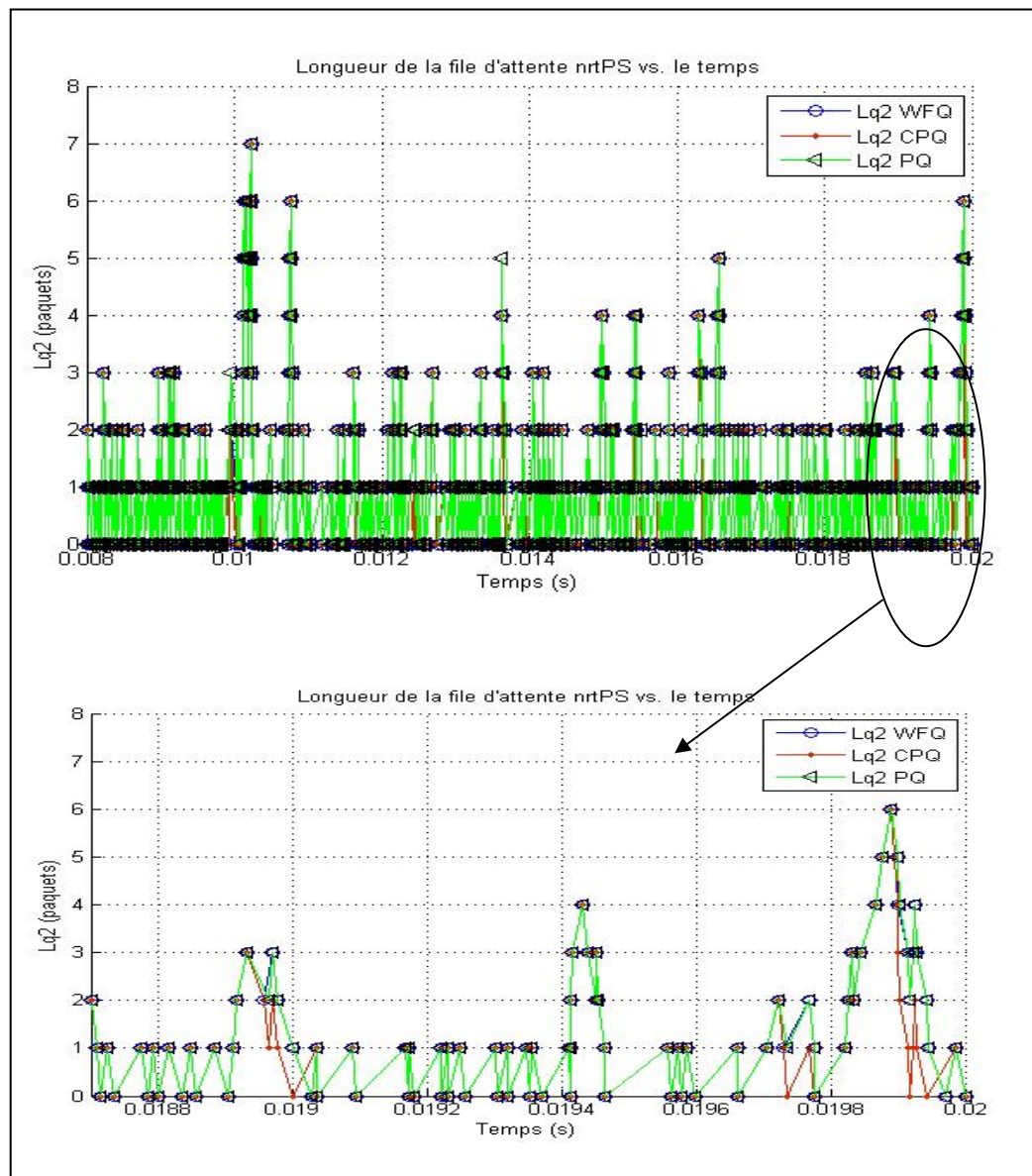


Figure 4.19 Longueur de la file d'attente nrtPS pour le scénario 4.

Le comportement général des trois solutions pour ce scénario montre une certaine amélioration apportée par CPQ (le graphe en rouge). L'application de notre algorithme donne l'opportunité aux paquets nrtPS d'être servis avant leur tour. Les paquets rtPS se montre donc courtois et ce comportement apporte une optimisation de la QoS pour le trafic moins prioritaire.

Étude du délai d'attente dans la file d'attente de la voix

L'application des trois algorithmes PQ, WFQ, et CPQ garantit un bon délai pour la classe de trafic rtPS (Figure 4.20). Particulièrement, nous observons un écart entre les résultats donnés par les trois solutions. Le graphe en vert concerne l'algorithme PQ. Il représente le meilleur délai pour le trafic voix, suivi du graphe bleu correspondant à WFQ. Finalement, CPQ (graphe en rouge) offre un délai plus long en le comparant avec les deux autres, que nous considérons comme une conséquence de la courtoisie d'un ensemble de paquets de la voix, en cédant leur tour de service aux paquets de données.

Dans ce scénario, nous remarquons que l'écart entre les résultats est plus visible par rapport aux scénarios précédents. Ceci est la conséquence de l'augmentation du taux de trafic des deux classes de service. Ce qui augmente le pourcentage des paquets aptes à être courtois, et rend le nombre de solliciteurs de courtoisie dans le système plus important.

Les délais étant bons pour la voix expliquent l'absence de perte de paquets relative à la classe rtPS, un résultat que nous avons déjà constaté lors de l'examen des résultats numériques (voir Tableau 4.5).

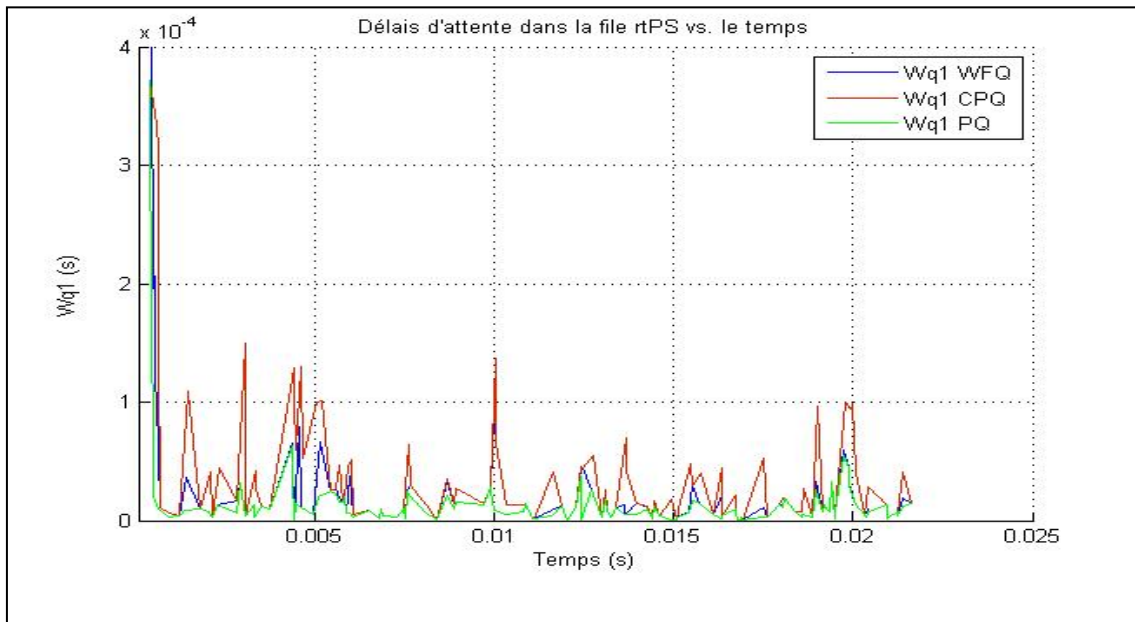


Figure 4.20 Délai d'attente des paquets rtPS (scénario 4).

Étude du délai d'attente dans la file nrtPS

Tenant compte des constations relatives à l'étude du délai d'attente des paquets rtPS, nous attendons à ce que notre solution apporte une performance pour le délai d'attente nrtPS. Évidemment, en analysant la Figure 4.21 et la Figure 4.22, nous concluons que l'algorithme de courtoisie diminue le temps d'attente dans la file des données.

La Figure 4.22 représente une portion de graphe parmi celles qui se répètent le plus souvent pour les trois approches. Notre solution contribue à l'amélioration de la QoS pour le trafic de basse priorité en terme de délai d'attente (voir graphe en rouge), alors que les deux autres solutions offrent un niveau moins bon de QoS pour ce même facteur.

Notons que le temps d'attente supplémentaire causé aux paquets rtPS influe positivement sur le service des paquets nrtPS. Ce temps-là sera alloué à la file d'attente nrtPS, ce qui explique la diminution que connaît le temps d'attente dans la file d'attente des données.

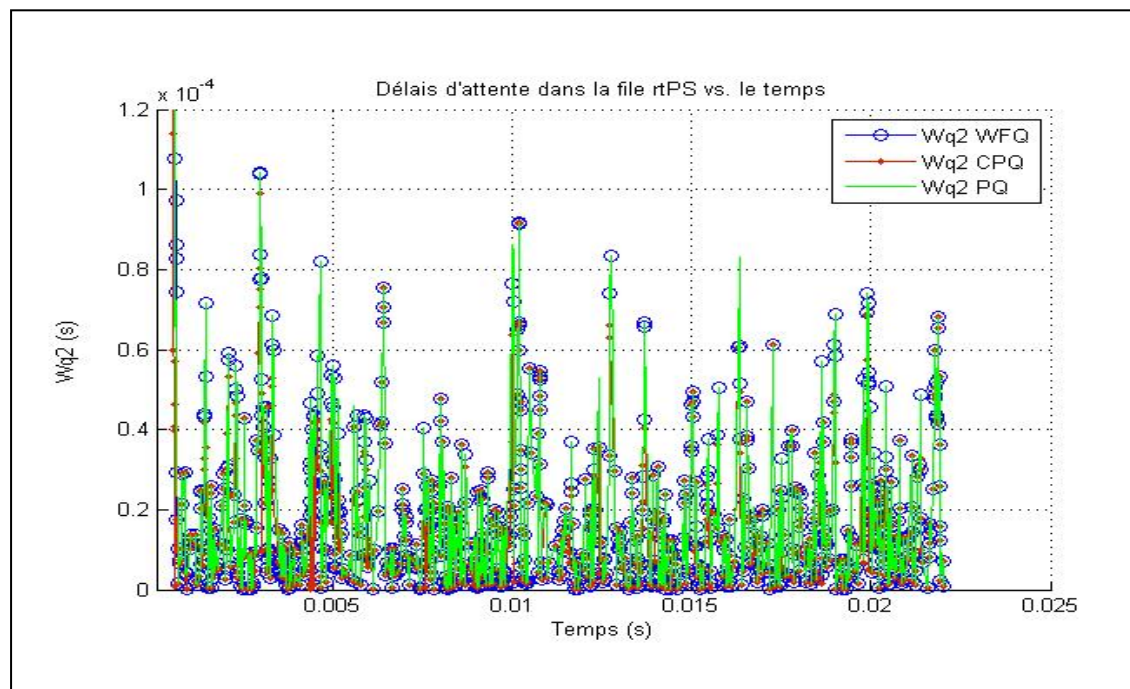


Figure 4.21 Délai d'attente dans la file nrtPS (scénario 4).

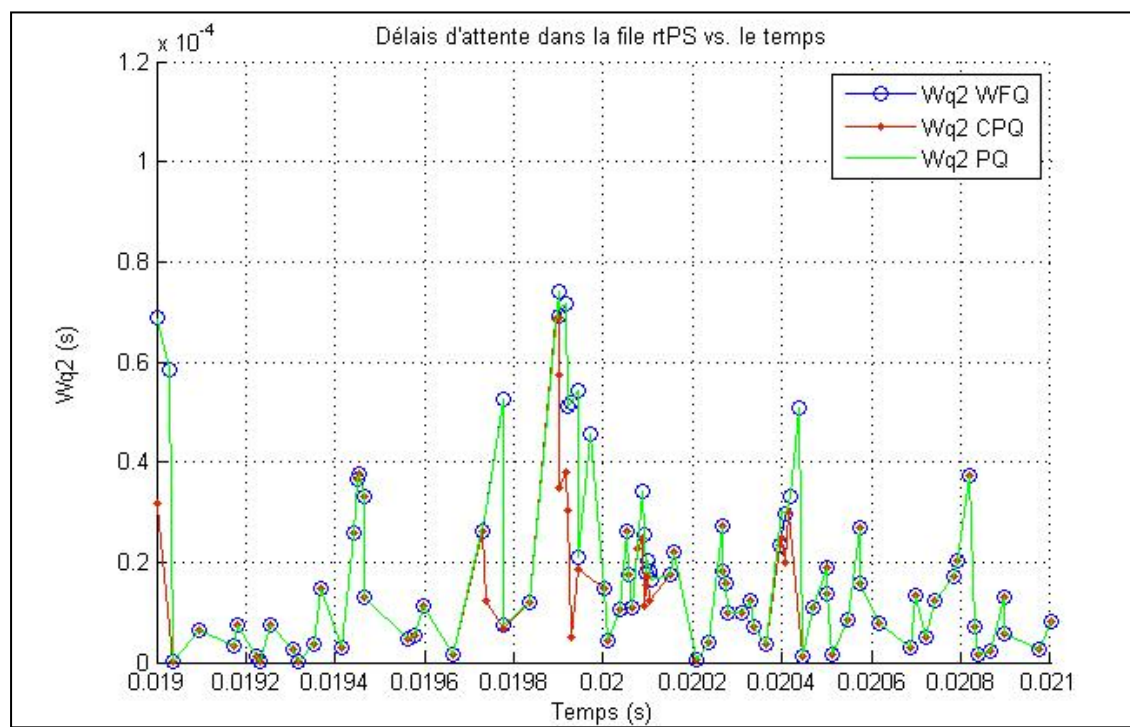


Figure 4.22 Délai d'attente dans la file nrtPS : Figure partielle.

Étude de la perte de paquet :

La Figure 4.23 illustre la perte de paquets de données dans le système de files d'attente conjointement avec les arrivées et les départs nrtPS pour les trois solutions PQ, WFQ et CPQ. Les résultats numériques révèlent une quantité de perte de paquets égale à 12 paquets pour CPQ et WFQ et équivalente à 14 pour PQ. Or, la Figure 4.23 montre que cette perte d'information arrive juste au début de la simulation. Nous n'allons donc pas considérer ce résultat comme fiable.

En conclusion, nous allons considérer que dans ce scénario les trois solutions ne causent pas de perte de paquets pour aucun type de trafic. Les résultats numériques pour cet élément ne seront pas pris en considération.

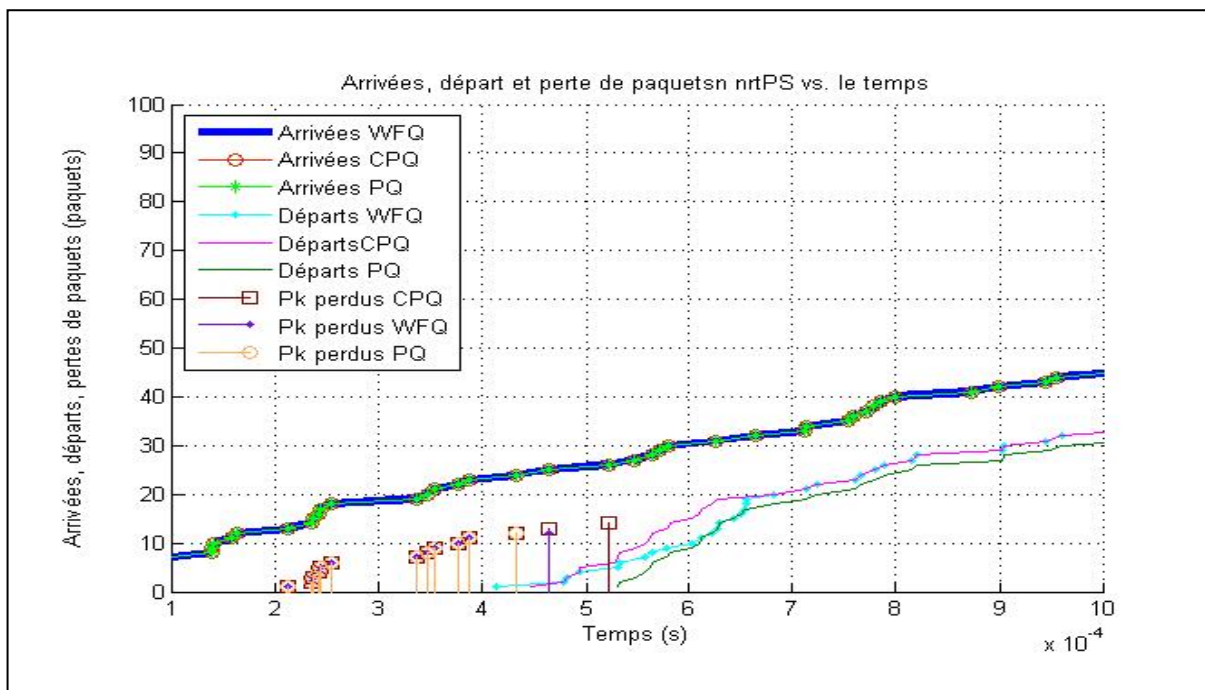


Figure 4.23 Pertes de paquets nrtPS (scénario 4).

4.3.5 Scénario 5: Étude de l'effet de λ_1 et λ_2 et de la taille de l'échantillon

Dans les scénarios précédents, nous avons exploré l'effet du taux du trafic rtPS, du taux du trafic nrtPS, ainsi que l'effet de l'augmentation des quantités de l'information relative aux données et à la voix, sur notre système de files d'attente en appliquant les trois algorithmes CPQ, PQ et WFQ. Dans le présent scénario nous allons augmenter davantage le nombre des paquets des deux types de trafic de telle sorte à fixer le taux d'inter arrivée moyen rtPS à 0.0001s et celui de nrtPS à 0.00001s. De plus, nous allons augmenter la taille de notre échantillon en simulons 10000 paquets au lieu de 1000 tel que nous l'avons déjà fait dans le scénario de référence.

Le but de ce scénario est d'étendre notre intervalle de confiance en explorant le comportement des trois algorithmes considérés dans ce travail durant une durée plus longue de transmission d'informations et en la présence d'une quantité assez importante de trafic rtPS et de trafic nrtPS. Nous souhaitons donc savoir laquelle des trois solutions résiste dans une telle situation.

Résultats numériques :

Le Tableau 4.6 résume les résultats numériques les plus importants pour les trois algorithmes CPQ, WFQ et PQ. Les données de ce tableau illustrent que PQ est optimal pour le délai d'attente moyen dans la file rtPS, ainsi que pour sa longueur moyenne, alors que CPQ est performant pour le délai d'attente moyen dans la file nrtPS, ainsi que pour sa longueur moyenne. WFQ donne des résultats qui se situent au milieu de ceux obtenus des deux précédents.

Nous remarquons que le délai d'attente moyenne dans la file d'attente de la voix donné par CPQ est assez grand par rapport à ceux donnés par les deux autres approches. Ceci s'explique du fait qu'un grand nombre de paquets rtPS cèdent leur tour de service pour l'autre classe, ce qui augmente leur délai d'attente dans leur file, mais sans atteindre le délai

maximal toléré pour la classe rtPS ce qui s'interprète par le nombre de perte de paquets de la voix qui est nul dans ce scénario. Le nombre ayant profité de la courtoisie est de l'ordre de 82.49% de l'ensemble de paquets nrtPS arrivés.

Tableau 4.6 Résultats numériques pour le scénario 5

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	4.964	0.13289	0.10695
Nombre de paquets moyen dans la queue nrtPS (data)	6.2788	6.879	6.8855
Délais moyens dans la queue rtPS (sec)	5.006e-005	1.3857e-006	1.1153e-006
Délais moyens dans la queue nrtPS (sec)	6.3893e-005	7.1723e-005	7.1798e-005
Nombre de paquets rtPS arrivés	900	900	900
Nombre de paquets nrtPS arrivés	9102	9100	9100
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	890	1115	1116
Nombre de paquets nrtPS bénéficiant de la courtoisie	7509		
% de paquets ayant bénéficié de la courtoisie	82.49%		

Résultats graphiques

Étude de la longueur de la file d'attente rtPS

La longueur de la file d'attente rtPS est illustrée dans la Figure 4.24. Le graphe en rouge représente le résultat obtenu de CPQ, celui en vert est relatif à PQ, quant au graphe bleu, il correspond au résultat de WFQ.

PQ est la solution qui garde le nombre de paquets le plus petit pour la file d'attente de classe prioritaire. WFQ vient en deuxième rond concernant ce même résultat alors que CPQ donne le nombre moyen le plus élevé concernant le nombre de paquets de la voix dans leur file d'attente. Ce comportement est dû au fait que PQ favorise la voix et sert ses paquets jusqu'à ce que la file correspondante sera vide. WFQ sert les paquets selon la valeur de leur finish time, étant donné que les paquets de la voix sont plus petits que ceux de FTP (les données) alors il est clair que de nombreux paquets rtPS seront servis en premier. Ce qui explique la raison pour laquelle la taille de la file d'attente de la voix pour cette solution coïncide avec l'approche PQ. Toutefois, pour le cas de CPQ, en considérant les résultats numériques, nous déduisons que plus de 80% des paquets de FTP sont servis à la place des paquets de la voix, sachant que tous les paquets arrivant à la file d'attente de la classe prioritaire vont attendre jusqu'à ce que le temps supplémentaire alloué pour la classe moins prioritaire soit écoulé. CPQ cause donc une augmentation considérable de la taille de la file d'attente rtPS.

Un écart important se montre dans le délai donné par CPQ par rapport à ceux obtenus par les deux autres solutions pour la file d'attente rtPS. La raison est qu'un nombre très important des paquets de la voix attendent un temps supplémentaire considérable dans leur file d'attente. Notons que cette attente additionnelle n'influe pas négativement sur la QoS de la voix car le taux de perte de paquets est nul pour cette classe de service.

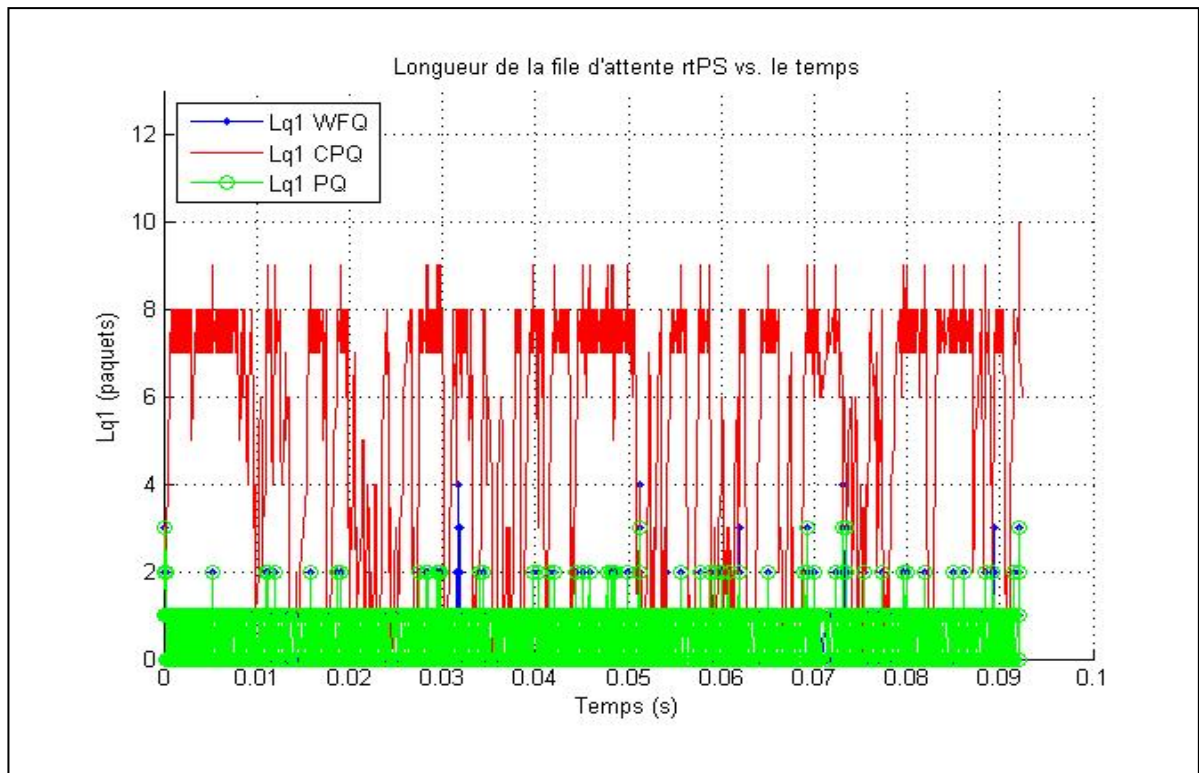


Figure 4.24 Longueur de la file d'attente rtPS (scénario 5).

Étude de la file d'attente nrtPS

Les résultats donnés par les trois algorithmes de gestion de files d'attente PQ, CPQ et WFQ relativement à la taille du buffer nrtPS sont illustrés dans la Figure 4.25 et la Figure 4.26. Cette dernière figure symbolise l'allure générale des résultats, elle représente un agrandissement d'une portion de la Figure 4.25.

L'examen de la Figure 4.26 implique que la performance de la taille de la file de données est obtenue en appliquant notre approche, à savoir CPQ. Cette diminution est dû au nombre de paquets de la classe de moindre priorité présents dans leur queue, car de nombreux paquets de cette file ont eu l'opportunité d'être servis à la place des paquets de la classe de haute priorité.

WFQ donne parfois des résultats meilleurs que ceux correspondants à PQ, car cette solution sert les paquets en appliquant une certaine équité, contrairement à PQ qui ne traite les paquets nrtPS que si la file d'attente rtPS est vide.

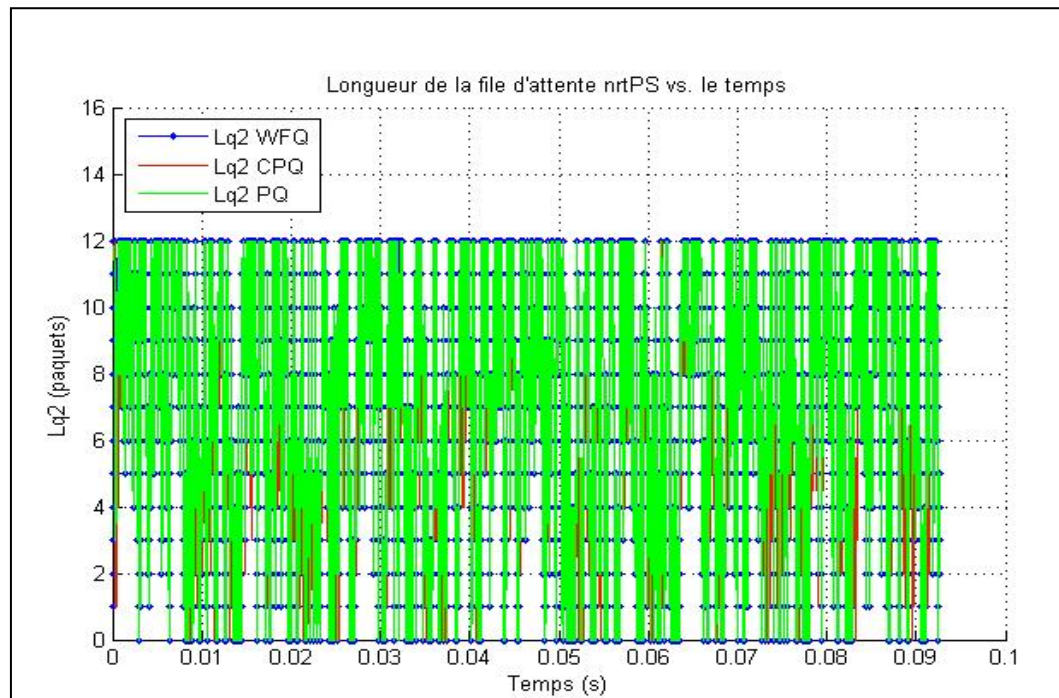


Figure 4.25 longueur de la file d'attente nrtPS (scénario 5).

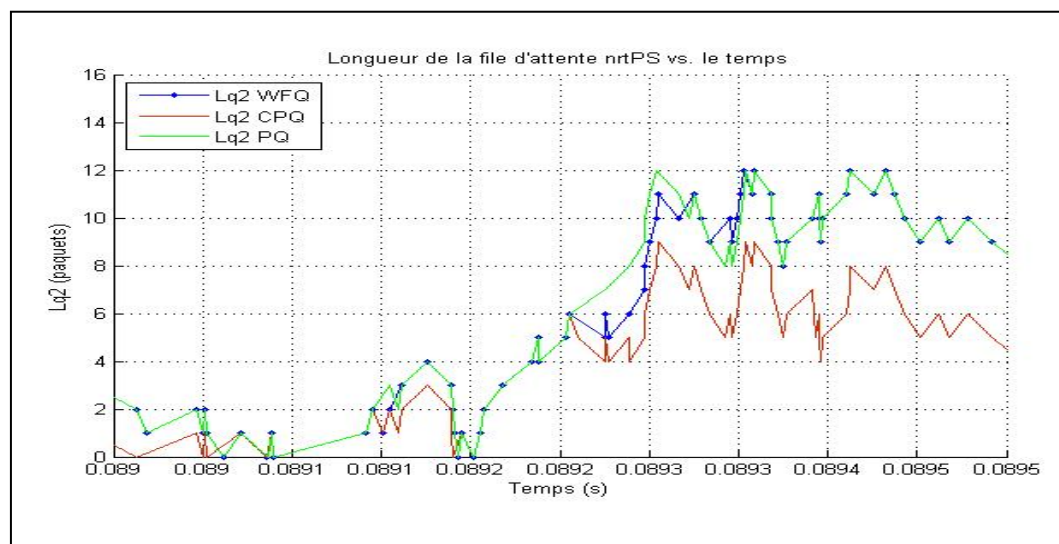


Figure 4.26 longueur de la file d'attente nrtPS : résultats partiels.

Étude du délai d'attente dans la file d'attente rtPS

Comme le nombre de paquets nrtPS bénéficiant de la courtoisie est assez important dans ce scénario, alors le temps d'attente supplémentaire relatif à la transmission de paquets courtois doit être aussi assez considérable, c'est la cause pour laquelle le délai d'attente dans le buffer de la voix se montre beaucoup plus grand dans le cas d'application de CPQ que dans les deux autres solutions, à savoir PQ et WFQ. La Figure 4.27 appuie ce que nous venons de citer. Elle montre aussi que le délai d'attente relatif à WFQ est parfois supérieur à celui qui correspond à PQ. Ceci est une conséquence de la sélection des paquets à servir selon la valeur de leur finish time qui doit être la plus petit possible. Certain paquets nrtPS vont être donc servis malgré que la queue rtPS ne soit pas vide.

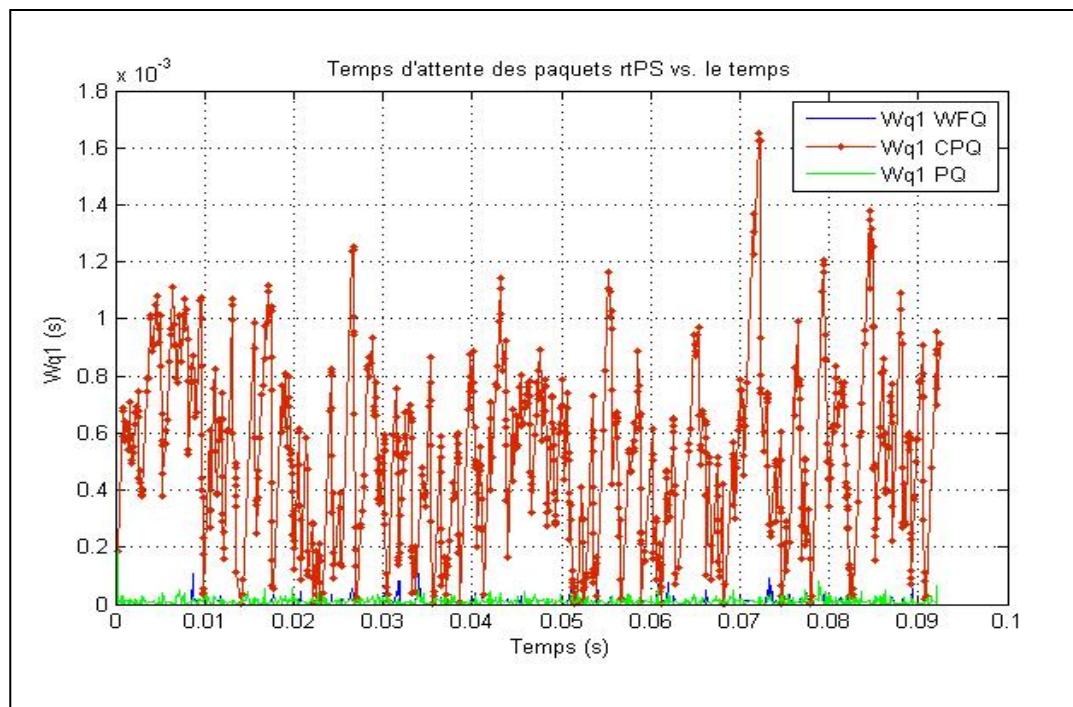


Figure 4.27 Délai d'attente dans la file de la voix (scénario 5).

Étude du délai d'attente dans la file d'attente nrtPS

La Figure 4.28 montre quelques graphes correspondants aux résultats obtenus pour le délai d'attente dans la file de la classe de basse priorité, à savoir nrtPS. PQ offre les délais les plus longs. Par contre, CPQ est la solution qui donne les meilleurs résultats pour ce facteur. Le temps où les paquets de la voix temporisent dans leur file, la file des données va être servie pendant ce temps là. Ce mécanisme ne va que diminuer le délai d'attente pour les paquets nrtPS bénéficiaires de la courtoisie de la classe prioritaire.

WFQ a un comportement qui converge vers celui de PQ, mais parfois cette solution assure des délais restreints par rapport à PQ. Cet algorithme tient à garantir l'équité de service des paquets des différentes classes de service.

Étude du taux de perte de paquets nrtPS

La perte de paquets relative à la classe de trafic nrtPS est illustrée par la Figure 4.29. Cette perte est représentée par les trois graphes en bas de la figure en question. Celui en jaune représente les résultats de CPQ. Celui en bleu concerne WFQ et celui en noir est relatif à PQ. Nous constatons que la perte de paquets causée par notre solution, CPQ, est inférieure à celles données par WFQ et PQ. La courtoisie diminue donc le rejet des paquets de données. Ce rejet peut avoir lieu à cause d'un débordement de la file d'attente des data. Comme nous l'avons déjà vu, CPQ offre l'opportunité de réduire la taille de file d'attente nrtPS pour atteindre la valeur minimale par rapport aux deux autres approches, ce qui minimise la probabilité d'avoir le buffer de données plein. Par conséquent, diminuer la probabilité de rejet de paquets.

Pour conclure, notons que l'algorithme de courtoisie offre l'opportunité de minimiser le délai d'attente pour les données, réduit la taille de la file d'attente nrtPS et surtout il diminue le taux de perte de paquets. Ce résultat est considéré comme étant très important et bénéfique, car le système ne sera pas contraint de retransmettre l'information perdue.

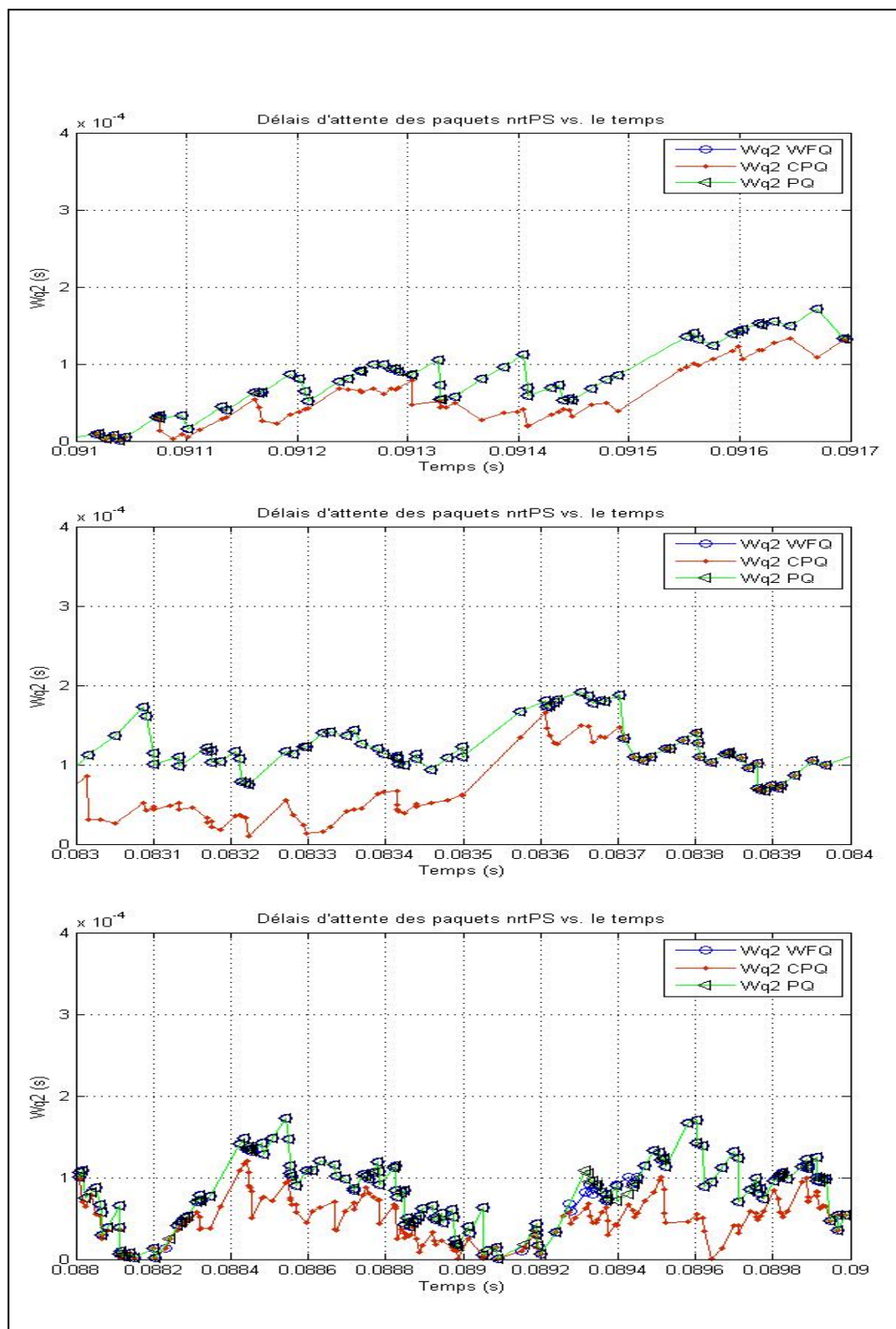


Figure 4.28 Délai d'attente dans la file nrtPS (scénario 5).

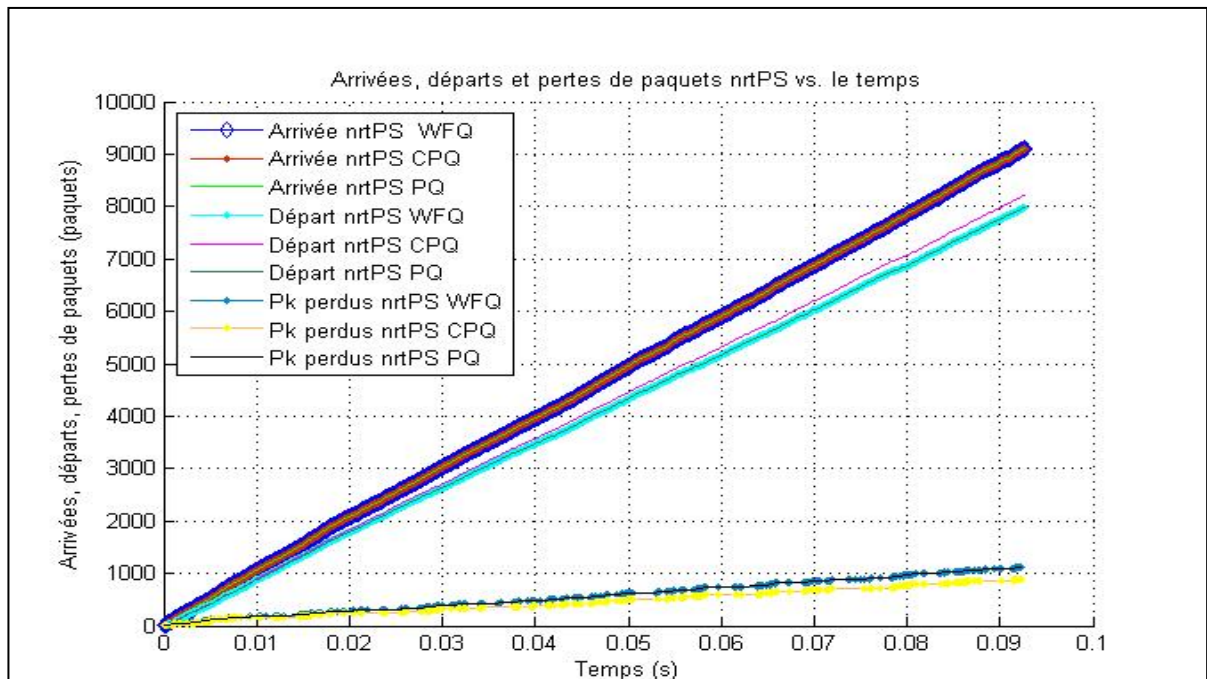


Figure 4.29 pertes de paquets de la classe nrtPS (scénario 5).

4.3.6 Scénario 6 : Étude de l'impact de l'augmentation de R1

L'objectif de ce scénario est de déterminer l'apport de notre solution dans le cas de l'augmentation de la taille de la rafale R1 correspondant à la voix. Cette valeur a été fixée dans les scénarios précédents à 2. Dans ce scénario nous allons la fixer à 6.

Résultats numériques :

Le Tableau 4.7 résume les résultats les plus importants pour les trois algorithmes CPQ, WFQ et PQ.

CPQ demeure toujours l'algorithme qui améliore la QoS de la classe des données. Il est évident que quand la taille d'une rafale de la voix augmente, le nombre des paquets nrtPS privilégiés diminue. Ceci est dû au fait que le nombre de paquets qui seront servis à la place

des paquets de la voix, se calcul en soustrayant la longueur de la file d'attente de la voix et R1 de la taille maximale du buffer, ce qui implique que plus R1 est grand, plus le nombre des paquets courtois est petit. (voir l'équation 3.17)

Reste à dire que même si nous avons augmenté la taille de la rafale de la voix, les résultats donnés par CPQ sont très satisfaisants. Le pourcentage des paquets nrtPS qui sont privilégiés par l'algorithme de courtoisie dépasse 78 % de la totalité des paquets de type data arrivant vers le système de file d'attente considéré dans cette simulation.

En examinant le Tableau 4.6 et le Tableau 4.7, nous constatons que l'augmentation de R1 engendre la convergence des résultats obtenus par CPQ vers les résultats donnés par WFQ.

Tableau 4.7 Résultats numériques relatifs pour le scénario 6

Résultats	CPQ	WFQ	PQ
Nombre de paquets moyen dans la queue rtPS (voix)	2.376	0.13289	0.10695
Nombre de paquets moyen dans la queue nrtPS (data)	6.4904	6.879	6.8855
Délais moyens dans la queue rtPS	2.4304e-005	1.3857e-006	1.1153e-006 s
Délais moyens dans la queue nrtPS	6.6736e-005	7.1723e-005	7.1798e-005 s
Nombre de paquets rtPS perdus	0	0	0
Nombre de paquets perdus nrtPS	988	1115	1116
Nombre de paquets nrtPS bénéficiant de la courtoisie	7164		
% de paquets ayant bénéficiés de la courtoisie	78.7%		

Résultats graphiques

Étude de la longueur de la file d'attente rtPS

Tout comme le scénario 5, celui-ci donne une longueur de la file d'attente rtPS, $Lq1$, plus considérable dans le cas de l'application de CPQ (Figure 4.30), seulement que dans ce cas la longueur maximale ne dépasse pas 6 paquets alors qu'en scénario 5 elle atteint les 10 paquets (voir Figure 4.24) Donc nous pouvons conclure que plus $R1$ augmente plus le comportement de CPQ convergent vers WFQ.

Le comportement de CPQ dans ce scénario est semblable à celui constaté dans le scénario 5. La Figure 4.31 donne une vision plus claire de nos résultats pour $Lq1$.

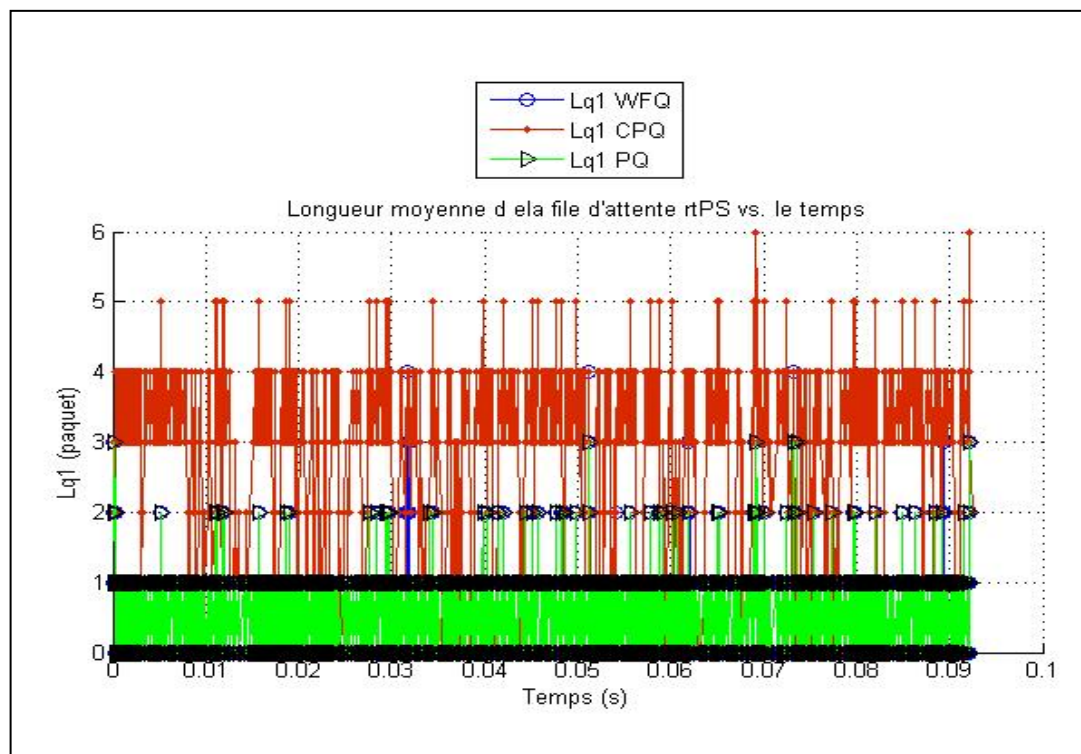


Figure 4.30 Longueur de la file d'attente rtPS (scénario 6).

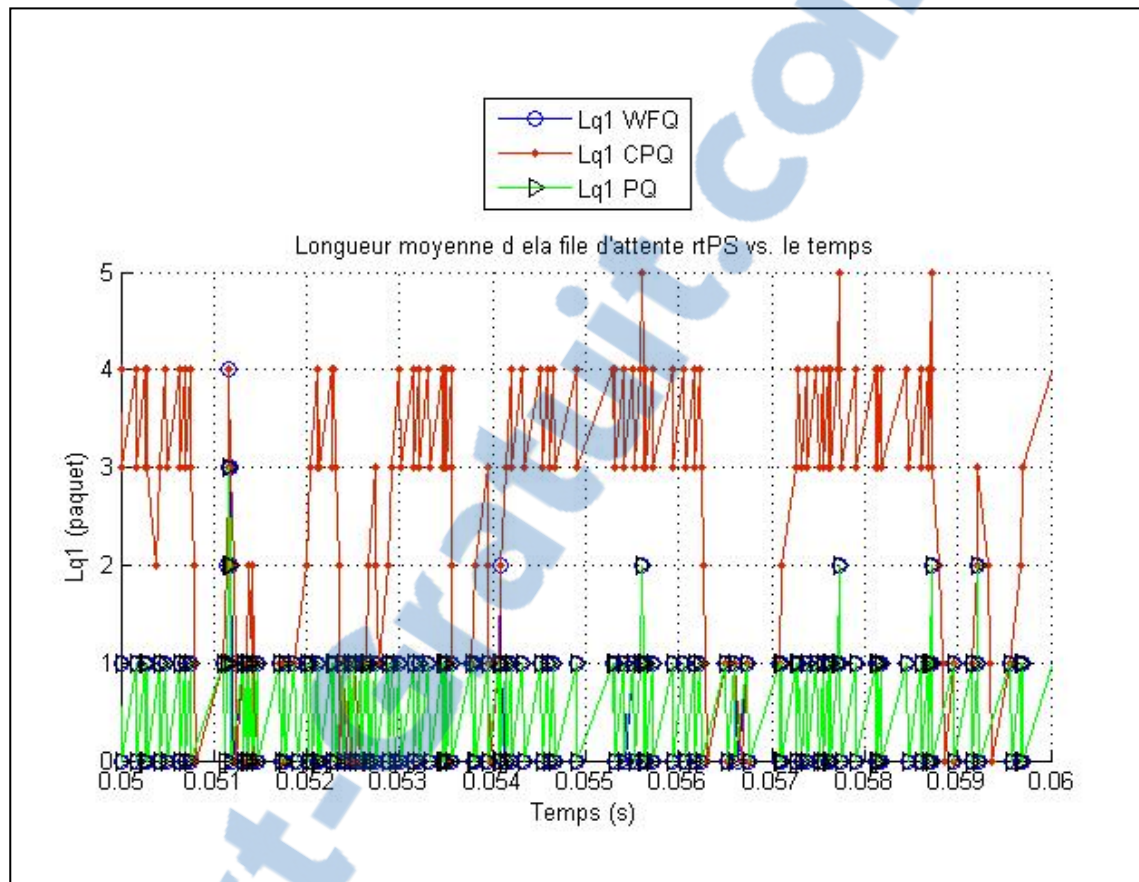


Figure 4.31 Longueur de la file d'attente rtPS : Figure partielle.

La longueur de la file d'attente rtPS est la plus importante dans le cas de l'application de CPQ. WFQ donne des résultats légèrement bons par rapport à notre solution alors que PQ donne les meilleures longueurs pour le buffer du trafic rtPS.

Étude de la longueur de la file d'attente nrtPS

L'augmentation de R1 de 2 à 6 a entraîné une diminution du nombre de paquets nrtPS transmis à la place des paquets rtPS, ce qui a influencé les résultats graphiques donnés par CPQ concernant les données.

En comparant ces résultats avec ceux offerts par le scénario 5 (voir la Figure 4.26 et la Figure 4.33), nous constatons que la longueur de la file d'attente nrtPS obtenue en appliquant CPQ convergent plus vers celle résultante par l'exécution de WFQ .

Tel que pour la longueur de la file d'attente des paquets de la voix, celle relative aux paquets de données suit le même comportement que dans le scénario 5. Autrement dit, sa valeur est plus considérable dans le cas d'application de PQ et moins importante dans le cas d'application de CPQ. Seulement, notons que l'apport de CPQ est moins bon que celui obtenu dans le scénario précédant car la valeur de R1 a été augmentée.

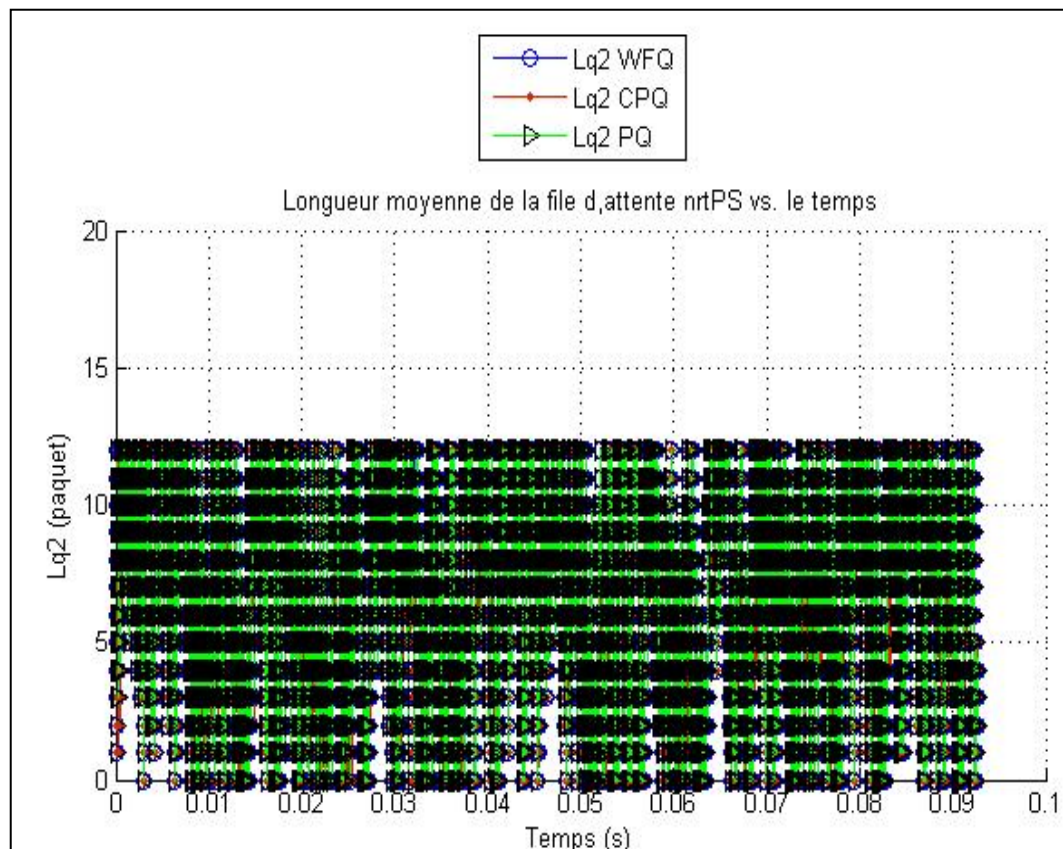


Figure 4.32 Longueur de la file d'attente nrtPS (scénario 6).

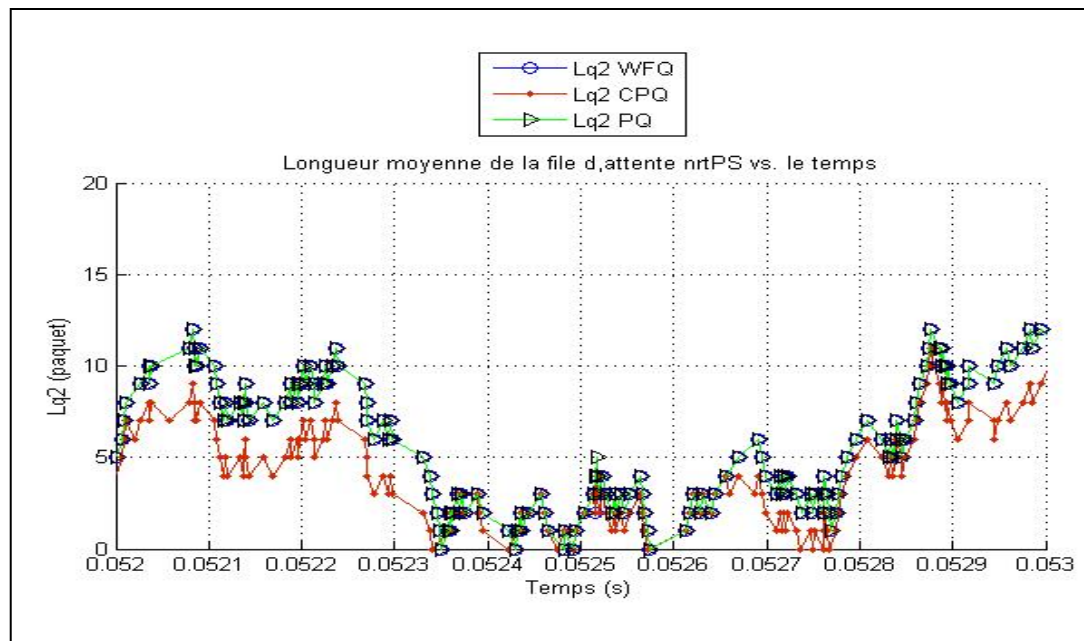


Figure 4.33 Longueur de la file d'attente nrtPS : Portion du graphe global.

Étude du délai d'attente dans la file rtPS

Le délai d'attente dans la file de la voix n'atteint pas 1.2 ms dans ce scénario dans le cas de l'application de CPQ (voir Figure 4.34) alors que dans le cas de $R1 = 2$ et en gardant les autres paramètres identiques (scénario 5), nous constatons que ce délai dépasse 1.6 ms (voir Figure 4.25). Ceci est dû à la diminution de la taille de la file d'attente rtPS, $Lq1$, dans ce scénario par rapport au précédent. La diminution de $Lq1$ est une conséquence de la réduction du nombre de paquets courtois à cause de l'augmentation de $R1$. Comme nous l'avons mentionné ci-dessus dans cette section, $R1$ contribue au calcul du nombre de paquets courtois. Son augmentation diminue le nombre de paquets rtPS aptes à céder leur tour de service.

PQ reste la meilleure solution pour favoriser le trafic de la voix, suivie de WFQ (voir Figure 4.35).

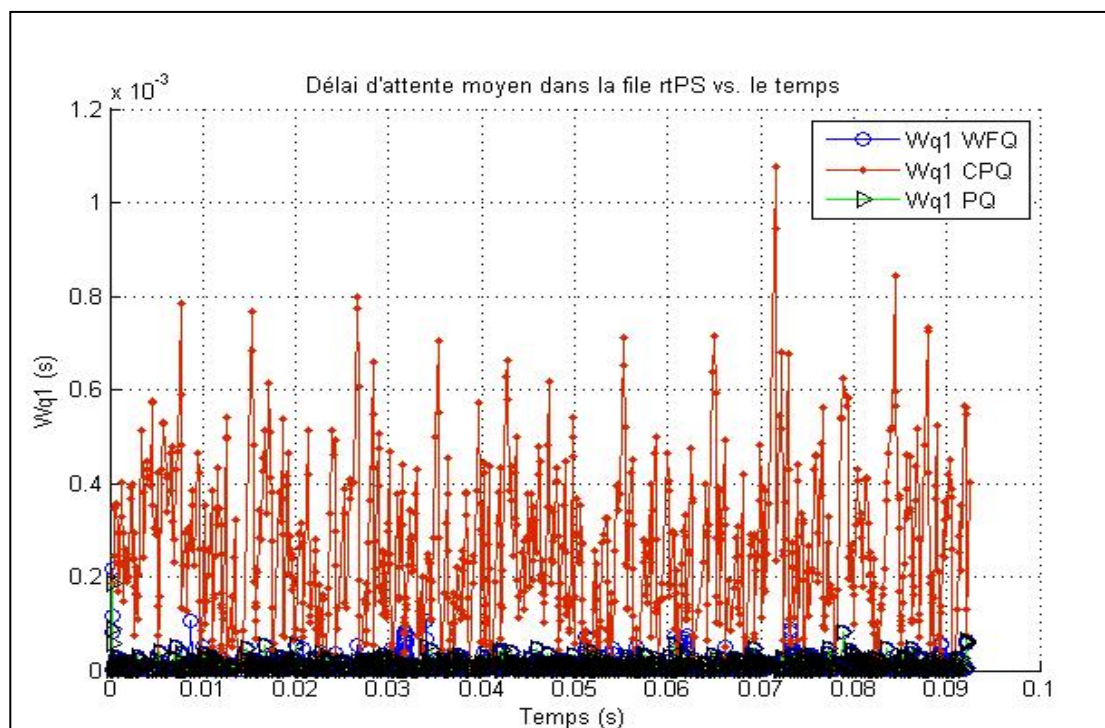


Figure 4.34 Délai d'attente dans la file rtPS (scénario 6).

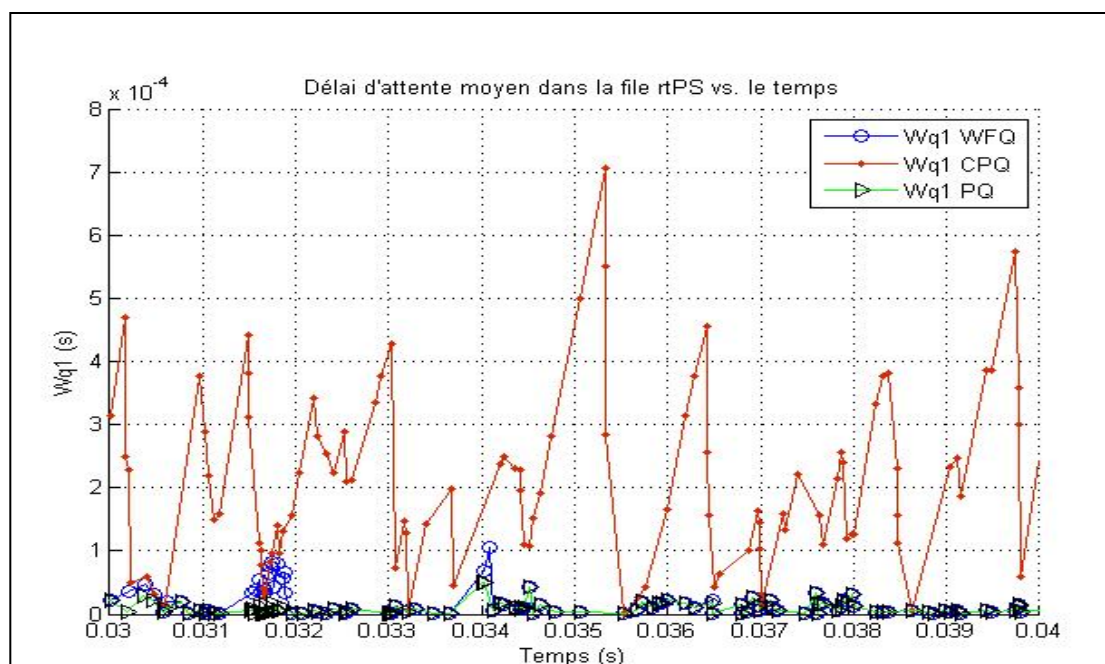


Figure 4.35 Délai d'attente dans la file rtPS (scénario 6).

Étude du délai d'attente dans la file nrtPS

La Figure 4.37 représente une partie de la Figure 4.36. Toutes les deux illustrent le délai d'attente dans la file de data, noté $Wq2$.

Les résultats donnés par ce scénario gardent CPQ comme meilleure solution pour $Wq2$. Alors que PQ est celle qui donne les délais les plus longs. Seulement, nous constatons que le passage d'une valeur de $R1 = 2$ à une autre valeur $R1 = 6$ diminue la performance de CPQ relative à $Wq2$. La Figure 4.28 montre des délais plus petits que la Figure 4.37. Nous déduisons donc que l'augmentation de la valeur de $R1$ influe sur le comportement de CPQ qui commence à converger vers celui de WFQ.

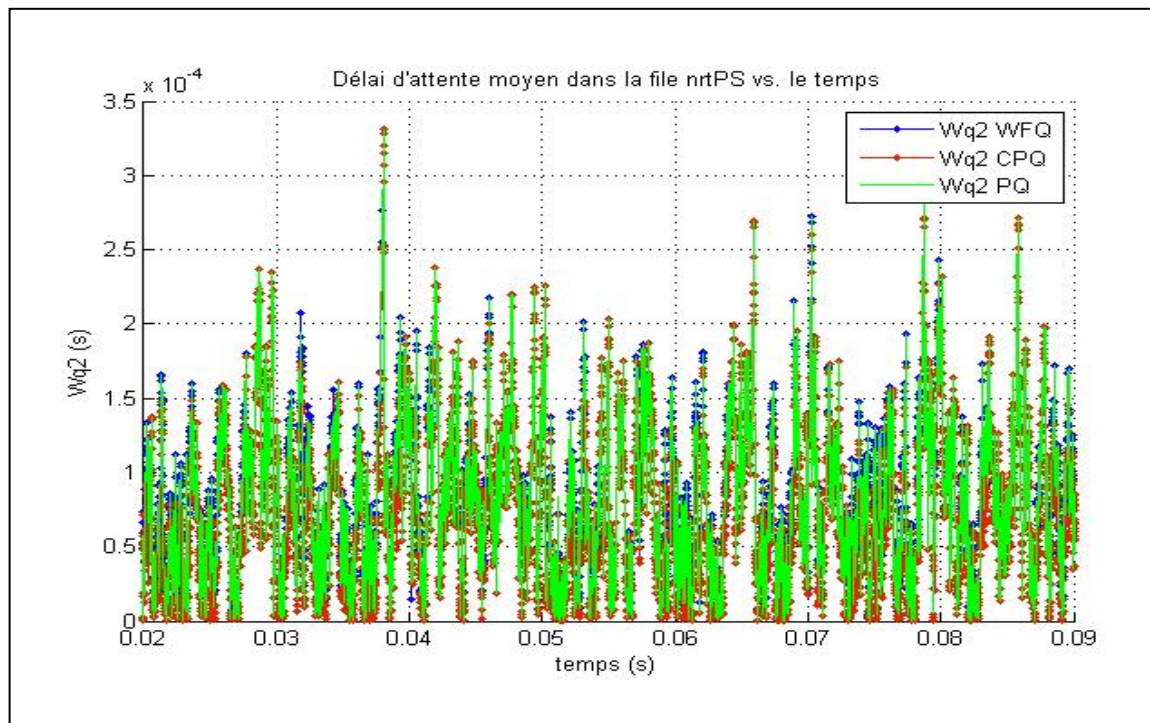


Figure 4.36 Délai d'attente dans la file nrtPS (scénario6).

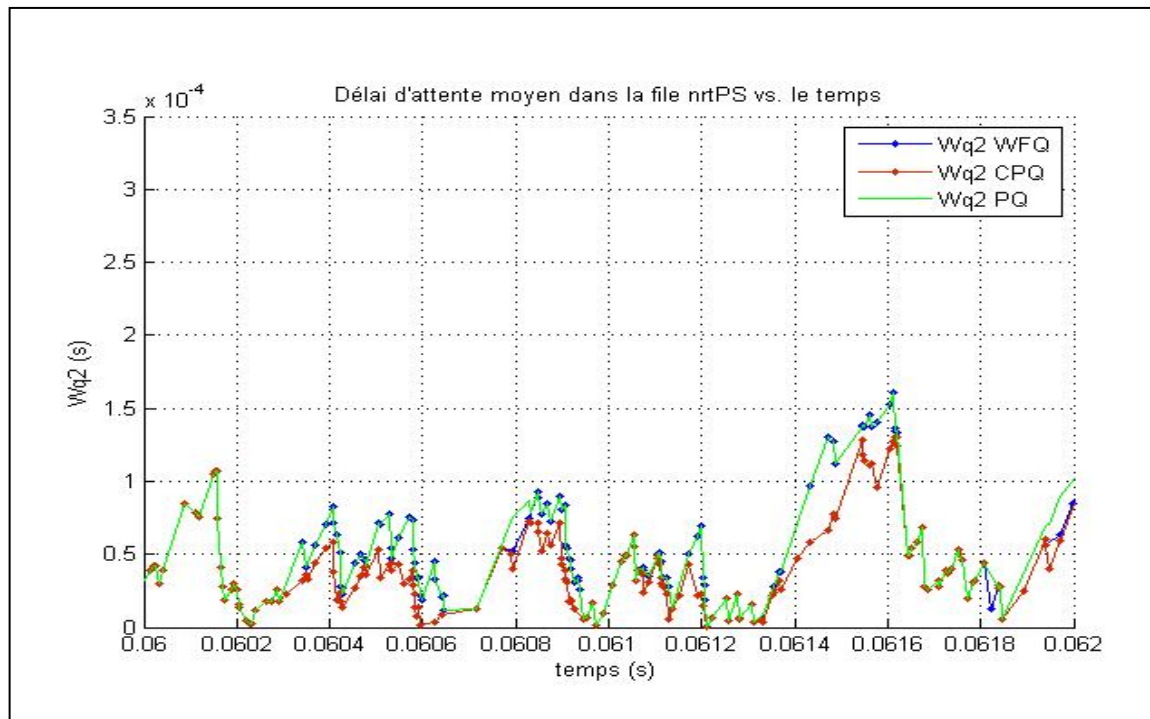


Figure 4.37 Délai d'attente dans la file nrtPS : Figure partielle.

Étude de la perte de paquets

La perte de paquets relative à la classe nrtPS est proportionnelle avec la valeur de $R1$. Plus $R1$ augmente, plus le nombre de paquets perdus devient important. Nous expliquons ce comportement par le fait que $R1$ contribue dans le calcul du nombre de paquets courtois, une valeur plus petite augmente le nombre de ces paquets, ce qui permet de servir plus de paquets nrtPS, par conséquent la probabilité de perte de paquets va diminuer. Ce résultat est illustré par le graphe rose de la Figure 4.38. Réciproquement, tel que le montre le graphe marron de cette même figure, l'augmentation de la valeur de $R1$ va élever la probabilité de perte de paquets nrtPS en augmentant le nombre des paquets rejetés appartenant à ce type de trafic.

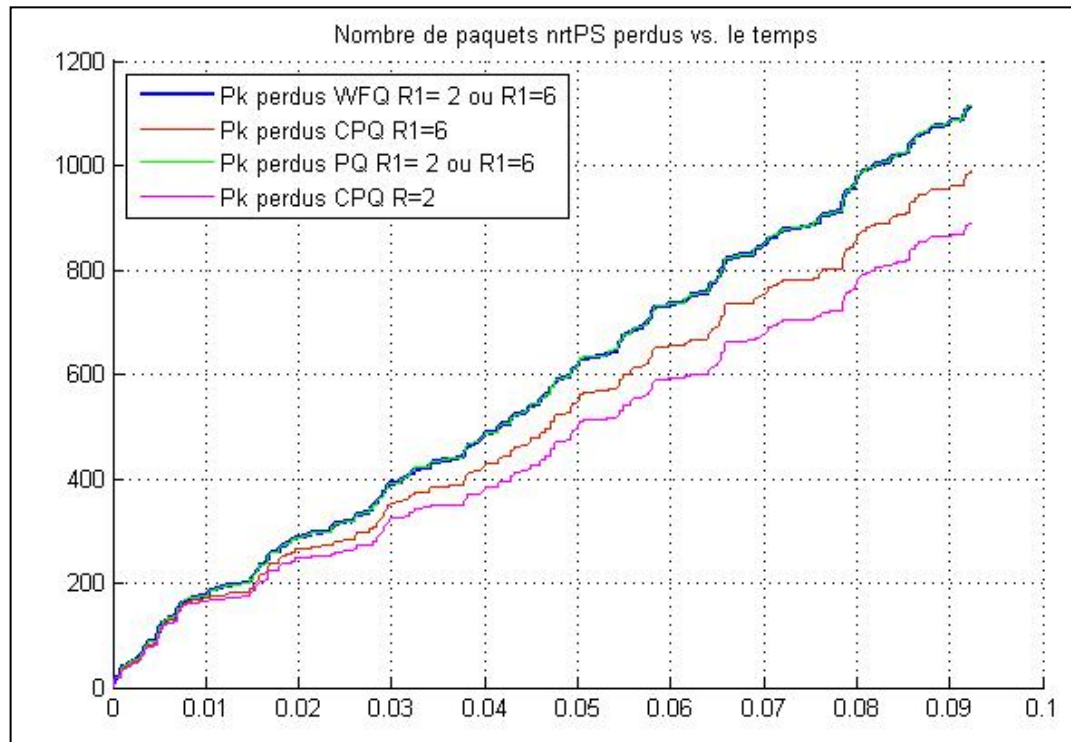


Figure 4.38 Nombre de paquets perdus vs. Le temps : R=2 et R=6.

4.4 Apport de l'algorithme de courtoisie

Dans cette section nous allons présenter l'apport positif de notre approche. Pour ce faire, nous allons tracer le graphe relatif au pourcentage de paquets nrtPS ayant bénéficié de la courtoisie versus le taux moyen d'arrivée nrtPS.

Dans le but d'obtenir des résultats fiables et représentatifs, nous allons considérer le rapport $r_{\lambda_2} = \lambda_2 / (\lambda_1 + \lambda_2)$. Cette fraction représente le taux moyen du trafic nrtPS arrivant au système de files d'attente par rapport au taux moyen du trafic global. De cette manière, nous déterminons quand est-ce que notre algorithme est recommandé, et pour quel taux d'arrivée de trafic nrtPS il serait bon de l'appliquer.

Les données utilisées pour le tracé du graphe sont extraites des résultats numériques des scénarios 1 à 5. Nous n'avons pas pris en considération le scénario 6 car il a une valeur de $R1$ différente aux autres. Nous pouvons donc construire le tableau suivant :

Tableau 4.8 Apport de CPQ en fonction de taux de trafic nrtPS

Délai maximal pour rtPS	10 ms	10 ms	20 ms	20 ms	20 ms
λ_1 (paquet/sec)	2500	250	5000	5000	10000
λ_2 (paquet /sec)	10000	10000	20000	40000	100000
r_{λ_2}	80 %	97.5%	80%	88.8%	90.9%
Apport de CPQ	7.5758%	2.1561%	13.26%	22.24%	82.49%

Comme nous le voyons dans le Tableau 4.8, deux délais d'attente maximaux ($D_{\max_rtPS}$) pour la voix sont pris en compte dans nos scénarios. Nous allons donc procéder au tracé de deux graphes. Chacun va correspondre à un délai. Rappelons que le délai d'attente maximal pour les paquets de la classe rtPS correspond au temps d'attente maximal des paquets de la voix dans leur file d'attente. Après l'écoulement de ce temps, le paquet en question sera rejeté. La raison d'appliquer ce mécanisme est du fait que la voix est sensible aux délais et que le délai de bout en bout ne doit pas dépasser 200 ms.

La Figure 4.39 illustre l'apport positif de notre algorithme pour les deux cas de figure : $D_{\max_rtPS} = 10$ ms (en bleu) et $D_{\max_rtPS} = 20$ ms (en rouge).

En examinant la figure ci-dessous nous constatons que pour un $D_{\max_rtPS}$ égal à 10 ms, plus la portion des paquets nrtPS est importante dans le réseau, plus les paquets rtPS se montrent moins courtois. La cause d'avoir un tel résultat est le fait qu'un paquet rtPS ne soit pas capable d'attente plus que 10 ms dans sa file, ce qui implique une diminution du temps

ξ_1 représentant le temps d'attente supplémentaire des paquets courtois dans leur file pour permettre au trafic de donnée d'être servi. La réduction de ξ_1 entraîne la réduction du *coeff_{courtoisie}* qui représente le nombre de paquets de données qui peuvent être servis avant leur tour. Ceci explique bien la réduction de leur pourcentage représenté dans la Figure 4.39.

Concernant le graphe relatif à $D_{\max_rtPS}$ équivalent à 20 ms, nous constatons que l'apport positif de notre algorithme augmente avec la croissance du taux de trafic nrtPS dans le réseau. Le fait d'élever la valeur du délai maximal d'attente des paquets de la voix dans leur file permet d'augmenter la valeur de ξ_1 , car tant que ce délai n'est pas atteint les paquets peuvent attendre davantage dans leur file. *coeff_{courtoisie}* étant relative à ξ_1 , sa valeur va donc croître à son tour. Par conséquent, le pourcentage des paquets nrtPS privilégiés par CPQ va augmenter.

Nous constatons que notre algorithme est hautement recommandé dans le cas de la présence d'un trafic nrtPS important par rapport au trafic rtPS. Ce dernier ne doit pas avoir un temps d'attente maximal dans la file trop restreint.

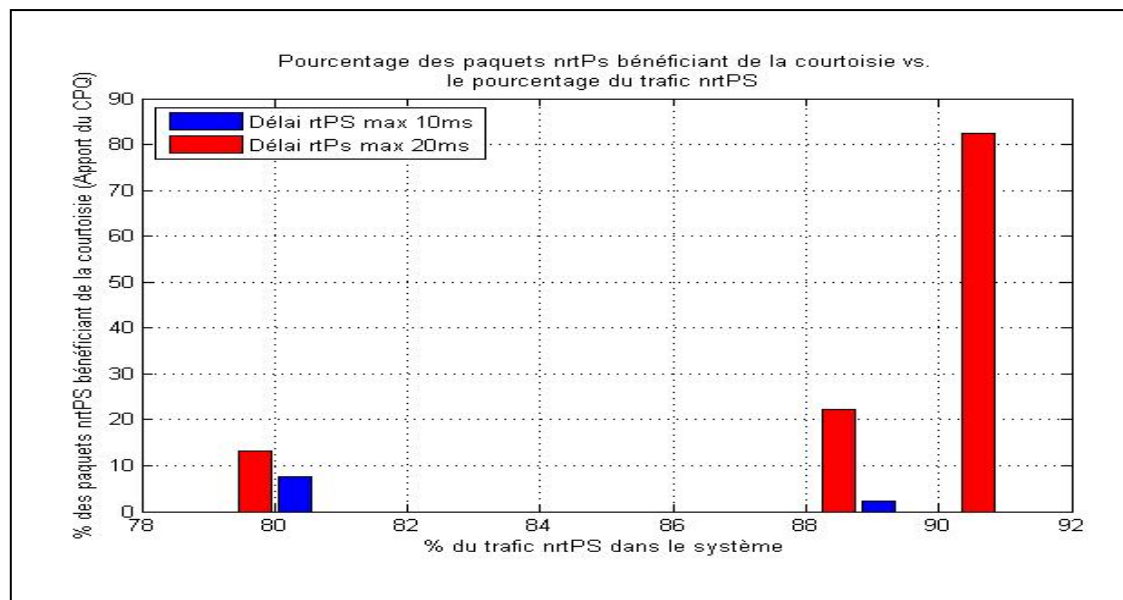


Figure 4.39 Apport de l'algorithme de courtoisie.

4.5 Conclusion

Dans ce chapitre nous avons procédé à la validation de notre modèle de gestion de file d'attente (CPQ) à travers la simulation. Nous avons donc implémenté l'algorithme conçu dans le chapitre précédent dans le simulateur Matlab afin de réaliser nos scénarios.

Le but de notre solution est d'optimiser la performance de la QoS dans les réseaux WiMAX fixes en améliorant le service des classes moins prioritaires. Notre apport consiste à permettre aux paquets de basse priorité d'être servis à la place des paquets de haute priorité quand ces derniers peuvent attendre davantage dans leur file d'attente sans affecter leur QoS. Ceci a diminué la taille de la file d'attente de la classe de basse priorité, ainsi que son délai d'attente, ainsi que le taux de pertes de paquet relatif à cette classe.

Dans ce chapitre nous avons considéré les deux classes de trafic rtPS et nrtPS. La classe rtPS est considérée comme la plus prioritaire et l'autre est la moins prioritaire. Nous avons réalisé un ensemble de scénarios où nous avons implémenté un système de gestion de files d'attente M/G/1 avec priorité non préemptive au sein d'une station de base d'un réseau WiMAX pour le canal montant (uplink). Pour chaque scénario, nous avons appliqué les algorithmes d'ordonnancement CPQ, PQ et WFQ séparément, en considérant les mêmes paramètres de bases, tels que les taux d'arrivées des paquets de types rtPs et nrtPS, la taille des files d'attente et leur poids, la taille maximale d'une rafale du trafic voix, le délai d'attente maximal d'un paquet rtPS, etc.

Nous avons considéré un scénario de référence représentant le taux d'arrivées idéal pour les deux types de trafic qui maintient le système dans un état de stabilité. D'autres scénarios ont été effectués afin d'examiner la robustesse de notre algorithme face au changement des paramètres qui peuvent affecter son comportement.

Les résultats obtenus avec CPQ ont été comparés avec les deux approches PQ et WFQ. Ils montrent que l'apport de notre approche est similaire à celui de WFQ dans le cas d'un taux

de trafic de la voix assez élevé par rapport à FTP. Dans ce cas le nombre de paquets courtois diminue, autrement la taille de la file d'attente rtPS va dépasser le seuil de remplissage relatif à la perte de paquets tolérée. Par contre, une quantité de trafic de données assez importante par rapport à la quantité de la voix présente dans le réseau recommande hautement l'utilisation de CPQ car comme on l'a vu dans les différents scénarios, l'apport de CPQ est très encourageant. La courtoisie des paquets rtPS est arrivée à privilégier plus de 80 % des paquets de basse priorité (voir la Figure 4.39).

Notons que le délai d'attente maximal toléré pour les paquets de la voix dans leur file d'attente, noté $D_{\max_rtPS}$, influence le nombre de paquets nrtPS bénéficiant de la courtoisie. Il est clair que plus un paquet rtPS est apte à attendre dans sa file sans atteindre la valeur $D_{\max_rtPS}$ plus le nombre de paquet nrtPS profitant de la courtoisie augmente.

Finalement, nous déduisons que notre approche est apte à apporter une forte amélioration dans les réseaux WiMAX disposant de plusieurs connexions FTP par rapport aux connexions de la VoIP avec un délai d'attente $D_{\max_rtPS}$ raisonnable pour les paquets de la voix. Autrement, notre solution se comporte tel que l'algorithme d'ordonnancement WFQ.

CONCLUSION

Les réseaux WiMAX offrent aux fournisseurs la possibilité d'étendre leurs services même dans des endroits isolés. C'est une solution performante en termes de débit et de portée et à faible coût en la comparant avec les réseaux filaires. Il suffit de mettre en place une station de base BS, et munir les utilisateurs du réseau d'une simple antenne Indoor ou Outdoor. La couverture d'une BS peut atteindre les 50 Km.

Dans ce type de réseaux, la garantie de la QoS demeure un point crucial, Nombreuses sont les recherches menées pour son optimisation. Cette problématique est loin d'être simple, elle est plutôt complexe, plusieurs pistes de recherches en dérivent. Certaines focalisent sur les mécanismes de contrôle d'admission des appels, d'autres tentent de résoudre les problèmes reliés à la mobilité, la coexistence du WiMAX avec d'autres types de réseaux prend aussi sa part d'intérêt, tout comme le problème de mise à l'échelle des solutions proposées, certes, nombreuses sont celles qui traitent le volé d'ordonnancement dans ces réseaux.

Quoique de multiples pistes aient été explorées afin de présenter des algorithmes s'adaptant le plus aux exigences des applications du 802.16, la quasi-totalité défavorise le trafic nrtPS par rapport au trafic rrtPS. Il est évident que le service à temps réel est très sensible au délai, mais cela n'empêche pas d'offrir plus d'avantages pour le service nrtPS. Réduire le taux de perte de paquets, ainsi que diminuer le délai relatif au trafic non real-time dans les réseaux WiMAX va optimiser les performances de ces systèmes. Cela dit, il ne faut point dégrader la QoS de la classe prioritaire.

Notre solution consiste à développer un système de gestion des différents types de trafic dans les réseaux 802.16. Contrairement à la majorité des solutions qui existe dans la littérature, notre algorithme cherche à améliorer le service des trafics moins prioritaires tout en répondant aux exigences des trafics plus prioritaires. Le principe de notre approche consiste à substituer les paquets de haute priorité par ceux de moindre priorité pour les transmettre.

Ceci n'est possible que si le taux de pertes maximal de paquets de la classe qui cède sa place n'est pas atteint, par contre, celui correspondant aux paquets bénéficiant de la courtoisie doit être dépassé. En outre, l'attente supplémentaire des paquets substitués ne doit pas excéder le temps d'attente maximal toléré. Ce temps-ci est proportionnel au nombre de paquets attendus pour atteindre le seuil de remplissage de la file d'attente de la classe courtoise relatif au taux de perte de paquets maximal toléré.

Notre mécanisme s'applique au canal montant de la station de base, comme il peut se déployer au sein d'une station d'abonné jouant le rôle d'un point d'accès pour d'autres SS. Dans notre étude, nous avons focalisé sur les deux types de trafic rtPS et nrtPS. Un cas général a été présenté dans ce travail pour faire allusion que notre solution peut s'appliquer pour plusieurs types d'appels, non seulement que nous pouvons implanter une file d'attente pour les trois classes rtPS, nrtPS et BE, mais aussi nous pouvons appliquer une différenciation au sein de la même classe en créant une queue pour les appels courants, une autre pour les nouveaux appels et une troisième pour les handoff en cas d'un réseau WiMAX mobile. Il peut même être étendu pour couvrir le cas des réseaux hétérogènes WiMAX/WiFi. En conséquence, la mise à l'échelle ne pose aucun problème pour notre solution.

Dans ce travail nous avons proposé un modèle mathématique correspondant au système de file d'attente de notre réseau. Ce modèle a permis, entre autres, le calcul des délais d'attente moyen dans les files d'attente, leurs longueurs moyennes, le temps d'attente supplémentaire dans la file de la classe courtoise, la taille maximale d'une rafale, ainsi que la priorité des paquets. Pour des raisons de complexité, notre analyse n'a pas été poussée au point d'obtenir des résultats numériques. Néanmoins, nous avons validé notre étude par simulation, en utilisant Matlab et en comparant nos résultats avec les deux algorithmes d'ordonnancement PQ et WFQ. Dans certains cas nos résultats sont similaires à ceux de WFQ et meilleurs que ceux de PQ, mais dans d'autres cas ils sont très performants par rapport aux deux. Nos expériences montrent que l'application de notre algorithme est propice dans le cas de présence d'augmentation de trafic nrtPS, le cas où plus de 80% des paquets de la classe de basse priorité ont bénéficié de la courtoisie donne des résultats très encourageants. Par

contre, l'existence d'un taux de trafic très important appartenant à la classe prioritaire rend le comportement de l'algorithme de courtoisie semblable à celui de WFQ.

Notre solution répond aux objectifs que nous avons fixés au début de notre étude. Non seulement qu'elle assure une équité d'attente parmi les différentes classes en introduisant le mécanisme de calcul de priorité pour chaque paquet, mais aussi elle arrive à réduire le délai d'attente pour les paquets de moindre priorité et leurs taux de pertes de paquets. De plus, elle respecte les contraintes de notre modèle, tant qu'elle ne détériore pas la QoS de la classe prioritaire.

RECOMMANDATIONS

Du fait que l'algorithme de courtoisie n'apporte pas ses fruits à tous les coups, nous recommandons le développement d'un mécanisme pour son déclenchement dynamique au moment propice en nous basant sur les critères qui nécessitent son application, tels que le taux des deux types de trafic ainsi que la taille des rafales.

La taille des rafales étant fixée dans nos expérimentations, il serait aussi intéressant d'appliquer la formule correspondante que nous avons développée dans notre modèle mathématique afin d'obtenir des résultats plus proches de la réalité.

Achever l'analyse mathématique en poussant l'étude jusqu'à l'obtention des valeurs numériques est un point important à développer afin de pouvoir confronter les résultats de la simulation avec les résultats numériques.

Finalement, nous souhaiterons rendre cette solution plus robuste en introduisant un mécanisme de fragmentation de paquets de données. Ceci va réduire le taux de pertes de paquets de la classe prioritaire.

Selon (Davidson and Fox 2002) même si le système de file d'attente fonctionne convenablement en favorisant le trafic de la voix, cette classe de trafic peut parfois voir sa QoS se dégrader. Supposons qu'un paquet prioritaire arrive vers une file d'attente vide suite à la transmission de tous ses paquets. Dans le cas où le serveur soit occupé par le traitement d'un paquet de données, le paquet de voix sera mis en attente dans sa file.

L'exemple suivant a été extrait de (Davidson and Fox 2002) pour un lien de 64 kb/s et des paquets FTP de 1500 octets, le paquet de la voix mis en attente doit attendre un délai de 187.5 ms, calculé comme suit :

$$\textit{Serialization delay} = 1500 \text{ bytes} \times 8 \text{ bits/bytes} / (64\,000 \text{ bits/sec}) = 187.5 \text{ ms}$$

Or, les paquets de la voix sont envoyés chaque 20 ms dans le cas de PCM, avec un délai de bout en bout acceptable de 150 ms.

Un mécanisme de fragmentation des paquets de données va permettre de minimiser le délai d'attente de la voix. Pour ce faire, il est recommandé de subdiviser un paquet de 1500 octets en un ensemble d'unités de telle sorte à en envoyer une chaque 10 ms. La taille d'un fragment se donnera donc comme suit :

$$\textit{Fragmentation size} = (0.01 \text{ secondes} \times 64\,000 \text{ bps}) / (8 \text{ bits /byte}) = 80 \text{ bytes}$$

Donc sur un lien de 64 kbps, un fragment de 80 octets prend 10 ms pour être transmis. Cependant, il va falloir permettre au paquet prioritaire en attente d'être transmis juste après le premier fragment, autrement cette solution n'apportera aucun avantage.

Cette solution peut être introduite à notre approche afin de réduire le délai d'attente des paquets prioritaire, ce qui minimise leur taux de pertes de paquets. Cependant, cet ajout « nous laisse toujours du temps pour la courtoisie ».

LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES

- Belghith, A. and L. Nuaymi (2008). Comparison of WiMAX scheduling algorithms and proposals for the rtPS QoS class. EW 2008. 14th European Wireless Conference: 1-6.
- Bose, S. K. (2002). An introduction to queueing systems New York : Kluwer Academic/Plenum Publishers
- Chen, K.-C. and J. R. B. d. Marca (2008). Mobile WiMAX Hoboken, N.J. : John Wiley, 2008.
- Davidson, J. and T. Fox (2002). Deploying Cisco Voice over IP Solutions, Cisco Press.
- Ekelin, C. (2006). Clairvoyant non-preemptive EDF scheduling. 18th Euromicro Conference on Real-Time Systems. Germany 23-29.
- Gross, D. and C. M. Harris (1998). Fundamentals of queueing theory. New York : Wiley.
- Ho, S. W. (2004). "Adaptive Modulation (QPSK, QAM)." Retrieved nov. 2007, 2007, from <http://www.intel.com/netcomms/technologies/wimax/303788.pdf>.
- Kadoch., M. (2004). Protocoles et réseaux locaux l'accès Internet École de technologie supérieure.
- Kalai Kalaichelvan, L. H. (2007). Wireless Broadband System Parts ALTHOS.
- Kleinrock, L. (1975). Queueing systems New York : Wiley
- Lera, A., A. Molinaro, et al. (2007). "Channel-Aware Scheduling for QoS and Fairness Provisioning in IEEE 802.16/WiMAX Broadband Wireless Access Systems "Network, IEEE 21(5): 34-41.
- Lin, J.-C., C.-L. Chou, et al. (2008). Performance Evaluation for Scheduling Algorithms in WiMAX Network. AINAW 2008. 22nd International Conference on Advanced Information Networking and Applications - Workshops 25-28 March 2008
- Lo, S.-C. and Y.-Y. Hong (2008). A novel QoS scheduling approach for IEEE 802.16 BWA systems ICCT 2008. 11th IEEE International Conference on Communication Technology 46-49.

- Mehrjoo, M. and X. Shen (2008). An efficient scheduling scheme for heterogeneous traffic in IEEE 802.16 wireless metropolitan area networks. IST 2008. International Symposium on Telecommunications 263-267.
- Meyer, J. (2005). 3G Wireless with WiMAX and Wi-Fi: 802.16 and 802.11, McGraw-Hill.
- Nuaymi, L. (2007). WiMAX: Technology for Broadband Wireless Access. England, John Wiley & Sons.
- Pareek, D. (2006). WiMAX : Taking Wireless to the MAX. Boca Raton, FL, Auerbach Publications.
- Prasath, G. A., C. P. Fu, et al. (2008). QoS scheduling for group mobility in WiMAX ICCS 2008. 11th IEEE Singapore International Conference on Communication Systems 1663-1667.
- Qu, G., Z. Fang, et al. (2008). A new QoS Guarantee Strategy in IEEE802.16 MAC Layer ICPCA 2008. Third International Conference on Pervasive Computing and Applications. 1: 157 - 162
- Ruangchaijatupon, N., L. Wang, et al. (2006). A Study on the Performance of Scheduling Schemes for Broadband Wireless Access Networks. ISCIT '06. International Symposium on Communications and Information Technologies: 1008 - 1012.
- Sauter, M. (2006). Communication Systems for the Mobile Information Society. England, John Wiley & Sons.
- Sweeney, D. (2006). WiMax operator's manual : building 802.16 wireless networks. Berkeley, CA, Apress.
- Systems, C. (2005). "Comment utiliser la gestion des ressources radio pour fournir des services fiables et sécurisés de réseaux sans fils." White Paper Cisco Systems, Inc., from http://www.cisco.com/web/CH/fr/assets/docs/Ressources_radio.pdf.
- Thaliath, J., M. M. Joy, et al. (2008). Service class downlink scheduling in WiMAX COMSWARE 2008. 3rd International Conference on Communication Systems Software and Middleware and Workshops: 196 - 199
- Thomas Hardjono, L. R. D. (2005). Security in Wireless LANs and MANs, Artech House
- Vinay, K., N. Sreenivasulu, et al. (2006). Performance evaluation of end-to-end delay by hybrid scheduling algorithm for QoS in IEEE 802.16 network
- IFIP International Conference on Wireless and Optical Communications Networks: 5 pp.-5

- Wang, Y., S. Chan, et al. (2008). Priority-Based fair Scheduling for Multimedia WiMAX Uplink Traffic. ICC '08. IEEE International Conference on Communications 301-305.
- Wenfeng Du, Zhen Ji, et al. (2008). Waiting Queue Based Bandwidth Allocation Architecture for nrtPS and BE Services in IEEE 802.16 PMP Network WiCOM '08. 4th International Conference on Wireless Communications, Networking and Mobile Computing 1-5.
- Xie, X., H. Chen, et al. (2008). Simulation Studies of a Fair and Effective Queueing Algorithm for WiMAX Resource Allocation ChinaCom 2008. Third International Conference on Communications and Networking in China: 281-285.