

TABLE DES MATIÈRES

	Page
INTRODUCTION.....	12
CHAPITRE 1 LES RÉSEAUX SANS FIL AD HOC.....	18
1.1 Introduction.....	17
1.2 La fonctionnalité de la couche MAC dans Ad hoc	19
1.2.1 Les méthodes d'accès au canal	19
1.2.2 Plusieurs canaux (Muti-canaux)	22
CHAPITRE 2 LES MECANISMES DE SUPPORT DE LA QUALITÉ DE SERVICE DANS LES RÉSEAUX AD HOC.....	25
2.1 Introduction.....	25
2.2 Classification des méthodes supportant la QoS dans Ad hoc	25
2.3 Classification des approches de QoS	26
2.4 La classification basée sur les couches (MAC & Réseau).....	27
2.5 La QoS dans la couche MAC.....	28
2.5.1 MACA/PR.....	28
2.5.2 Le protocole RTMAC (Real-Time medium access control protocol)	29
2.6 Les mécanismes de contrôle de la Qos dans la couche MAC.....	30
2.6.1 L'ordonnancement équitable (Fair Scheduling).....	31
2.6.2 Le Fair Queueing (FQ).....	31
2.6.3 Le Service Fluide GPS (Generalized Processor Sharing).....	32
2.6.4 Le PGPS ou WFQ.....	32
2.6.5 Self Clocked Fair Queuing (SCFQ).....	33
2.6.6 Start-Time Fair Queuing (STFQ).....	34
2.6.7 FAIR QUEUEING distribué dans les réseaux Ad hoc	34
2.7 Les approches de QoS proposé avec la plate forme 'détection multi-usagers'	36
CHAPITRE 3 PROPOSITION D'UN ORDONNANCEMENT ÉQUITABLE AVEC DÉTECTION MULTI-USAGERS	37
3.1 Introduction.....	37
3.2 Modèle de ordonnancement équitable dans la couche MAC.....	37
3.3 Description du modèle	37
3.3.1 Modèle	38
3.3.2 Matrice et Graphe de dépendance de flots.....	38
3.3.3 Start-time Fair Queuing (SFQ)	42
3.3.4 Construction du regroupement.....	44
3.3.5 Gestion du trafic.....	45
3.3.6 Comment satisfaire les deux critères équité & débit en même temps ?.....	46
3.3.7 Théorie des jeux.....	47
3.3.8 Formulation de la combinaison équité et débit avec la théorie des jeux	48

3.4	Les différentes objectives de l'ordonnancement et leurs implémentations	53
3.4.1	Ordonnancement basé sur Start time Fair Queuing (STFQ).....	53
3.4.2	Ordonnancement basé sur 'Time out priority' (TOP')	53
3.4.3	Ordonnancement basé sur 'Throughput maximization' (TM).....	54
3.4.4	Ordonnancement basé sur l'arbitrage Nash (NASH)	54
3.4.5	Ordonnancement basé sur la maximisation de la somme	55
CHAPITRE 4 IMPLÉMENTATION ET ÉTUDE DES PERFORMANCES.....		56
4.1	Introduction.....	56
4.2	Étude de performances.....	57
CONCLUSION.....		64
LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES.....		66

LISTE DES TABLEAUX

	Page
Tableau 3.1	Matrice de dépendance de flots exemple 140
Tableau 3.2	Matrice de dépendance de flots exemple 241
Tableau 4.1	Les paramètres de simulation.....56

LISTE DES FIGURES

	Page
Figure 1.1	Exemple de l'architecture Ad hoc18
Figure 1.2	Chronogramme d'une émission par le protocole CSMA/CA.....20
Figure 1.3	L'échange RTS/CTS.....21
Figure 1.4	CDMA.....23
Figure 1.5	FDMA.....24
Figure 1.6	TDMA.....24
Figure 2.1	Classification des approches de QoS26
Figure 2.2	Classification basées sur les couches.....27
Figure 2.3	La discipline de service GPS32
Figure 2.4	Les modèles d'accès36
Figure 3.1	Graphe des noeuds exemple 1.....40
Figure 3.2	Graphe de dépendance de flots exemple 1.....40
Figure 3.3	Graphe des noeuds exemple 2.....41
Figure 3.4	Graphe de dépendance de flots exemple 2.....42
Figure 3.5	Exemple de SFQ44
Figure 3.6	Arbre de jeux.....44
Figure 3.7	Exemple de domaine de négociation de Pareto50
Figure 3.8	Solutions Optimales de la théorie des jeux52
Figure 4.1	Délai des paquets voix des différents modèles d'ordonnancements 159
Figure 4.2	Délai des paquets voix des différents modèles d'ordonnancements 259
Figure 4.3	Débit des différents modèles d'ordonnancements 160
Figure 4.4	Débit des différents modèles d'ordonnancements 2.....60

Figure 4.5	Délai des paquets voix du modèle d'ordonnancement WSUM 1.....	61
Figure 4.6	Délai des paquets voix du modèle d'ordonnancement WSUM 2.....	62
Figure 4.7	Débit du modèle d'ordonnancement WSUM 1	62
Figure 4.8	Débit du modèle d'ordonnancement WSUM 2	63

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

ABR	Associativity Based Routing
AC	Access Category
ACAR	Average Call Acceptance Ratio
ACF	Admission Control Failure
ACK	Acknowledgement
AIFS	Arbitration Inter Frame Space
AIMD	Additive Increase Multiplicative Decrease
AODV	Ad hoc On-Demand Distance-Vector
(Q)AP	(QoS) Access Point
AR	Admission Report
BE	Best Effort
CA	Collision Avoidance
CBR	Constant Bit Rate
CEDAR	Core Extraction Distributed Ad hoc Routing
CFP	Contention Free Period
CP	Contention Period
CSMA	Carrier Sense Multiple Access
CW	Contention Window
CW_{max}	Contention Window maximum
CW_{min}	Contention Window minimum
DBASE	Distributed Bandwidth Allocation/Sharing/Extension
DCF	Distributed Coordination Function
DIFS	Distributed Inter Frame Space
DMAC	Distributed Mobility Adaptive Clustering
DSDV	Destination Sequenced Distance-Vector Routing
DSR	Dynamic Source Routing
ECN	Explicit Congestion Notification

EDCF	Enhanced DCF
EQ	Enhanced QoS
FQMM	Flexible QoS Model for Mobile Ad hoc Networks
HC	Hybrid Coordination
HCF	Hybrid Coordination Function
HSR	Hierarchical State Routing
IEEE	Institute of Electrical and Electronics Engineers
IFS	Inter Frame Space
INSIGNIA	In-Band Signaling Support for QoS
IP	Internet Protocol
MAC	Medium Access Control
MANET	Mobile Ad hoc NETWORK
MPLS	Multiprotocol Label Switching
MSDU	MAC Service Data Unit
NAV	Network Allocation Vector
OLSR	Optimized Link State Routing
PF	Persistence Factor
PHY	Physical Layer
PIFS	PCF Inter Frame Space
QoS	Quality of Service
RES	Reservation
MRSVP	Mobile Resource Reservation Setup Protocol
PRTMAC	Proactive Real Time MAC
RTS/CTS	Request To Send/Clear To Send
SIFS	Short Inter Frame Space
SWAN	Stateless Wireless Ad hoc Networks
TC	Traffic Category
TCP	Transport Control Protocol
TDMA	Time Division Multiplexing Access
TORA	Temporally Ordered Routing Algorithm

TXOP	Transmission Opportunity
UDP	User Datagram Protocol
VBR	Variable Bit Rate
WLAN	Wireless Local Area Network
WRP	Wireless Routing Protocol

INTRODUCTION

Les réseaux mobiles Ad hoc se présentent comme la prochaine évolution des réseaux de communication sans fil en même temps qu'une alternative intéressante aux réseaux classiques dans le cas où la mise en œuvre de ces derniers s'avère trop coûteuse, temporaire ou tout simplement, impossible. Malgré leurs avantages et applications, ces réseaux obéissent à une logique d'organisation et de fonctionnement différente des réseaux classiques. Ils sont dépourvus de toute infrastructure de commutation fixe, laissant aux unités mobiles le soin de s'interconnecter, de s'organiser et d'établir la communication entre elles. Les réseaux ad hoc introduisent un certain nombre de nouvelles problématiques qui constituent autant de défis pour les chercheurs du domaine.

Le problème fondamental à résoudre dans ces réseaux est la question de la qualité de service (QoS) : comment offrir et assurer aux différents utilisateurs la QoS requise, dans un contexte de communication sans infrastructure où la position des nœuds et la qualité des liaisons sans fil varient de façon aléatoire? L'une des approches proposées dans la littérature pour atteindre cet objectif consiste à faire coopérer les différentes couches du modèle OSI (Open System Interconnection) : l'approche inter-couches (Cross Layer) qui semble une solution prometteuse.

Le nouveau concept, modèle de réception multi-usagers utilise plusieurs canaux. Plusieurs transmissions peuvent être autorisées entre les nœuds et le nœud peut avoir plusieurs réceptions à la fois. Cette technologie plus avancée peut de façon significative améliorer la performance d'utilisation de la bande passante. Les signaux multiples des différents mobiles peuvent être reçus en même temps, ce qui permet de réduire le délai d'un bout à l'autre de transfert de paquet. La réception multi usagers évite l'interférence entre les signaux reçus et cette caractéristique peut augmenter l'utilisation de la bande passante avec un grand facteur. Cette technologie permet d'éviter les collisions car elle peut résoudre le problème des nœuds cachés.

Ce document constitue notre recherche sur les mécanismes de support de la qualité de service dans les réseaux Ad hoc. Nous suggérons un modèle d'ordonnancement équitable distribué au niveau de la couche MAC pour améliorer l'équité et le débit.

Le réseau considéré est un réseau de nœuds capable de détecter simultanément plusieurs utilisateurs, 'la détection multi-usagers'.

Problématique

Les avancées dans le domaine des technologies de l'information et de la communication, mises de l'avant par la popularité de l'Internet, ont eu pour effet de multiplier et diversifier les trafics requérant une certaine forme de Qualité de Service (QoS). La prise en charge de ces trafics et le support de la QoS demandée devient donc un enjeu incontournable pour les applications à temps réel (applications multimédias). Toutefois, de nombreuses contraintes, intrinsèques au médium de communication sans fil ainsi qu'au paradigme Ad hoc, viennent complexifier cette prise en charge : le médium radio utilisé est peu fiable et la bande passante disponible est souvent limitée, la mobilité des nœuds du réseau rend la topologie entièrement dynamique, et finalement l'énergie disponible ainsi que la puissance de traitement et de stockage des unités mobiles est relativement faible. Enfin, la prise en charge de la qualité de service est un problème réparti sur l'ensemble du chemin qui relie la source à la destination.

Malgré les nombreuses approches recensées dans la littérature, force est de constater qu'aucune d'entre-elles ne résout complètement le problème du support de la qualité de service dans les réseaux mobiles Ad hoc.

Plusieurs des objectifs (flexibilité, adaptabilité de la QoS, mobilité) nécessaires au déploiement de ces réseaux sont souvent conflictuels et parfois même diamétralement opposés. Ce faisant, il appert que la meilleure solution pour offrir la QoS dans les réseaux Ad hoc sera celle qui sera en mesure d'allier ces différents objectifs d'une manière cohérente et efficace. C'est à la fois le principal problème et le défi le plus imposant auquel font face les chercheurs du domaine. De plus, comme il est remarquer dans plusieurs recherches, si les

unités composant le réseau sont trop mobiles, il est difficile, voire même impossible de garantir une quelconque forme de qualité de service.

Les principaux problèmes qui gênent la réalisation d'un vrai modèle de qualité de service dans les réseaux Ad hoc sont :

L'absence d'infrastructure et la topologie dynamique

Le réseau Ad hoc mobile est autonome et dynamique. En comparant avec les réseaux sans fil avec infrastructure, il n'y a pas la relation maître-esclave dans un réseau mobile Ad hoc. Les nœuds communiquent entre eux pour établir la connexion, donc chaque nœud agit comme un routeur. Par conséquent, dans un réseau mobile ad-hoc, un paquet peut circuler à partir d'une source à une destination, soit directement, soit par un certain ensemble des nœuds intermédiaires.

Bande passante limitée

L'interférence affecte considérablement l'utilité des réseaux Ad hoc en introduisant les multi-sauts et la bande passante devient très limitée.

L'énergie limitée

La mobilité de la communication est omniprésente et les nœuds mobiles sont munis des batteries d'énergies limitées. Ce contexte donne lieu à un modèle de réseau d'énergie limité dans laquelle l'énergie limite la communication dans un réseau plus large.

État d'information imprécise

L'état des liens change continuellement et l'état des flux change à chaque moment.

Medium utilisé

La majorité des technologies actuellement utilise un seul canal partagé entre les nœuds. Cette méthode d'accès au canal réduit l'efficacité du réseau Ad hoc, au lieu d'avoir plusieurs transmissions à la fois, il y a une seule qui est permise. Ce problème oblige à penser d'utiliser l'accès multicanaux.

Trafic en mode diffusion

Le trafic dans réseaux Ad hoc est en mode diffusion, (trafic distribué) donc tous les nœuds partagent la même responsabilité dans le réseau. Ce trafic dans les réseaux Ad hoc oblige à avoir une coordination entre les nœuds afin d'avoir une meilleur vision sur l'état du trafic.

Problème de l'équité

Certains nœuds peuvent avoir plus de priorité que d'autres.

Intérêt du projet

Dans les réseaux Ad hoc, la qualité de service relève plusieurs problématiques, elle doit subir différentes contraintes par flux, par lien et par nœud et les paramètres de QoS se diffèrent d'une application à autre. Par exemple, les applications multimédia ont des fortes contraintes en bande passante, délai et gigue. Les applications militaires ont des contraintes de sécurité. Les applications d'urgence et de secours requièrent la disponibilité continue du réseau (énergie). D'autres applications comme les communications d'un groupe dans une salle de conférence requièrent un minimum de consommation de l'énergie et de batterie.

Notre intérêt dans ce projet est de contribuer dans cette direction de recherche afin d'améliorer la QoS pour les applications multimédia, fournir plus de bande passante aux applications et réduire le délai de bout en bout.

Les architectures des plates formes

Actuellement dans la couche Mac dans les réseaux Ad hoc, il y a deux concepts qui sont utilisés, le premier se base sur un canal partagé par tous les nœuds dans la portée de transmission tel qu'IEEE 802.11 où une seule transmission peut être autorisée entre deux nœuds et les autres nœuds diffèrent leur transmission et le nœud peut avoir une seule réception à la fois. Le deuxième concept utilise plusieurs canaux. Plusieurs transmissions peuvent être autorisées entre les nœuds et le nœud peut avoir une seule réception à la fois. Le nouveau concept, modèle de réception multi-usagers proposé dans [2], utilise plusieurs

canaux. Plusieurs transmissions peuvent être autorisées entre les nœuds et le nœud peut avoir plusieurs réceptions à la fois. Cette technologie plus avancée peut de façon significative améliorer la performance d'utilisation de la bande passante. À savoir, que l'évolution technologique actuelle permet l'intégration la réception multi-usagers basée sur CDMA sur les puces de mobiles. Cette nouvelle technologie peut être implémentée pour le nouveau design de MAC dans les réseaux Ad hoc.

Il y a trois avantages principaux pour l'utilisation de la réception multi-usagers dans les réseaux Ad Hoc. Premièrement, les signaux multiples des différents mobiles peuvent être reçus en même temps, ce qui permet de réduire significativement le délai d'un bout à l'autre de transfert d'un paquet. Deuxièmement, la réception multi-usagers évite l'interférence entre les signaux reçus et cette caractéristique peut augmenter l'utilisation de la bande passante fréquence avec un grand facteur [3, 4,5]. Troisièmement, cette technologie permet d'éviter les collisions et peut résoudre le problème des nœuds cachés. Bien que la réception multi-usagers soit connue depuis longtemps, la plupart des études se concentrent sur la couche physique [5]. À notre connaissance, il n'y a aucune étude publiée sur l'application de cette technologie dans les réseaux Ad Hoc.

Le dernier concept (modèle de transmissions et réceptions multiples) utilise plusieurs canaux. Plusieurs transmissions peuvent être autorisées entre les nœuds et le nœud peut avoir plusieurs transmissions ou plusieurs réceptions à la fois. Cette technologie plus avancée peut de façon significative améliorer la performance de l'utilisation de la bande passante et aussi régler le problème de submersion de réceptions qui peut affecter les nœuds dans le modèle avec réception multi-usagers à cause de l'asymétrie entre les transmissions et les réceptions. Cette nouvelle technologie sera la dernière plate forme à étudier afin d'implémenter nos modèles qu'on propose.

Notre mémoire consiste à élaborer des mécanismes efficaces et évolutifs de gestion de la qualité de service dans les réseaux Ad hoc utilisant la détection multi-usagers.

De manière plus spécifique, nos objectifs dans ce projet visent à concevoir et implémenter un ordonnancement équitable dans la couche MAC avec un nouveau concept ‘détection multi-usagers’ pour améliorer l’équité et débit.

Le but de ces modèles d’ordonnancement est d’améliorer l’équité et le débit mais on remarque qu’il y a une compétition entre ces deux paramètres, l’amélioration d’un paramètre peut dégrader la performance de l’autre. L’utilisation du théorème des jeux coopératifs (Game theory) [6] permet de combiner les deux paramètres en même temps pour obtenir une solution optimale.

Parmi ces solutions, on a trois solutions principales :

1. La solution optimale est celle qui permet d’égaleriser les fonctions de préférence de chaque joueur (la préférence est l’utilité u). Cette fonction est associée avec la solution de Raiffa [6].
2. La solution optimale est celle qui permet de maximiser la somme de fonctions de préférence de chaque joueur. Cette fonction est associée avec la solution de Thomson [6].
3. La solution optimale est celle qui permet de maximiser le produit de fonctions de préférence de chaque joueur. Cette fonction est associée avec la solution de Nash [6].

Cette partie sera détaillée dans le chapitre 3.

CHAPITRE 1

LES RÉSEAUX SANS FIL AD HOC

1.1 Introduction

Pour que les réseaux mobiles ad hoc puissent espérer prendre une place prépondérante dans l'environnement ubiquitaire qui se dessine à l'horizon, ils devront être en mesure d'offrir des mécanismes supportant des trafics hétérogènes et nécessitant des requis de services des plus diversifiés. À l'instar des réseaux filaires ou même sans fil traditionnels, les trafics que les réseaux ad hoc (Figure 1.1) devront supporter varient du simple courriel qui ne nécessite aucun requis de service particulier aux applications temps réel comme la vidéoconférence qui sont très sensibles à toute variation de délai de transmission et dont le transport nécessite une bande passante relativement élevée. Pour supporter des conditions de trafic si variées, il est nécessaire d'offrir différents niveaux de qualité de service, chacun de ces niveaux répondant à un requis de service particulier.

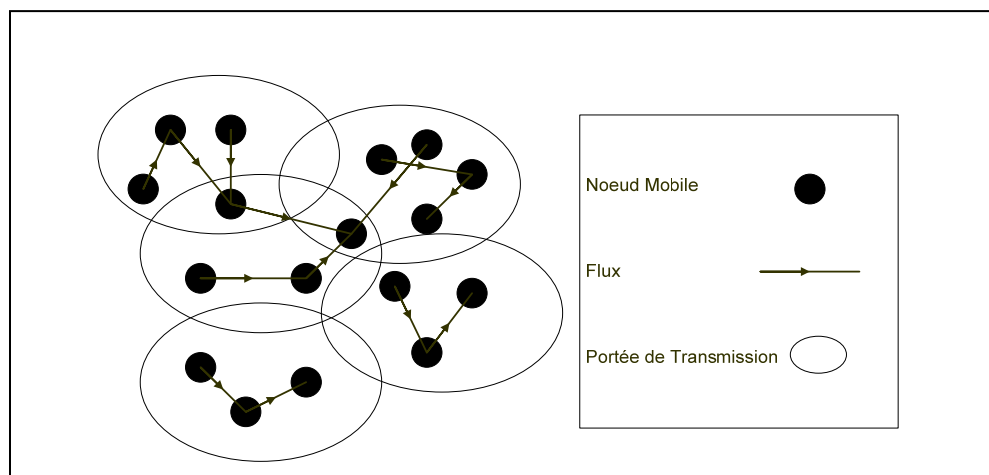


Figure 1.1 Réseau sans fil Ad hoc.

Tiré de Sarkar et Ba (2008)

Nous présentons dans cette partie d'abord un résumé sur la fonctionnalité des protocoles de la couche MAC.

1.2 La fonctionnalité de la couche MAC dans les réseaux Ad hoc

Un problème crucial dans les réseaux ad hoc est le partage de la bande radio. Si les réseaux sans fil (WLAN Wireless LAN) proposent des solutions, ces dernières ne sont pas complètement adaptées aux réseaux ad hoc. Il convient de mettre en place de nouvelles techniques qui tiennent d'abord compte du caractère naturellement ouvert d'un réseau ad hoc. Ensuite, il convient d'utiliser au mieux la bande radio partagée ; pour ce faire, il convient de trouver le meilleur compromis entre la réutilisation spatiale et le problème des nœuds cachés. Les réseaux ad hoc qui sont actuellement expérimentés utilisent souvent des interfaces radio 802.11 [7] et donc un protocole CSMA pour partager la bande radio.

Par ailleurs, les réseaux ad hoc comme les autres réseaux doivent parfois être en mesure d'offrir des services ou fonctionnalités qui dépassent la simple fourniture de connectivité. Ces services peuvent être requis par des applications particulières, par exemple si celles-ci nécessitent de la qualité de service. On va décrire au dessous les différentes méthodes d'accès au canal.

1.2.1 Les méthodes d'accès au canal

Un seul canal partagé

Actuellement, la plupart des technologies sans fil proposées (et aussi utilisées dans les réseaux Ad hoc) sont en accès partagé de type CSMA/CA (Carrier Sense Multiple Access/Collision Avoidance), ce qui impose qu'un nœud ne puisse simultanément émettre et recevoir, ni recevoir lorsque ces voisins émettent. Il résulte de ce schéma, que le débit moyen pour la transmission d'un seul flux est divisé. Cette inefficacité ne peut être réduite car le modem radio travaille sur un seul canal.

Le protocole CSMA/CA

Le protocole CSMA (Carrier Sensing Multiple Access) [7] (CSMA / CA) est un mécanisme utilisé par la couche MAC dans la norme IEEE 802.11 dans les réseaux sans fil. Le (CSMA / CD) est une technique bien étudiée dans les normes 802.x des réseaux filaires. Cette technique ne peut pas être utilisée dans le cadre de réseaux sans fil parce que le taux d'erreur est très élevé et les collisions dégradent le débit. En outre, la détection des collisions dans le réseau sans fil n'est pas toujours possible. La technique adoptée ici est donc d'éviter les collisions.

Un nœud qui souhaite transmettre des données doit d'abord écouter le canal pour un temps prédéterminé pour vérifier si un autre nœud est en transmission sur le canal. Si le canal est détecté "libre", le nœud est autorisé à entamer le processus de transmission. Si le canal est ressenti «occupé», le nœud ajourne sa transmission pour une période de temps aléatoire.

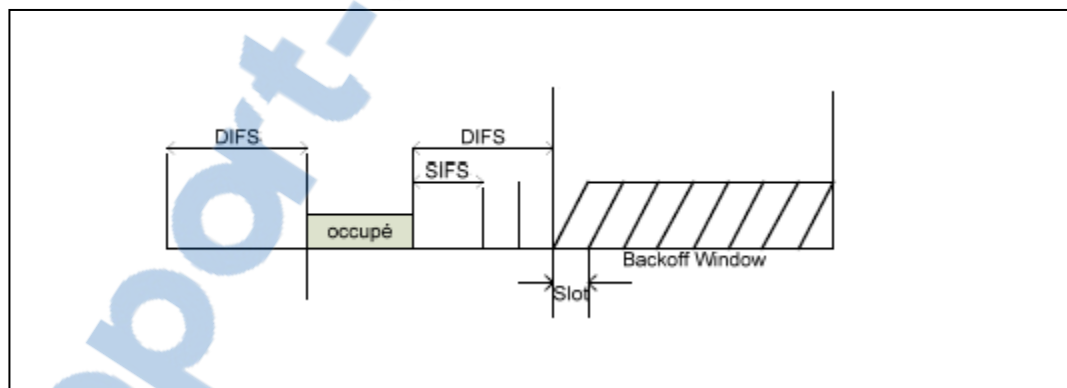


Figure 1.2 Accès au canal avec CSMA/CA.

Le mécanisme RTS/CTS

RTS / CTS (Request to Send / Clear to Send) [7] est un mécanisme utilisé pour réduire les collisions des trames produites par le problème du terminal caché (Figure 1.3).

Un nœud qui souhaite envoyer des données envoie une requête d'envoi de trame (RTS). Le nœud de destination répond avec (CTS). Toute autre nœud recevant le RTS ou CTS devraient s'abstenir d'envoyer des données pendant un temps donné (pour résoudre le problème du nœud caché). Le temps ou le nœud doit attendre avant d'essayer l'accès au médium est inclus dans les RTS et CTS.

En règle générale, l'envoi de RTS / CTS se produit lorsque la taille du paquet dépasse ce seuil. Si la taille du paquet que le nœud veut transmettre est plus grande que le seuil, le RTS / CTS est déclenchée. Sinon, la trame de données est envoyée immédiatement.

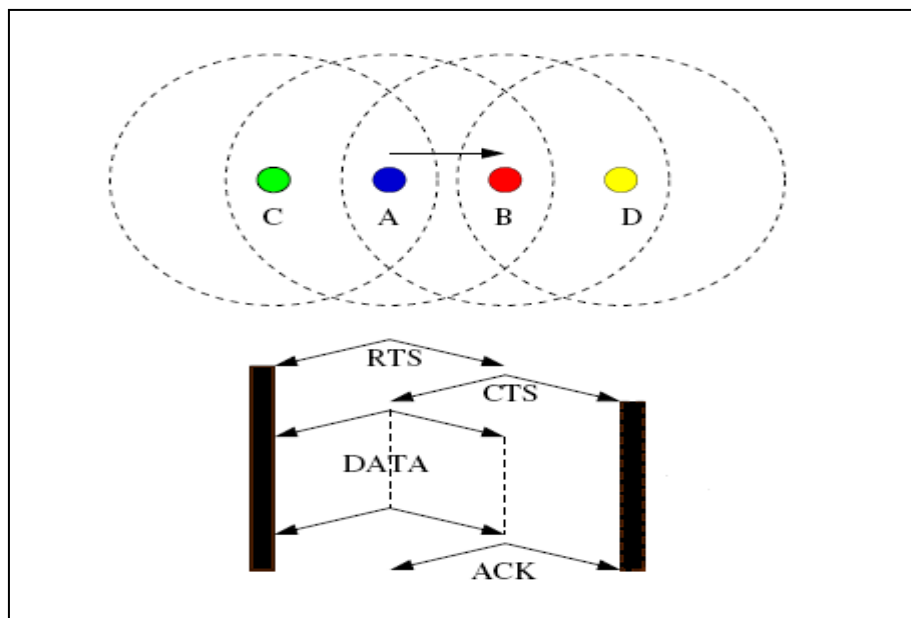


Figure 1.3 L'échange RTS/CTS.

1.2.2 Plusieurs canaux (Muti-canaux)

La méthode d'accès multiple permet à plusieurs nœuds connectés sur le même support de transmission multi-point de transmettre sur ce support et de partager sa capacité.

La méthode d'accès multiple est basée sur une technique de multiplexage, qui permet à plusieurs flux de données ou des signaux de partager le même canal. Le multiplexage est fourni par la couche physique.

La méthode d'accès multiple est également basée sur un protocole d'accès multiples et mécanisme de contrôle d'accès au média (MAC). Ce protocole traite l'adressage, l'affectation des canaux à différents utilisateurs, et d'éviter les collisions.

Cette méthode permettrait de réduire l'inefficacité du seul canal partagé dans les réseaux Ad hoc.

Il y a trois types d'accès, Code Division Multiple Accès (CDMA), Frequency Division Multiple Accès (FDMA), Time Division Multiple Accès (TDMA).

CDMA: Code Division Multiple Accès

Le principe de base du CDMA [8] est que chaque station se voit attribuer un code unique utilisé pour distinguer entre les messages. En CDMA, ni la bande de fréquence et le temps sont utilisés pour distinguer une station. Plutôt, les stations se distinguent par le code qui leur était assigné. CDMA permet l'utilisation simultanée l'espace des fréquences et du temps par tous les signaux et c'est le choix préféré d'aujourd'hui pour une interface air à une station sans fil.

En CDMA, chaque station émet dans tout le spectre attribué. Ainsi, de nombreux signaux différents peuvent occuper l'espace de la même fréquence en même temps. CDMA et la

propagation spectre sont intimement liés. En fait, CDMA peut être considéré comme juste une autre façon de voir l'étalement de spectre.

CDMA transmet sur la plage de fréquences disponibles. Il ne peut pas attribuer une fréquence spécifique à chaque utilisateur sur le réseau de communication. Cette méthode, appelée multiplexage. Cela permet à plusieurs utilisateurs de communiquer sur le même réseau à un moment donné et chaque utilisateur a été attribué à des fréquences spécifiques (Figure 1.4).

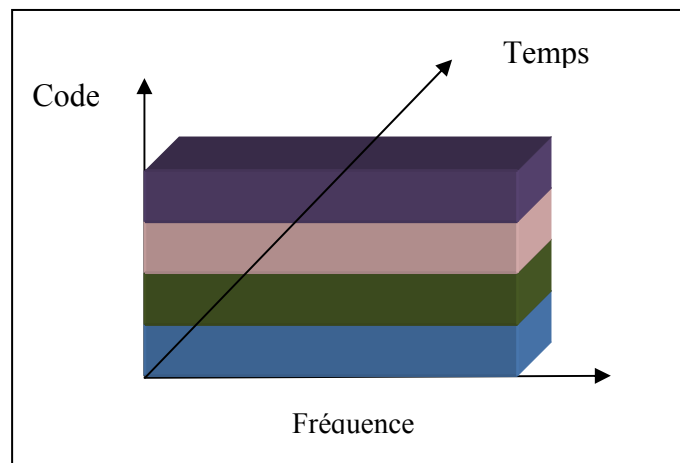


Figure 1.4 CDMA.

FDMA : Frequency Division Multiple Accès

Le principe de base de FDMA [8] est que chaque nœud est attribué à une bande de fréquence unique utilisée par ce nœud pour communiquer. Dans l'approche FDMA, chaque signal utilise une petite partie de la bande passante du canal partagé.

FDMA prend plusieurs sources de données et les déplace de sorte qu'ils sont adjacents dans la fréquence. En combinant les largeurs de bande de chaque canal et en les additionnant, vous pouvez obtenir une idée approximative de la bande passante totale du canal partagé. Ce canal est utilisé pour envoyer l'ensemble du groupe des données multiplexées (Figure 1.5).

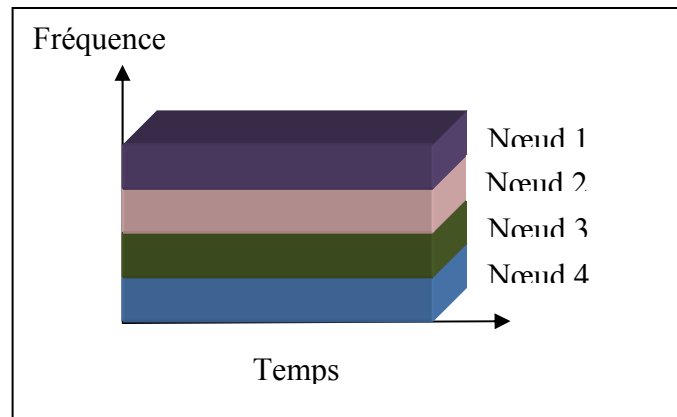


Figure 1.5 FDMA.

TDMA : Time Division Multiple Accès

Le principe de base TDMA [8] est que chaque nœud est attribué à un intervalle de temps unique dans lequel le nœud communique. Dans l'approche TDMA, chaque signal utilise la totalité de la bande passante pour de courtes périodes de temps.

TDMA prend plusieurs sources de données et les déplace de telle sorte qu'ils sont adjacentes dans le temps. En combinant la durée de chaque canal et les ajoutant ensemble, vous pouvez obtenir une idée approximative du temps total du canal partagé. Ce canal est utilisé pour envoyer l'ensemble du groupe des données multiplexées (Figure 1.6)

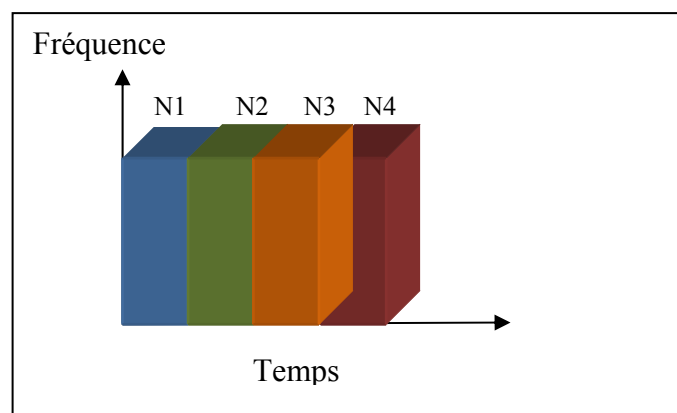


Figure 1.6 TDMA.

CHAPITRE 2

LES MECANISMES DE SUPPORT DE LA QUALITÉ DE SERVICE DANS LES RÉSEAUX AD HOC

2.1 Introduction

Pour que les réseaux mobiles Ad hoc puissent espérer prendre une place prépondérante dans l'environnement ubiquitaire qui se dessine à l'horizon, ils devront être en mesure d'offrir des mécanismes supportant des trafics hétérogènes et nécessitant des requis de services des plus diversifiés. À l'instar des réseaux filaires ou même sans fil traditionnels, les trafics que les réseaux ad hoc devront supporter varient du simple courriel qui ne nécessite aucun requis de service particulier aux applications temps réel comme la vidéoconférence qui sont très sensibles à toute variation de délai de transmission et dont le transport nécessite une bande passante relativement élevée. Pour supporter des conditions de trafic si variées, il est nécessaire d'offrir différents niveaux de qualité de service, chacun de ces niveaux répondant à un requis de service particulier.

2.2 Classification des méthodes supportant la QoS dans Ad hoc

La majorité des travaux reliés aux réseaux mobiles ad hoc classent les méthodes ou mécanismes de qualité de service en trois grandes catégories :

- Les protocoles d'accès au médium (couche MAC) qui cherchent à ajouter des fonctionnalités aux couches basses de la pile de protocole afin de pouvoir supporter divers requis de service.
- Les protocoles de routage (couche réseau) supportant la qualité de service qui recherchent des routes ayant suffisamment de ressources disponibles pour satisfaire certaines contraintes de service.
- Les protocoles de signalisation offrant des mécanismes de réservation de ressources indépendamment du protocole de routage utilisé.

Généralement, ces composantes collaborent afin d'accomplir les objectifs spécifiques qui sont définis dans un modèle de qualité de service. Ces différentes catégories feront respectivement l'objet des sections qui suivent.

2.3 Classification des approches de QoS

Plusieurs critères sont utilisés pour classer les approches, en se basant soit sur l'interaction entre le protocole de routage et le mécanisme de QoS, soit sur l'interaction entre le niveau 2 et 3, soit le mécanisme de mise à jour des informations de routage (Figure 2.1).

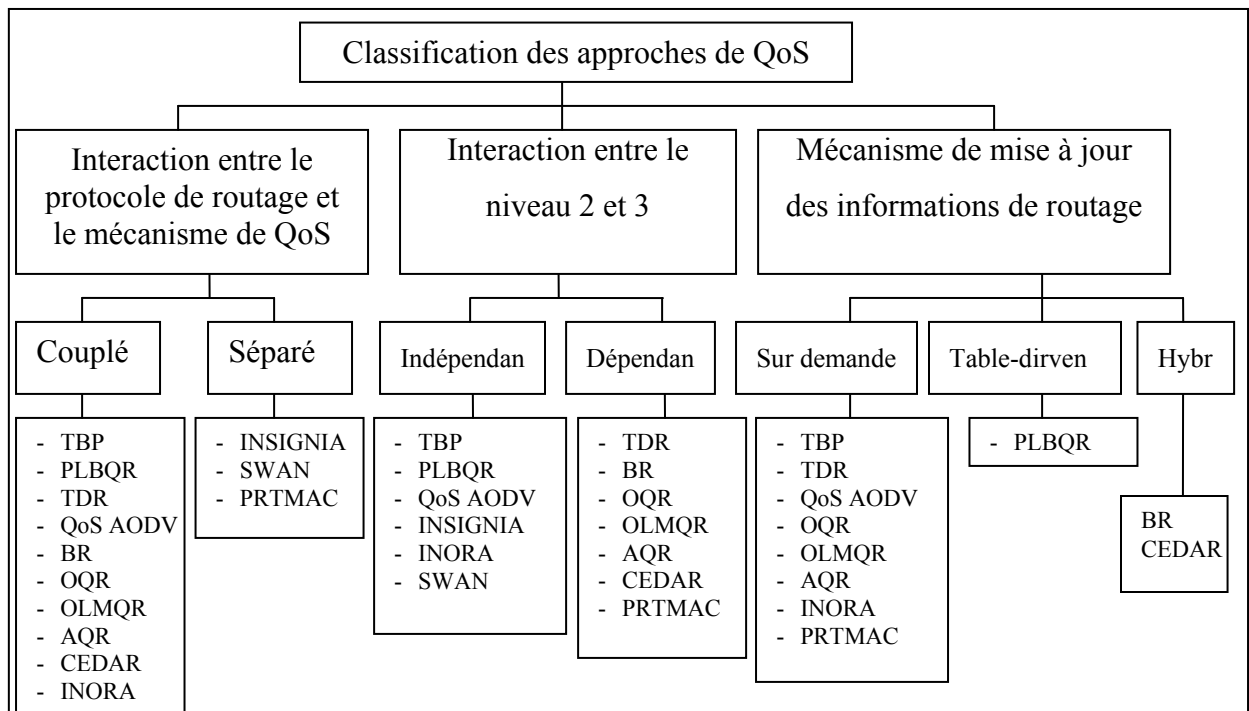


Figure 2.1 Classification des approches de QoS.

Tiré de Sarkar et Ba (2008)

Pour l'interaction entre le protocole de routage et le mécanisme de QoS, les approches de QoS peuvent être classifiées en deux catégories : couplé et séparé. Dans le cas couplé, le protocole de routage et le mécanisme de QoS interagissent entre eux pour délivrer une QoS garantie, si le protocole de routage change alors la QoS se dégrade. Dans le cas séparé, le mécanisme de QoS ne dépend pas d'un protocole de routage spécifique.

Pour l'interaction entre les niveaux 2 (MAC) et 3 (réseau), les approches de QoS peuvent être classifiées en deux catégories : indépendant et dépendant. Dans le cas indépendant, le niveau 3 ne dépend pas du niveau MAC. Dans le cas dépendant, l'approche nécessite l'assistance de la couche MAC pour fournir la QoS au protocole de routage.

Pour le mécanisme de mise à jour des informations de routage, les approches peuvent être classifiées en trois catégories : Table-driven, sur demande et les approches hybrides. Dans le cas de l'approche Table-driven, chaque nœud maintient une table de routage qui aide à la transmission des paquets. Dans l'approche sur demande, il n'y a pas de tables qui sont maintenues dans les nœuds, par contre le nœud source demande la création du chemin. Enfin, les approches hybrides couplent les deux approches Table-driven et sur demande.

2.4 La classification basée sur les couches (MAC & Réseau)

Les solutions de QoS existantes peuvent aussi être classifiées en se basant sur la couche où elles opèrent (Figure 2.2). Quelques solutions sont expliquées dans les sections suivantes.

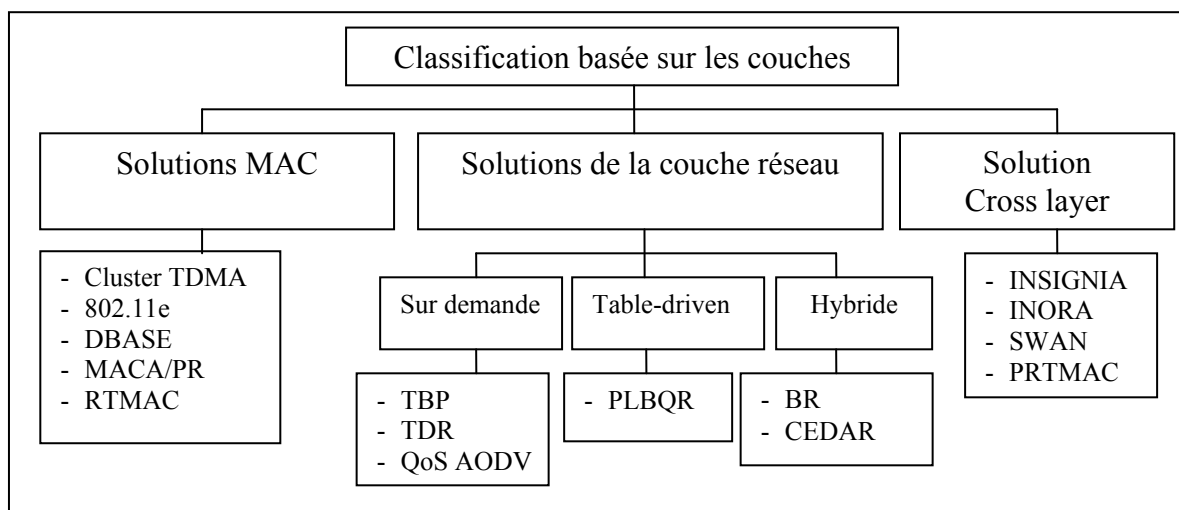


Figure 2.2 Classification basées sur les couches.

Tiré de Sarkar et Ba (2008)

2.5 La QoS dans la couche MAC

Plusieurs protocoles au niveau de la couche MAC ont été proposés pour les réseaux sans fil ayant pour but de fournir la garantie de la QoS pour un temps réel. Portant la majorité des protocoles de contrôle d'accès conçus pour des réseaux sans fil utilisent une composante aléatoire (généralement, pour la résolution de conflits d'accès).

Le principal objectif de ces protocoles, dans un contexte de qualité de service, est de réduire le plus possible les délais d'accès pour les flux temps réel. Parmi ces protocoles on trouve MACA/PR (Multiple Access Collision Avoidance with Piggyback Reservation), 802.11, et RTMAC.

2.5.1 MACA/PR

Multi Access Collision Avoidance with Piggyback Reservation (MACA / PR) [7] est un protocole MAC pour les réseaux Ad hoc. MACA / PR permet une transmission rapide et fiable des données non temps réel afin de garantir une bande passante pour le trafic en temps réel.

Pour la transmission des données non temps réel avec MACA / PR, le nœud avec un paquet à envoyer doit d'abord attendre une «fenêtre» libre dans le tableau de réservation (RT), qui enregistre tous les paquets à envoyer et recevoir. Il attend alors pendant un temps aléatoire supplémentaire. S'il sent que le canal est libre, il procède à une demande d'envoi par RTS, CTS et ACK. Si le canal est occupé, il attend que le canal devient inactif et répète la procédure ci-dessus.

Pour la transmission de paquets en temps réel, le comportement de la MACA / PR est différente. Pour transmettre le premier paquet de données d'une connexion en temps réel, l'expéditeur S lance RTS-CTS, puis avec ACK si le CTS est reçu. Notez que si l'expéditeur ne parvient pas à recevoir les accusés de réception de plusieurs émetteurs, il redémarre la

connexion avec le dialogue RTS-CTS à nouveau. MACA / PR ne retransmettre pas en temps réel des paquets après la collision.

Pour réserver une bande passante pour le trafic en temps réel, les informations de planification en temps réel est transporté dans les en-têtes des paquets et ACK. Ainsi, à travers les informations de réservation (piggybacked), la bande passante est réservée pour garanti le trafic en temps réel.

2.5.3 Le protocole RTMAC (Real-Time medium access control protocol)

RTMAC [9] se compose en deux parties :

1. Son protocole MAC est une extension en temps réel d'IEEE 802.11-DCF, il utilise deux mécanismes :
 - MAC pour le trafic 'best effort',
 - Protocole de réservation pour la circulation en temps réel.
2. Utilise le protocole de routage avec QoS fournissant :
 - Une réservation de bout en bout,
 - Une libération de ressources.

Différentes techniques sont utilisées pour fournir le trafic 'best effort' en temps réel :

- Les messages de contrôle ResvRTS, ResvCTS et ResvACK sont utilisés pour le trafic en temps réel.
- Les messages de contrôle RTS/CTS/ACK sont utilisés pour le trafic 'best effort'.

Pour donner la priorité aux paquets en temps réel, le temps d'attente pour transmettre les paquets ResvRTS est réduit deux fois.

2.6 Les mécanismes de contrôle de la QoS dans la couche MAC

Pour satisfaire les besoins des applications temps-réel dans les réseaux filaires ou sans fil, une première solution suggère d'allouer plus de bande passante aux flux temps-réel afin de résoudre les problèmes de congestion, perte de paquets et de délai puisque la bande passante devient disponible à faible coût. Cependant, plusieurs applications temps-réel telles que la Voix sur IP, la vidéo conférence, etc. ont besoin de garanties particulières pour le délai, gigue et taux de perte et non seulement la bande passante. La proposition d'allouer des ressources supplémentaires n'est pas un choix judicieux pour assurer la QoS d'une part, ce choix conduit à un surdimensionnement du réseau réduisant ainsi son taux d'utilisation en temps normal. D'autre part, avec l'émergence des applications multimédia et l'utilisation croissante des réseaux à commutation de paquets tels que l'internet, la saturation rapide de la totalité des ressources du réseau est prévue ce qui ramène de nouveau à la situation de ressources limitées.

La deuxième solution consiste à fournir les mécanismes nécessaires pour assurer la QoS qui se résument principalement en trois fonctionnalités supplémentaires aux services fournis par le réseau :

1. Ordonnancement (Scheduling).
2. La gestion des tampons (gestion des files d'attente).
3. Les mécanismes de régulation de trafic (shaping et policing).

Le médium radio utilisé dans les réseaux sans fil est peu fiable et la bande passante disponible est souvent limitée, donc l'équité est une très bonne issue pour résoudre le problème d'accès au canal partagé.

Avec l'ordonnancement équitable (Fair Scheduling) dans la couche MAC, les différents flots qui partagent le même canal peuvent utiliser la bande passante de façon proportionnelle en fonction de leur poids, ce qu'on va l'expliquer par la suite. Dans notre recherche, on se concentre sur le mécanisme Fair Scheduling afin de fournir la QoS dans la couche MAC.

Notre proposition se base sur la conception et l'implémentation d'un algorithme d'ordonnancement équitable distribué.

2.6.1 L'ordonnancement équitable (Fair Scheduling)

Ces mécanismes d'ordonnancement équitables se basent sur les algorithmes de Fair Queuing qui sont les meilleures techniques pour partager les ressources (la bande passante) de façon équitable dans le but de satisfaire les exigences de la QoS.

2.6.2 Le partage équitable (Fair Queuing (FQ))

La discipline de service « à partage équitable » (Fair Queuing) est proposée [10] en 1987 et représente le travail initiateur pour le développement des techniques d'ordonnancement à partage équitable la bande passante. L'objectif est de fournir équitablement le même taux de service aux différents flux partageant les ressources.

Fair Queuing est un algorithme d'ordonnancement qui permet aux paquets de flux multiples de partager équitablement la capacité du lien. L'avantage par rapport au FIFO est que les paquets de données ne peuvent pas prendre plus que leur part de la capacité du lien.

Pour résoudre les problèmes du FQ, les auteurs ont proposé dans [11] un service en tourniquet bit-à-bit, qui est un modèle théorique, pour le partage équitable de la bande passante. De plus, la notion du poids a été introduite pour pondérer le service proportionnellement à la bande passante exigée par le flux. Cette proposition est appelée Weighted Fair Queueing (WFQ). Cette discipline de service est connue aussi sous le nom de GPS (Generalized Processor Sharing) et PGPS qui est la version « paquet » de GPS et qui correspond exactement à WFQ définie dans [12].

2.6.3 Le Service Fluide GPS (Generalized Processor Sharing)

Le GPS [12] est un modèle théorique de multiplexage de flux qui se charge de transmettre les paquets bit par bit en fonction de leurs poids associés appelés aussi taux de service. La Figure 2.3 montre le déroulement de la discipline de service GPS.

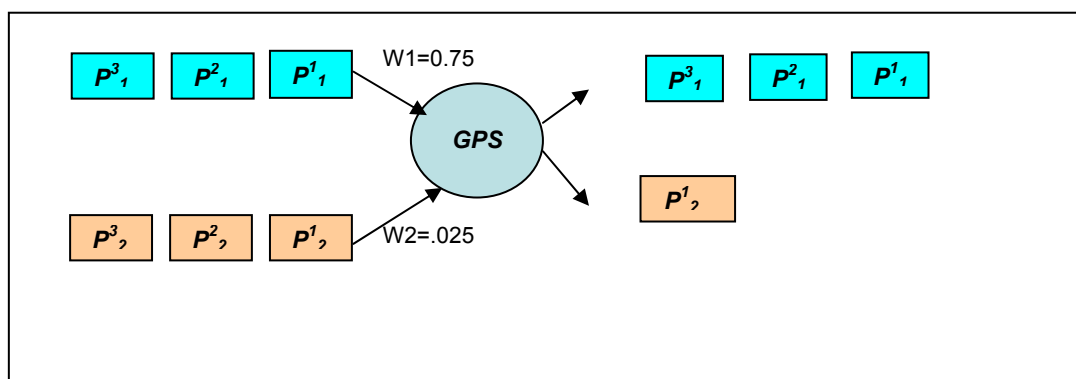


Figure 2.3 La discipline de service GPS.

GPS est une discipline fluide non implantable en réalité qui suppose que le serveur peut servir simultanément plusieurs flux. Dans un système réaliste qui considère des paquets, un seul flux peut être servi à un instant donné et un paquet doit être servi entièrement avant qu'un autre paquet puisse l'être. Il existe plusieurs techniques qui permettent d'émuler la discipline « idéalisée » GPS. L'algorithme WFQ (ou PGPS) est la version « paquet » la plus répandue.

2.6.4 Le PGPS ou WFQ

Le système WFQ est une émulation du système GPS proposé les auteurs dans [12]. Le serveur WFQ consiste à servir les paquets dans l'ordre croissant de leurs instants de fin de service dans le système GPS correspondant. Ces instants sont appelés temps virtuel de fin de service.



L'idée de l'algorithme WFQ est de calculer pour chaque paquet le moment où il serait terminé afin d'appliquer l'algorithme GPS. Ensuite, WFQ ordonnance les paquets dans l'ordre croissant de leur temps de fin.

$$F(i, k, t) = \max \{F(i, k-1, t), R(t)\} + P(i, k, t)/W(i) \quad (2.1)$$

avec :

- $F(i, k, t)$: Le temps de fin de service du $k^{\text{ième}}$ paquet du $i^{\text{ième}}$ flux.
- $P(i, k, t)$: Le temps d'arrivée du $k^{\text{ième}}$ paquet du $i^{\text{ième}}$ flux.
- $W(i)$: Le poids du $i^{\text{ième}}$ flux.
- $R(t)$: le numéro de série à l'instant t .

Ensuite, WFQ ordonnance les paquets dans l'ordre croissant de leur temps de fin.

2.6.5 Self Clocked Fair Queuing (SCFQ)

SCFQ [13] est une autre approximation de GPS proposée dans [11] qui résout le problème de la complexité de calcul du temps virtuel $V(t)$ de WFQ en le remplaçant par le temps de fin de service du paquet en cours indépendamment du flux auquel il appartient. Ainsi, Le temps virtuel de fin de service est :

$$F(i, k, t) = \max \{F(i, k-1, t), CF\} + P(i, k, t)/W(i) \quad (2.2)$$

avec CF est le temps de fin de service du paquet en cours de service.

Les garanties temporelles fournies par cette approximation sont moins bonnes que celles obtenues par WFQ. Ceci est dû à une déviation plus importante par rapport à GPS en termes d'estimation du temps virtuel.

2.6.6 Start-Time Fair Queuing (STFQ)

STFQ [14] est une autre forme d'approximation de GPS proposée pour résoudre le problème de la complexité de calcul du temps virtuel et émuler correctement GPS dans le cas d'un serveur (lien de sortie) à capacité variable. En effet, la discipline WFQ suppose que la capacité du serveur C est constante. Sur deux exemples, les auteurs montrent que WFQ n'émule pas parfaitement GPS quand la capacité du serveur est variable.

Avec STFQ, deux estampilles temporelles sont associées à chaque paquet : le temps virtuel de début de service et le temps virtuel de fin de service. Contrairement à WFQ et SCFQ, le service des paquets avec STFQ se fait dans l'ordre croissant des temps virtuels de début de service. De plus, le temps virtuel $V(t)$ est défini comme étant le temps virtuel du début de service du paquet en cours. Le temps virtuel de début de service est :

$$S_i^k = \max\{F_i^{k-1}, V(t_i^k)\} \quad \text{Avec } k > 1 \quad (2.3)$$

Avec, F_i^k le temps virtuel de fin de service, définie par l'équation (2.1).

Il est remarquable que la complexité de calcul soit réduite comparée à celle de WFQ puisqu'elle ne nécessite que la vérification de l'estampille temporelle de début de service du paquet courant. STFQ et SCFQ sont de même ordre de complexité.

2.6.7 Fair queueing distribué dans les réseaux Ad hoc

Les algorithmes de Fair Queuing [11, 12, 13, 14] étaient développés pour les réseaux filaires (câblés) donc ils ne peuvent pas être appliqués directement dans les réseaux sans fil, centralisés avec station de base ou distribués comme les réseaux Ad hoc. Au premier lieu, les travaux sur le Wireless Fair Queuing étaient concentrés sur les réseaux sans fil centralisé. La plupart de ces travaux, IWFQ[15], CIF-Q[16], SBFA [17] sont des modifications des algorithmes déjà développés dans les réseaux filaires et adaptés au principe des réseaux sans fil. Le principe de ces algorithmes est d'implémenter un algorithme de fair Queuing au niveau de la station de base. Cette station a le rôle de gérer le trafic de façon équitable.

Les réseaux Ad hoc sont des réseaux distribués où tous les nœuds ont le même rôle et la même responsabilité ce qui implique que les algorithmes de fair Queuing développés pour les réseaux sans fil centralisés ne sont pas adaptés pour les réseaux sans fil distribués.

Plusieurs recherches récentes ont été réalisées afin de fournir un algorithme de Fair Queuing distribués dans les réseaux Ad hoc. Dans [18] l'idée principale de ce mécanisme est de sélectionner un intervalle de backoff proportionnel à une étiquette (Tag) de fin de flux. L'étiquette de fin est le rapport entre la longueur de paquet et le poids d'une trame ou $B_i = [\text{Scaling_Factor} * L_i / \Omega_i] * \text{random number}$. B_i est le backoff, L_i est la longueur de paquet, Ω_i est le poids et le random est une variable aléatoire uniformément distribuée dans $[0.9, 1.1]$. Ce nombre aléatoire dont la moyenne est présenté pour empêcher une collision entre deux paquets ou plus. Avec la combinaison de la longueur de poids de paquet dans le calcul de backoff, le trafic avec différentes classes de sortie peuvent être traités différemment. Ce mécanisme est basé sur l'algorithme (SCFQ). Les auteurs eux-mêmes notent dans leurs simulations que les sorties des flux sont tout à fait sensibles au choix des longueurs et des poids des trames.

Dans [19, 20, 21] par exemple, les auteurs ont développé les mécanismes qui réalisent l'équité en essayant de maximiser la sortie des flux. Ils ont proposé une approche intéressante basée sur la notion d'un graphique de contention des flux qui tient compte de la topologie du réseau. Ils ont également développé un modèle de topologie indépendant pour le fair Queuing [22].

Le but de ce projet est de réaliser un Fair Queuing distribué dans les réseaux sans fil Ad hoc avec un canal partagé d'une façon efficace. Les protocoles proposés dans les articles [19, 20, 21, 22] sont basés sur cinq idées essentielles :

- Chaque flot a une étiquette, un flot est associé à une table $[f_i, T_i]$.
- Cette étiquette est calculée grâce à l'algorithme STFQ (Start time Fair Queuing).
- Le flot avec la plus petite étiquette a la priorité d'être servi le premier.
- Chaque flot a un backoff (une fonction simple qui calcul ce backoff).

- L'échange des informations des flots voisins (l'étiquette de chaque flot) se fait grâce aux messages du protocole CSMA/CA, RTS/CTS/DATA/ACK.

2.7 Les approches de QoS proposé avec la plate forme 'détection multi-usagers'

Bien que la réception multi usagers soit connue depuis longtemps, la plupart des études se concentrent sur la couche physique. Selon notre connaissance, il n'y a aucune étude publiée sur l'application de cette technologie dans les réseaux Ad Hoc sauf l'article [2] du groupe de recherche du laboratoire LAGRIT.

Le but de cette recherche est d'implémenter un nouveau design de la couche MAC dans les réseaux Ad hoc avec le concept 'détection multi-usagers' (modèle 3 de la Figure 2.4) et les résultats obtenus dans le chapitre 4 montrent une réduction significative du délai d'un bout à l'autre, augmentation de la bande passante avec un facteur très important et le taux de perte des paquets est très bas avec ce mode 'détection multi-usagers'.

En conclusion cette technologie peut fournir une grande amélioration à la qualité de service dans les réseaux Ad hoc.

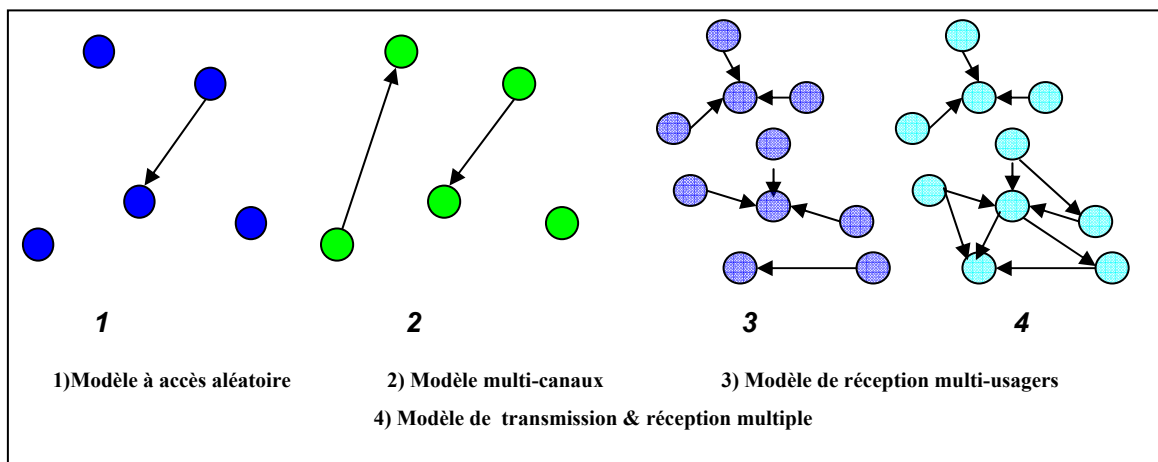


Figure 2.4 Les modèles d'accès.

CHAPITRE 3

PROPOSITION D'UN ORDONNANCEMENT ÉQUITABLE AVEC DÉTECTION MULTI-USAGERS

3.1 Introduction

Le processus d'élaboration de notre projet s'appuie sur, un modèle d'ordonnancement équitable au niveau la couche MAC, un modèle d'ordonnancement avec priorité pour le trafic avec multi-sauts pour améliorer la QoS dans le routage (QoS Routing) et la conception d'une interface qui devra permettre une coopération entre la couche d'accès au média MAC (Medium Access Control) et couche réseau.

3.2 Modèle de ordonnancement équitable dans la couche MAC

Dans cette section, on décrit la méthodologie que nous avons suivi afin de bâtir un modèle de ordonnancement équitable (fair scheduling) dans la couche MAC.

Au début, on donne une description du modèle utilisé. Par la suite, on décrit les trois fonctions pour bâtir cet ordonnancement. Puis on donne une formulation pour combiner l'équité et le débit avec la théorie des jeux.

3.3 Description du modèle

Dans la partie état de l'art on a montré que les recherches se focalisent pour implémenter un algorithme de fair queuing dans les réseaux Ad hoc. Le premier axe se base sur la modification de l'algorithme de backoff [23]. Le deuxième axe se base sur la notion du graphe de dépendance de flots [24] qui tient compte de la topologie du réseau (graphe de nœud).

On a trouvé que l'approche qui se base sur la notion du graphe de dépendance de flots donne une meilleure vision sur l'état de trafic et sur la dépendance entre les flots et permet aussi l'échange d'informations entre les nœuds et cet échange est indispensable dans les réseaux

Ad hoc où le trafic est en mode de diffusion. A partir de ce graphe de dépendance de flots, on peut bâtir et élaborer un algorithme d'ordonnancement.

Dans cette section, nous décrivons d'abord le modèle, la matrice de dépendance qui contient les dépendances de tous les liens actifs, le graphe de dépendance de flots et l'algorithme de Fair Queuing : Start-time Fair Queuing utilisé pour assigner les tags des paquets qui arrivent.

3.3.1 Modèle

Nous considérons dans un réseau Ad hoc des échanges des paquets entre les nœuds. Ces nœuds partagent un canal ou plusieurs canaux. Les transmissions sont localement diffusées et seulement les récepteurs qui sont dans la zone de transmission de l'expéditeur peuvent recevoir la transmission.

3.3.2 Matrice et Graphe de dépendance de flots

Nous utilisons le graphe des nœuds (Figure 3.1) pour représenter la topologie du réseau. Nous dérivons la matrice et le graphe (table 1) de dépendances de flots (Figure 3.2) à partir du graphe de nœud.

La matrice de dépendance contient les dépendances de tous les liens actifs MD

$$MD = [f_{i,j}] \text{ pour } 1 \leq i, j \leq N \quad (3.1)$$

$f_{i,j}$ est le flot entre le nœud i, j

N est le nombre des nœuds

Dans la couche MAC avec un seul canal partagé, la matrice de dépendance est remplie par cette condition :

$$md_{x,y} = \begin{cases} 0 & \text{si } i=j \\ 1 & \text{si } i \text{ et } j \text{ sont au plus à 2 saut.} \\ 0 & \text{si } i \text{ et } j \text{ sont à partir de 3 sauts} \end{cases} \quad (3.2)$$

Dans la couche MAC avec la détection multi-usagers, la matrice de dépendance est remplie par cette condition :

$$md_{x,y} = \begin{cases} 0 & \text{si } f_i=f_j \\ 1 & \text{si } f_i \text{ et } f_j \text{ ne peut pas transmettre en même temps} \\ 0 & \text{autre} \end{cases} \quad (3.3)$$

Le graphe de dépendance de flots est dérivé à partir de la matrice MD. Il est défini comme $G = (V, MD)$, où V dénote l'ensemble de tous les flots (sommets) et l'arête (le lien (f_i, f_j)) appartient à MD si et seulement si les flots f_i et f_j sont en conflit l'un avec l'autre.

Exemple 1 (dans MAC avec un seul canal partagé):

La Figure 3.1 montre un exemple graphe des nœuds, le tableau 3.1 montre la matrice de dépendance et la Figure 3.2 montre son graphe de dépendance de flots. Dans cette Figure, 3.2 on remarque qu'il y a une dépendance entre les flots f_0, f_1, f_2, f_3, f_4 et une dépendance entre f_4, f_5 .

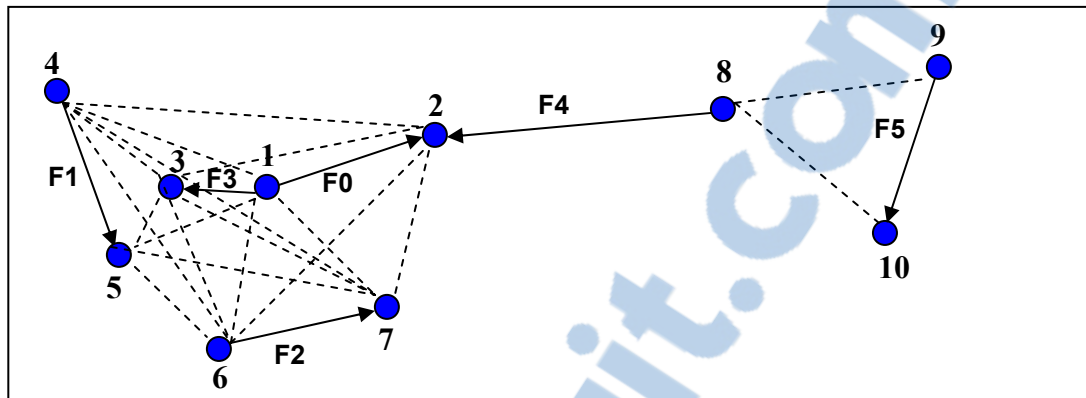


Figure 3.1 Graphe des noeuds exemple 1.

Tableau 3.1 Matrice de dépendance de flots exemple 1

	(1, 2)	(1, 3)	(4, 5)	(6, 7)	(8, 2)	(9, 10)
	F ₀	F ₃	F ₁	F ₂	F ₄	F ₅
F ₀	0	1	1	1	1	0
F ₃	1	0	1	1	1	0
F ₁	1	1	0	1	1	0
F ₂	1	1	1	0	1	0
F ₄	1	1	1	1	0	1
F ₅	0	0	0	0	1	0

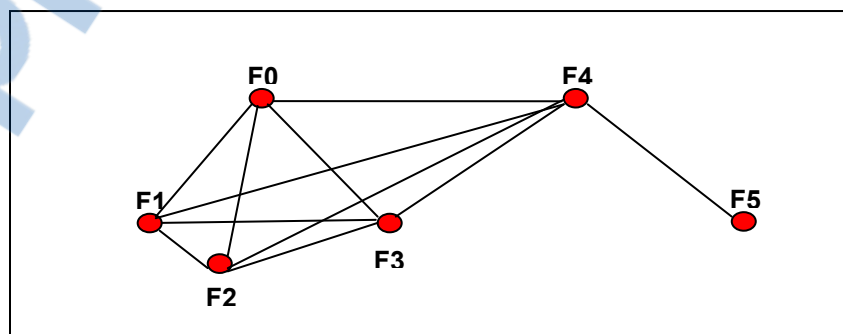


Figure 3.2 Graphe de dépendance de flots exemple 1.

Exemple 2 (dans la couche MAC avec une détection multi-usagers):

La Figure 3.3 montre un exemple de graphe des nœuds, le tableau 3.2 montre la matrice de dépendance et la figure 3.4 montre son graphe de dépendance de flots. Dans cette Figure, on remarque qu'il y a une dépendance entre les flots f_0 , f_3 , f_4 .

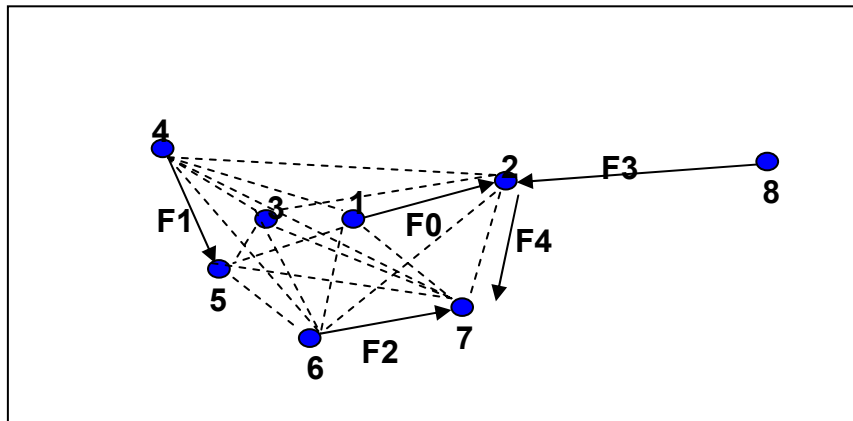


Figure 3.3 Graphe des nœuds exemple 2.

Tableau 3.2 Matrice de dépendance de flots exemple 2

	(1, 2)	(4, 5)	(6, 7)	(8, 2)	(9, 10)
	F_0	F_1	F_2	F_3	F_4
F_0	0	0	0	1	0
F_1	0	0	0	0	0
F_2	0	0	0	0	0
F_3	0	0	0	0	1
F_4	1	0	0	1	0

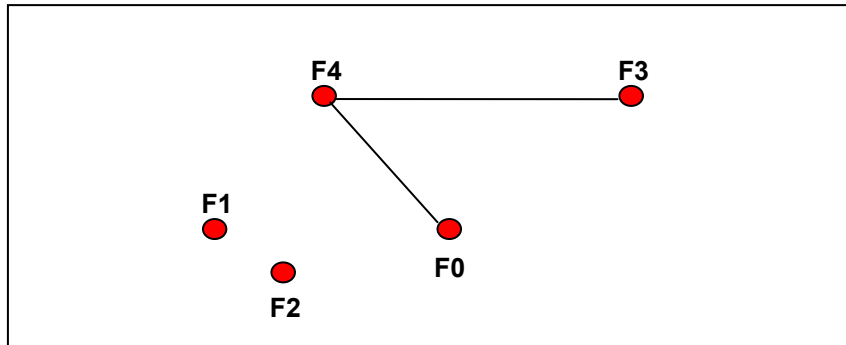


Figure 3.4 Graphe de dépendance de flots exemple 2.

3.3.3 Start-time Fair Queuing (SFTQ)

Beaucoup de recherches ont été réalisées sur les algorithmes d'ordonnancement équitable (fair queuing) pour accomplir une allocation équitable de bande passante sur un lien partagé. Les algorithmes d'ordonnancement équitable (fair queuing) dans la littérature essaient d'une manière caractéristique de se rapprocher de l'algorithme (GPS). Start-time Fair Queuing (SFTQ) est un parmi eux. SFTQ est vraiment efficace et alloue assez bien la bande passante de façon équitable sans tenir compte du contrôle d'admission et de la variation de serveur. STFQ est convenable pour notre modèle parce qu'il fournit l'équité, sans tenir compte de la variation dans la capacité de serveur.

Start-time Fair Queuing ordonnance les paquets dans l'ordre de tag de début. STFQ utilise deux étiquettes, un tag de début dénoté par $S(P_i^k)$ et un tag de fin dénotée par $F(P_i^k)$, qui sont associés à chaque paquet, où i dénote le nombre de flots k dénote le nombre de rond; les paquets sont ordonnés dans l'ordre des tags de début. Pour accomplir une la bande passante équitable avec SFQ, notre système alloue des tags de début en utilisant une horloge virtuelle; le temps virtuel $v(t)$ est défini comme la tag de début du paquet dans le service au temps t . Quand un paquet P_i^k arrive au temps t , il est étiqueté avec le tag de début. Le tag de début est calculé comme suivant :

Dans une période active : $S(P_i^k) = F(P_i^{k-1})$

Dans une période inactive : $S(P_i^k) = v(t)$

Au départ le temps virtuel du serveur est zéro. Pendant une période active, le temps virtuel $v(t)$ au temps t est défini pour être la tag de début du paquet en activité au temps t , $v(t) = S(P_i^k)$.

À la fin de la période occupée, le temps virtuel est mis au maximum du tag de fin alloué à n'importe quels paquets qui ont été servis d'ici là,

$$v(t) = \max_i^k F(P_i^k)$$

Le temps de fin est calculé comme :

$$F(P_i^k) = S(P_i^k) + L_i / W_i$$

Où L_i est la taille du paquet de flot i et W_i est son poids.

Les paquets sont servis dans l'ordre croissant des tags de début.

Dans cet exemple, nous expliquons le fonctionnement de SFQ. Il y a deux flots, F_1 , F_2 . $W_1=5$ est le poids de F_1 , $W_2=2$ est le poids de F_2 . La taille du Paquet =10.

Le début et fin du tag de chaque paquet sont :

$$S(P_1^1) = 0, F(P_1^1) = 10/5 = 2;$$

$$S(P_2^1) = 2, F(P_2^1) = 2 + 10/5 = 4;$$

$$S(P_3^1) = 4, F(P_3^1) = 4 + 10/5 = 6;$$

$$S(P_1^2) = 0, F(P_1^2) = 10/2 = 5;$$

$$S(P_2^2) = 5, F(P_2^2) = 5 + 10/2 = 10;$$

$$S(P_3^2) = 10, F(P_3^2) = 10 + 10/2 = 15;$$

Dans la Figure 3.5, on peut voir que l'ordre de transmission est : $P_1^1, P_2^1, P_3^1, P_1^2, P_2^2, P_3^2$.

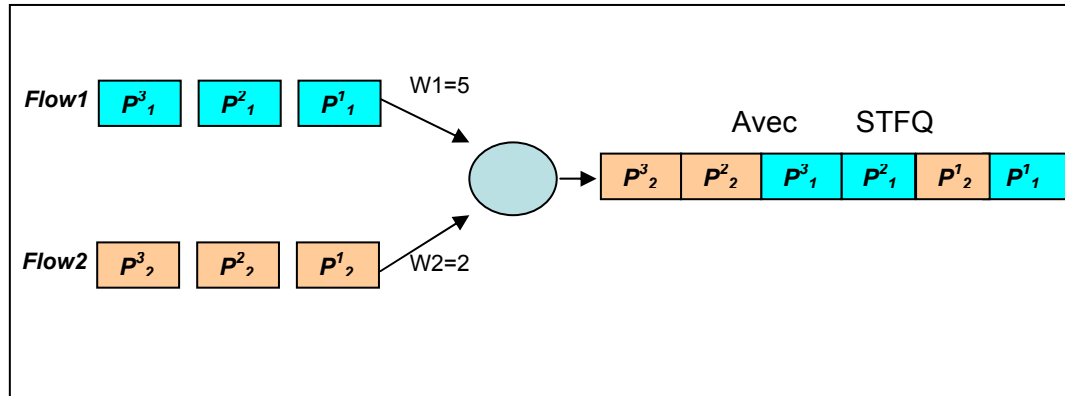


Figure 3.5 Exemple de STFQ.

3.3.4 Construction du regroupement

Un regroupement est un ensemble des liens (flots) qui peuvent transmettre en même temps. Avec la matrice de dépendance MD, on peut construire tous les regroupements possibles dans un voisinage d'un réseau.

On note que :

k est la cardinalité d'un regroupement.

On peut obtenir plusieurs regroupements pour une même MD. On note GR_K^Z le $Z^{\text{ème}}$ regroupement avec k cardinalité.

Chaque flot ($l_{i,j}$) a un paquet à transmettre pour chaque slot. Un tag de début dénoté par $S(P_{i,j})$ et un tag de fin dénoté par $F(P_{i,j})$, qui sont associés à chaque paquet, où i dénote le transmetteur et j dénote le récepteur; les paquets sont ordonnés dans l'ordre des tags de début.

Chaque regroupement GR_K^b a une Tag minimale Tag_K^Z :

$$Tag_K^Z = \min_{(i,j) \in GR_K^Z} S(P_{i,j}) \quad \text{Où} \quad l_{i,j} \in AL \quad (3.4)$$

Le tag désigne la priorité de chaque flot f , le flot avec le tag le plus petite a la priorité de transmettre le premier et le regroupement qui appartient ce flot a le minimum tag Tag_K^Z parmi d'autres regroupements.

La fonction Card désigne le nombre de flot de chaque regroupement.

$$Card \ GR_k^z = k \quad (3.5)$$

Par exemple dans l'exemple de la Figure 3.1 et 3.2 on peut avoir deux regroupements possibles : $GR_3^1 = (f1, f2, f4)$ et $GR_4^2 = (f0, f1, f2, f3)$

3.3.5 Gestion du trafic

Dans ce modèle, on peut avoir plusieurs regroupements mais la question qui se pose, comment peut on faire l'ordonnancement?

Il y a plusieurs possibilités. On a choisi deux possibilités. Avec la première, l'ordonnancement se fait avec la priorité donnée au regroupement qui a le tag le plus petit. La deuxième possibilité est de faire l'ordonnancement selon la cardinalité (Card) du regroupement. Ainsi, le regroupement avec le plus de flots sera transmit en premier.

On présente par la suite un algorithme qui se base sur la première proposition où l'ordonnancement se fait avec la priorité donné au regroupement qui contient le tag le plus petit.

On définit le trafic T ou l'ensemble des regroupements peut avoir cette condition :

$$\bigcup_{GR_k^z \in T} GR_k^z = AL \quad (3.6)$$

La condition (3.6) garantie que tous les liens actifs sont activés durant un cycle.

T représente l'ensemble de tous les regroupements possibles ordonnés selon le tag le plus petit afin d'obtenir un ordonnancement équitable ou selon leur cardinalité afin d'avoir plus de débit.

Dans l'algorithme suivant, on a pris la première priorité qui se base sur la priorité (tag) afin d'avoir un ordonnancement équitable. On peut avoir le même algorithme si on choisit la deuxième proposition qui se base sur la cardinalité du regroupement.

Algorithme:

On définit l'ensemble des flux (flow set), $\text{flow-set} = M = \{\text{tout les flots actives}\}$ et l'ensemble de dépendance 'dependence set', $\text{dep-set} = \{\text{flot en conflit avec d'autre flot}\}$ et n est le nombre des flots actives. On note f_i en conflit avec f_j par $f_i \neq f_j$.

L'ensemble dep-set est dérivé de la matrice de dépendance :

Au début, $\text{dep-set} = \{ \emptyset \}$,

Parcourir toute la matrice de dépendance :

si on trouve $d_{ij} = 1$ alors $f_i \neq f_j$

et $\text{dep-set} = \text{dep-set} \cup \{(f_i \neq f_j)\}$

Exemple: à partir de la table 3.2, $\text{dep-set} = \{(f_0 \neq f_4), (f_3 \neq f_4)\}$

Tout les cliques possible sont créés par toutes les combinaisons à partir du flow-set et dep-set .

A partir de la table 2 et avec l'algorithme défini au dessus, on peut obtenir 2 cliques possible: $\text{FC} = \{C_1, C_2\}$, $C_1 = (f_0, f_1, f_2, f_3)$ and $C_2 = (f_1, f_2, f_4)$.

3.3.6 Comment satisfaire les deux critères équité & débit en même temps?

Le modèle qu'on propose sert à fournir l'équité dans réseau Ad hoc, mais il affecte le débit. Le paramètre débit est très important pour la performance de réseau. L'intégration du paramètre débit dans la conception de notre modèle avec le paramètre équité, nous donne un ensemble de combinaisons par exemple dans la Figure 3.6, on a deux regroupements $\text{GR}_3^1 = (f_1, f_2, f_4)$, $\text{GR}_4^2 = (f_0, f_1, f_2, f_3)$ et deux actions (équité et débit). Le nombre de combinaisons

est quatre (4), parmi eux il y a la solution optimale qui peut satisfaire les deux avec les utilités demandée (équité et débit). D'après la littérature, la théorie des jeux peut être très utile pour trouver la solution optimale de ce type de problème.

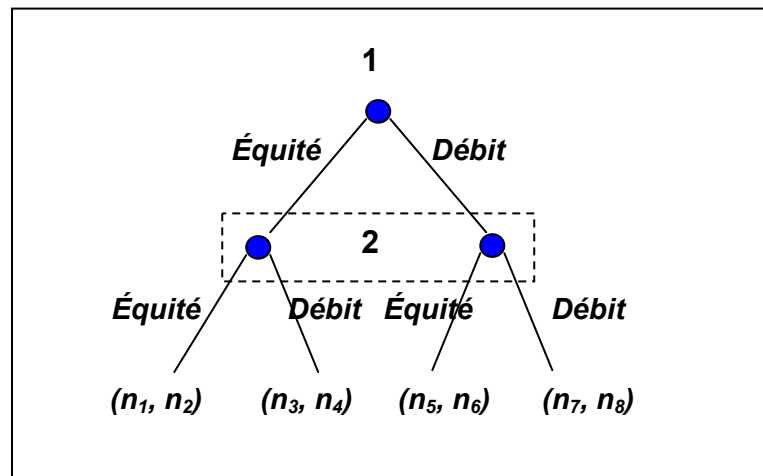


Figure 3.6 Arbres de jeux.

3.3.7 Théorie des jeux

On distingue deux branches dans la théorie des jeux [25], appelées jeux coopératifs et jeux non-coopératifs. L'idée générale dans les deux cas est qu'un certain nombre de joueurs participent à un jeu (à prendre ici dans le sens large du terme). Chacun peut entreprendre un certain nombre d'actions et l'on s'intéresse aux éventuels équilibres qui sont issus de leurs comportements. Dans les deux cas, on associe des utilités aux joueurs, qui représentent l'attrait relatif des différentes actions.

Les problèmes de cette deuxième famille impliquent, comme leurs noms l'indiquent, une coopération entre les joueurs (ce qui nous intéresse dans notre recherche). Ces derniers sont alors désignés dans la littérature sous le nom de 'bargainer'. Le résultat d'un marchandage porte le nom d'équité.

Plusieurs critères se trouvent dans la littérature. Ils ont d'abord été traités dans le cas de deux joueurs puis étendus pour la plupart à un nombre quelconque de joueurs. Un exemple typique

est la négociation autour d'un prix.

Le concept de Nash Bargaining Solution (NBS) est issu de la théorie des jeux coopératifs. Il est défini par un ensemble d'axiomes qui sont intéressants en matière d'équité. Ces axiomes s'appliquent à des utilités associées aux joueurs.

3.3.8 Formulation de la combinaison équité et débit avec la théorie des jeux

Dans cette section on décrit les définitions et les méthodes de la formulation de notre étude avec la théorie des jeux pour combiner l'équité et le débit afin de garantir la performance du réseau Ad hoc.

La stratégie du jeu

On peut présenter la stratégie du jeu comme un ensemble G :

$$G = \langle N, \{A_j\}_{j \in N}, \{U_j\}_{j \in N} \rangle \quad (3.6)$$

Où

- N est un ensemble de 2 joueurs ou plus ($N = \{1, 2, \dots, n\}$).
- A_i est un ensemble d'actions pour chaque joueur (dans notre cas, nous avons deux actions l'équité & débit).
- U_i est une fonction utilité (objectif) pour chaque joueur.

L'objectif de cette stratégie de jeux est la coopération entre les joueurs afin d'avoir une solution optimale qui arrange leur choix stratégique.

La fonction caractéristique

La fonction caractéristique décrit les possibilités de récompense pour chaque coalition (entre l'équité et le débit dans notre cas). Cette récompense (payoff), $V(C)$, représente l'utilité (U_i) totale que la coalition C peut obtenir quand leurs membres coopèrent.

Le domaine de négociation (Bargaining domain) et l'optimalité de Pareto

Dans les jeux coopératifs, le résultat final est l'intérêt principal [6]. Il est souvent convenable d'analyser le domaine de tous les résultats possibles, U , que l'on appelle domaine de négociation. Le résultat du jeu est défini par les valeurs d'utilités de joueurs (dans notre cas les joueurs sont les utilisateurs qui veulent avoir une équité dans le service et l'opérateur qui veut augmenter le débit) $u = \{u_1, \dots, u_n\} : u_i \in \mathbb{R}$. En général chaque utilité peut être exprimée dans de différentes unités. Nous considérons les domaines de négociations qui sont convexes, fermés et les ensembles bornés de $\mathbb{R}^n : \mathbb{R} \geq 0$. Un exemple avec un domaine où $n = 2$ est donné dans la Figure 3.7 (u_1 =équité et u_2 =débit). Le domaine de négociation dans notre cas est tous les regroupements possibles.

Une solution optimale Pareto est définie comme une solution qui garantit qu'il n'y a pas une autre solution dans laquelle un joueur peut augmenter son utilité sans affecter les autres. Dans la Figure 3.7 où toutes les solutions juste sur la bordure supérieure du graphe sont les solutions optimales de Pareto et elles forment la bordure Pareto. Le résultat du jeu dépend aussi du point de départ, $s = \{s_1, \dots, s_n\} : s \in U$, que l'on appelle quelques fois le point de désaccord ou le conflit. Le point de départ correspond à une hypothèse initiale que le résultat du jeu, u^* , ne peut pas être plus mauvais pour aucun joueur que le point de départ S , c'est à dire $u^* \geq s$. Le point de départ peut limiter le domaine de négociation (les regroupements possibles) comme il est montré dans la Figure 3.7.

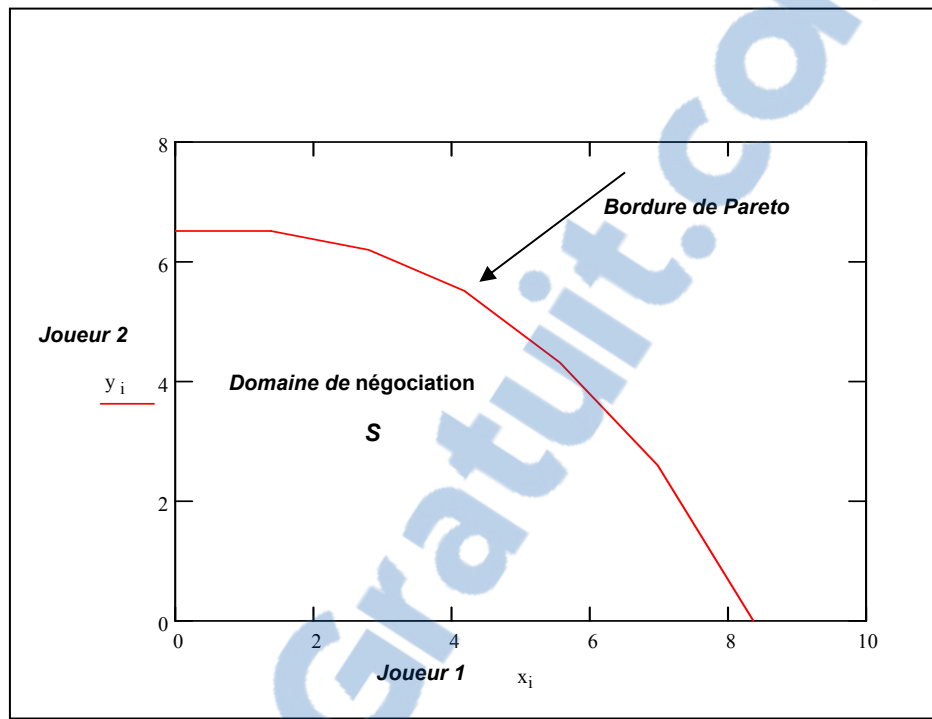


Figure 3.7 Exemple de domaine de négociation de Pareto.

Solutions optimales avec la théorie des jeux

Les solutions optimales sont fondée sur la maximisation du produit de fonctions de préférence de chaque joueur (préférence $\in \mathcal{U} = \{u_1, \dots, u_n\}$). Cette fonction est associée avec la solution de Nash.

À l'origine l'arbitrage de Nash a été défini par quatre axiomes :

- Symétrie : si le domaine de négociation est symétrique, $u_1 = u_2$ et le point de départ (s) sont sur cet axe, donc la solution est aussi sur cet axe.
- L'optimalité de Pareto : la solution est sur la bordure de Pareto.
- Invariance en respectant les transformations utilitaires : la solution pour n'importe quelle une transformation positive affine (U, s) dénoté par $V(U)$, $V(s)$ est $V(u^*)$ où u^* est la solution pour le système original.

- Indépendance des alternatives hors des propos : si la solution pour (U_1, s) est u^* et $U_2 \subseteq U_1$, $u^* \in U_2$, donc u^* est aussi la solution pour (U_2, s) . Autrement dit cet axiome montre que si U_1 est réduit, la solution originale et le point de départ sont toujours inclus et la solution du nouveau problème reste la même.

Il est important que la solution Nash puisse être aussi exprimée comme la maximisation du produit d'utilités normalisées

$$\max_{u \in U} \{u'_1 \cdot u'_2\} \quad (3.7)$$

où u'_i est l'utilité normalisée par sa valeur maximum dans le domaine utilitaire U . Ici on peut traiter les utilités normalisées comme la fonction privilégiée des joueurs (dans notre cas les joueurs sont les utilisateur qui veulent avoir une même priorité dans le service et l'opérateur qui veut augmenter le débit) qui signifie que dans la solution Nash chaque joueur est concerné seulement avec sa propre augmentation.

Cao [6] a étendu cette formulation à une famille de fonctions privilégiées qui tient aussi compte de l'augmentation de l'autre joueur. Cette famille est définie comme

$$\begin{aligned} v_1 &= u'_1 + \beta(1 - u'_2) \\ v_2 &= u'_2 + \beta(1 - u'_1) \end{aligned} \quad (3.8)$$

Lorsque $\beta = [-1, 1]$ la fonction privilégiée de joueur peut être défini comme une classe de solutions arbitrées :

$$u^* = \arg \left(\max_{u \in U} \{v_1 \cdot v_2\} \right) \quad (3.9)$$

Nous remarquons que lorsque :

- $\beta = 0$ la solution u^* correspond à l'arbitrage de Nash (voir Figure 3.8).
- $\beta \neq 0$ la fonction privilégiée du joueur tient compte de l'utilité de l'autre joueur.

Et en particulier si :

- $\beta = 1$ la fonction privilégiée traite avec le même poids l'augmentation du joueur et les pertes de l'autre joueur donc la solution égale les utilités normalisées des deux joueurs (Figure 3.7).
- $\beta = -1$ les deux augmentations ont le même poids donc la solution maximise la somme d'utilités normalisées (Figure 3.8).
- Ces deux solutions correspondent aux projets bien connus. La première solution ($\beta = 1$) est équivalent à la solution Raiffa-Kalai-Smorodinsky est défini comme l'égalité entre les utilités normalisées. La deuxième solution ($\beta = -1$) est équivalent à la solution modifiée de Thomson qui est défini comme la maximisation de la somme des utilités normalisées.

Il est aussi important de mentionner que dans l'espace privilégié, $\{V_1, V_2\}$, toutes les solutions mentionnées deviennent des solutions Nash.

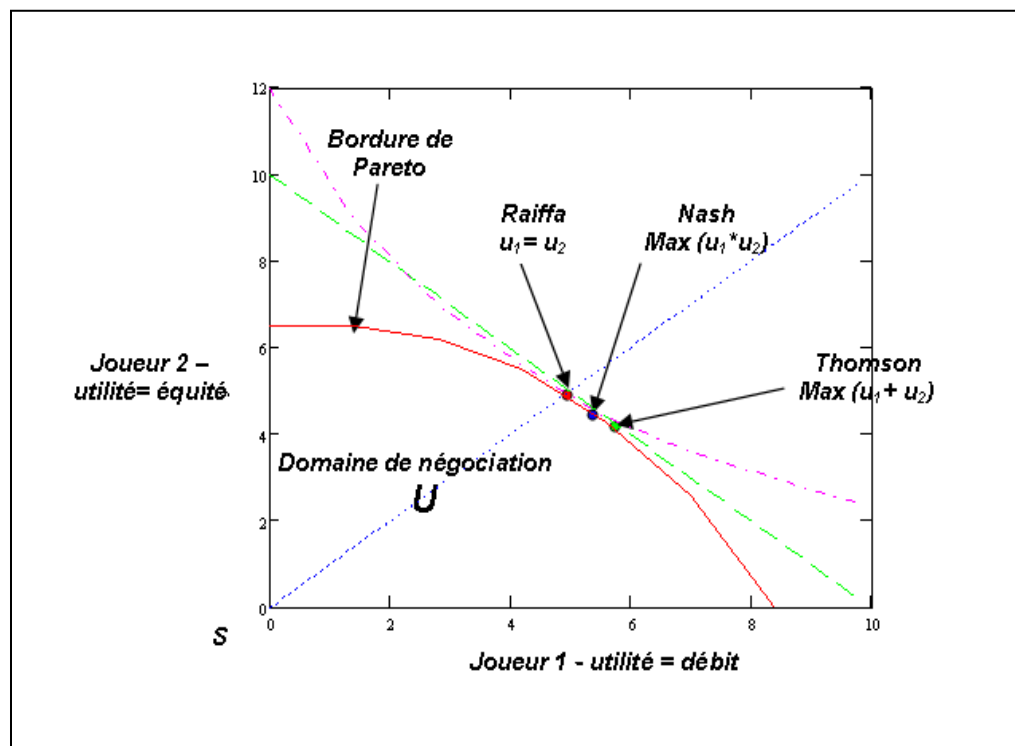


Figure 3.8 Solutions optimales de la théorie des jeux.

3.4 Les différentes objectives de l'ordonnancement et leurs implémentations

Dans cette section on décrit les différentes objectives de l'ordonnancement et leurs implémentations basées sur la plate-forme décrite dans la section précédente. Pour les formulations de différentes objectives de l'ordonnancement que nous présentons, le premier modèle se base sur STFQ. Le deuxième modèle est basé sur la priorité de délai d'attente (TOP) utilisé dans [2]. Le troisième modèle se base sur la maximisation de débit. Le quatrième et cinquième modèle se base sur les multi-objectifs. Ils comprennent le modèle basé sur l'arbitrage Nash (maximisation du produit de fonctions de préférence) et le modèle basé sur la solution de Thomson (maximisation de la somme de fonctions de préférence).

3.4.1 Ordonnancement basé sur l'algorithme 'Start time Fair Queuing (STFQ)'

L'implémentation basée sur l'algorithme STFQ est divisée en deux parties: le marquage (tagging) et l'ordonnancement. Le marquage maintient un retard sur le montant de service reçu par chaque flux en fonction de l'algorithme présenté dans la section 2. Après, dans chaque nœud, le mécanisme de programmation distribuée, décrit dans la section 2, sélectionne un clique qui contient le plus petit tag afin de préserver la planification équitable entre tous les flux, le clique sélectionnée définit la fonction du nœud (récepteur ou émetteur) et les paquets (flux) doivent être envoyés dans le suivant slot de transmission de données.

3.4.2 Ordonnancement basé sur 'Time out priority' (TOP)

La même métrique décrit dans [2] est utilisée dans notre modèle 'timeout priority' (TOP). Dans chaque nœud le mécanisme de programmation distribuée sélectionne un clique qui contient la plus petite valeur du 'timeout' avant que le paquet soit détruit afin de préserver la planification équitable entre tous les flux.

3.4.3 Ordonnancement basé sur le débit max ‘Throughput maximization’ (TM)

Le nombre maximum de flots, T , dans la clique est utilisé comme une métrique dans le modèle d’ordonnancement basé sur la maximisation du débit.

Dans chaque nœud le mécanisme de programmation distribuée sélectionne une clique qui contient le nombre maximum de flot dans chaque slot de transmission.

3.4.4 Ordonnancement basé sur l'arbitrage Nash (NASH)

Dans la théorie des jeux coopératives, la solution de l’arbitrage de Nash [6] appelée aussi NBS ‘Nash Bargaining Solution’ fournit une solution équitable entre les joueurs afin d’avoir une solution optimale qui arrange leur choix stratégique. Cette solution peut être réalisée par la maximisation du produit d'utilités normalisées. On applique ce concept pour faire une balance entre la satisfaction de la performance du délai de l'utilisateur et la satisfaction du débit de fournisseur du réseau.

Dans ce contexte l'inverse du plus petit ‘start tag’ (modèle STFQ) est utilisé comme la fonction d'utilité pour la satisfaction du délai pour chaque clique :

$$D' = 1 / S(P_i^k) \quad (3.8)$$

On définit la fonction d'utilité pour la satisfaction du débit le nombre maximum de flots T , pour chaque clique.

A la fin, la clique avec le produit des utilités maximum est sélectionnée :

$$\text{Max } (T * D') \quad (3.9)$$

3.4.5 Ordonnancement basé sur la maximisation de la somme des fonctions de préférence (Weighted sum 'WSUM')

Dans cette approche, on utilise les mêmes fonctions d'utilités, délai et débit définies dans l'ordonnancement basé sur l'arbitrage de Nash.

Néanmoins, dans ce cas, au lieu de maximiser le produit, nous sélectionnons la clique qui maximise la somme des utilités:

$$\text{Max } (T + \alpha D')$$
(3.10)

Cette formule peut être considérée comme une solution de Thomson [26] dans le cadre de la théorie des jeux coopératifs.

CHAPITRE 4

IMPLÉMENTATION ET ÉTUDE DES PERFORMANCES

4.1 Introduction

Dans ce chapitre, nous comparons et analysons la performance des algorithmes d'ordonnancements qu'on a présenté dans le chapitre 3 (TOP, STFQ, TOP', TM, NASH, WSUM).

Les résultats numériques sont obtenus à partir d'une simulation des événements discrets des modèles des réseaux Ad hoc avec les paramètres présentés dans le tableau 3.

Au début, les nœuds sont distribués aléatoirement dans un cercle. Le modèle de mobilité utilisé imite le comportement d'un homme et d'un véhicule en mouvement [27]. Le taux d'arrivée des paquets voix représente 20% du taux d'arrivée total.

Tableau 4.1 Les paramètres de simulation

Gain d'épandage	128
Modélisation de la zone du cercle	1000 m
La limite de la vitesse	50km/h
Maximum trans. energie	7w
Types de trafic	Voix & donnée
Le nombre de noeuds	60
Temps de simulation	100000 trames

4.2 Étude de performances

La moyenne du délai des paquets voix et le débit total (voix + paquets de données) sont deux mesures de performances sélectionnées pour la comparaison analytique. Le débit est défini comme le nombre total de paquets transmis au cours de la simulation ou le nombre de paquets par trame par nœud. La moyenne des délais des paquets voix est définie comme la moyenne des délais dans la queue de chaque nœud. Elle est exprimée en unités par la durée de la trame (10 ms).

Figures 4.1, 4.2, 4.3 et 4.4 décrivent les caractéristiques de performance en fonction du facteur de taux de génération des paquets, p . La capacité MUD limite le nombre de signaux CDMA reçus simultanément par conséquent, elle limite également le nombre de transmission des voisins de chaque nœud.

Figure 4.1 compare la moyenne des délais des paquets voix par rapport au facteur de la charge de trafic, p , et la figure 4.2 présente le délai par rapport à la capacité de MUD M pour $p = 0,8$ qui correspond à un réseau chargé.

De la même figure 4.3 compare le débit total (exprimé en nombre de paquets par trame par nœud) par rapport à la charge du trafic P et la figure 4.4 présente le débit (exprimé en nombre de bits par seconde par nœud). Par rapport à la Capacité MUD M pour $p = 0,8$. Les résultats sont présentés pour les 5 algorithmes d'ordonnancement: TOP, TOP', STFQ, TM, et Nash. Voici les principales conclusions de l'analyse des résultats:

- A. Le modèle TOP' donne une petite amélioration dans le délai de la voix par rapport au modèle TOP tandis que le débit total est comparable dans les deux cas. Cela indique que la connaissance de voisinage limité de l'état du réseau dans le modèle TOP ne détériore pas de manière significative les performances.

- B. La comparaison des méthodes basées sur la plate-forme présenté dans le chapitre 3 confirme qu'il y a un compromis clair entre le délai de la voix et le rendement du débit total. En d'autres termes l'amélioration de la mesure de délai est toujours faite au détriment du débit et vice versa.
- C. Le modèle STFQ apporte une performance significative au délai de la voix par rapport modèle TOP spécialement pour un réseau chargé. En particulier, pour $p = 0,8$ le délai dans le modèle STFQ a été réduit de 11%, tandis que dans le cas de $p = 1,6$ la réduction des délais a atteint 47%. Bien que ces gains sont réalisés au détriment de réduction de débit moyen ($\sim 5\%$), ces résultats indiquent que la STFQ protège très bien le trafic voix, même dans des grandes charges qui n'est pas le cas pour d'autres modèles.
- D. Comme prévu, le modèle TM fournit les meilleures performances du débit total, pour $p=0,8$, il dépasse le modèle TOP de $\sim 10\%$ et le modèle STFQ de $\sim 14\%$. Néanmoins, ces gains sont réalisés au détriment d'une augmentation sensible du délai de la voix qui est de l'ordre de 12% et 25%, respectivement.
- E. Comme prévu aussi, le modèle multi-objectif NASH offre un compromis entre la minimisation et la maximisation de délai voix et le débit. Il convient de noter que par rapport à TOP et TM ce compromis est une très bonne opération puisque le gain de débit avec le modèle TOP dépasse en comparant avec l'augmentation du délai.

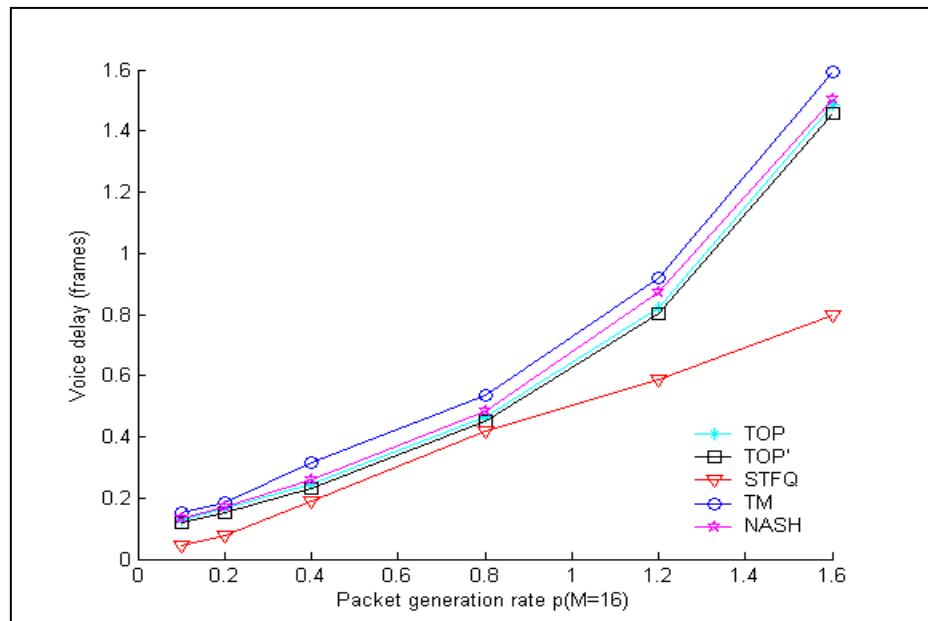


Figure 4.1 Délai des paquets voix des différents modèles d'ordonnancements 1.

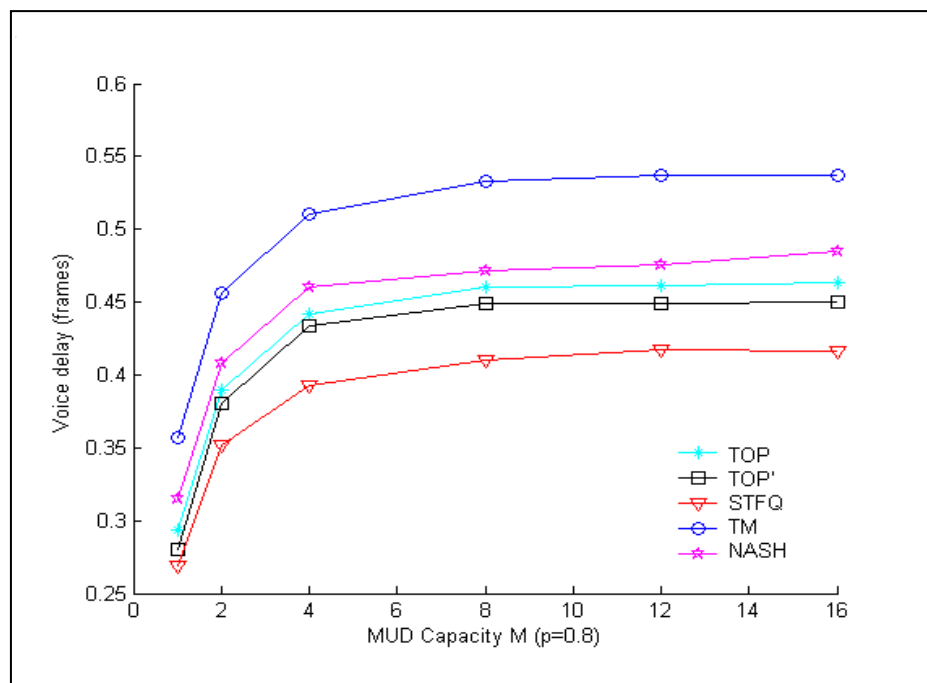


Figure 4.2 Délai des paquets voix des différents modèles d'ordonnancements 2.

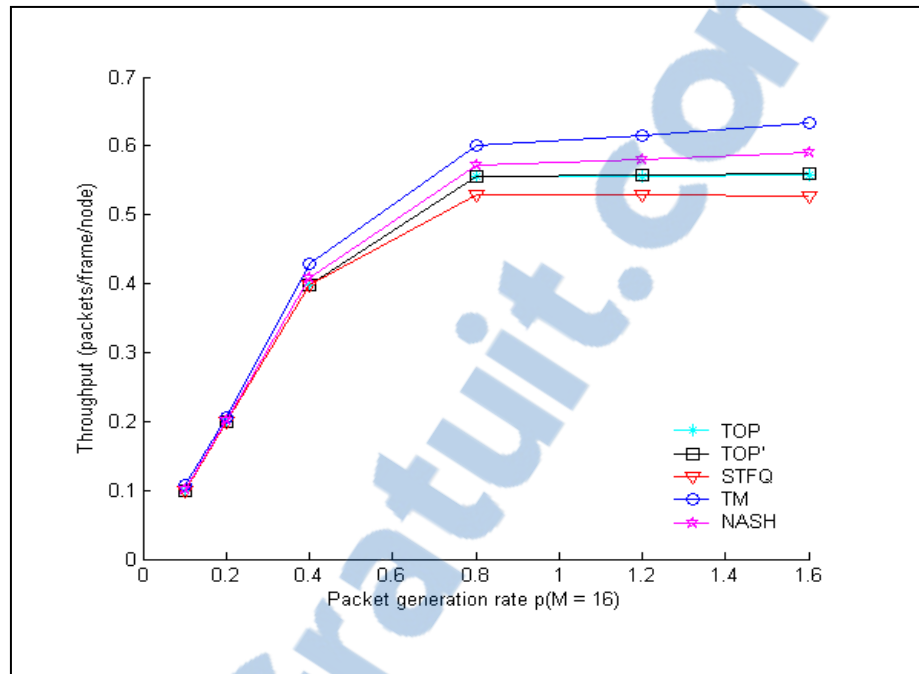


Figure 4.3 Débit des différents modèles d'ordonnancements 1.

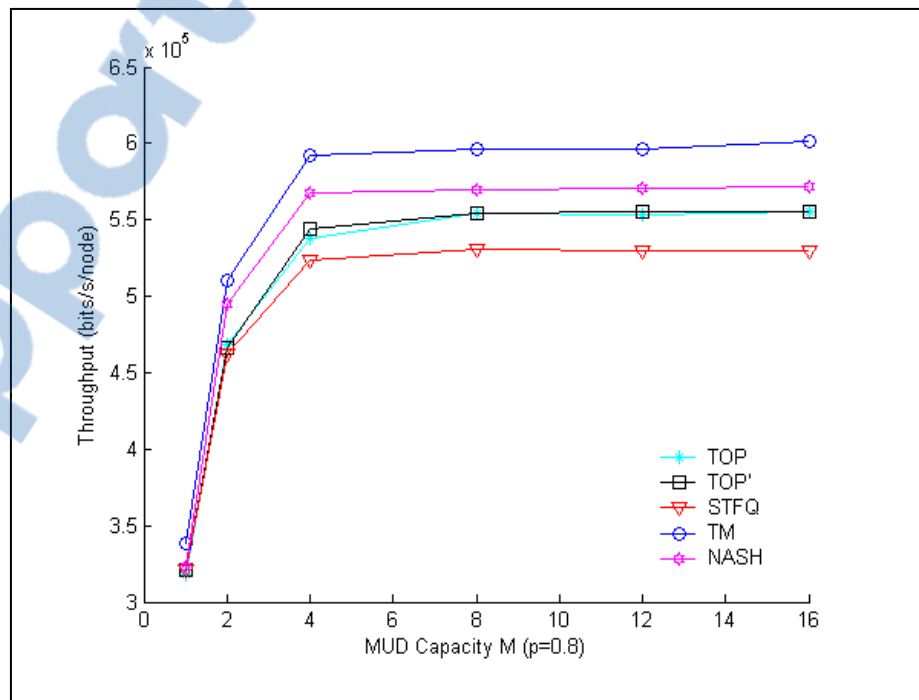


Figure 4.4 Débit des différents modèles d'ordonnancements 2.

La performance de la seconde fonction multi-objective, WSUM, est illustrée dans les figures 4.5, 4.6, 4.7, 4.8 avec le paramètre de poids différents. On compare le modèle WSUM avec les objectifs pertinents: STFQ, Nash et TM. Les résultats indiquent que le modèle WSUM donne la flexibilité de parvenir à un compromis arbitraire entre STFQ et TM. D'autre part la solution de Nash évite l'optimisation de la valeur du poids.

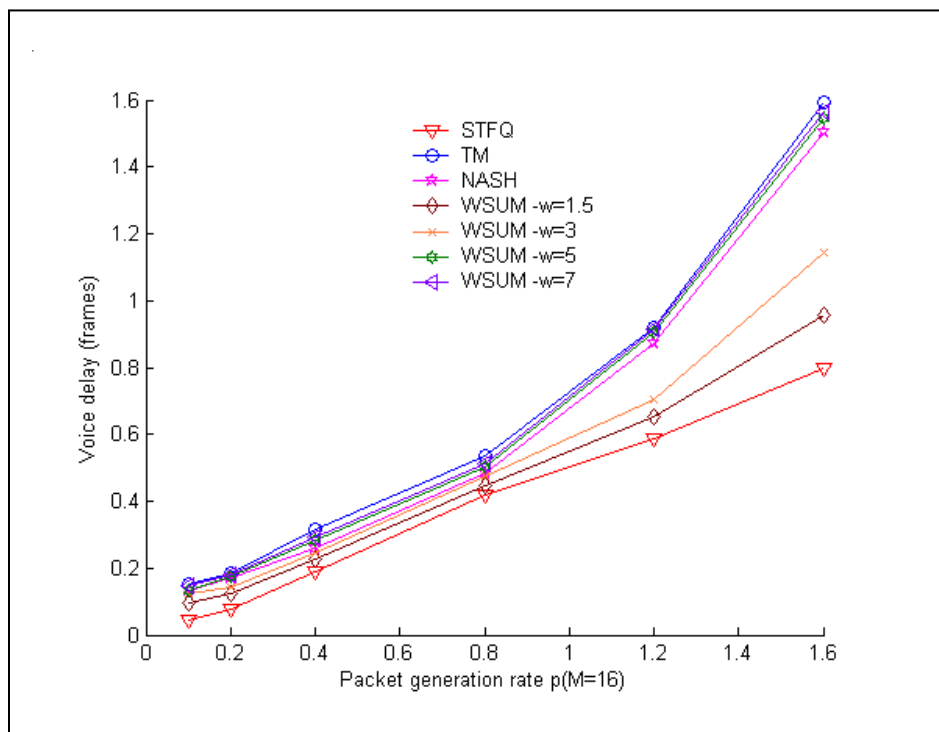


Figure 4.5 Délai des paquets voix du modèle d'ordonnancement WSUM 1.

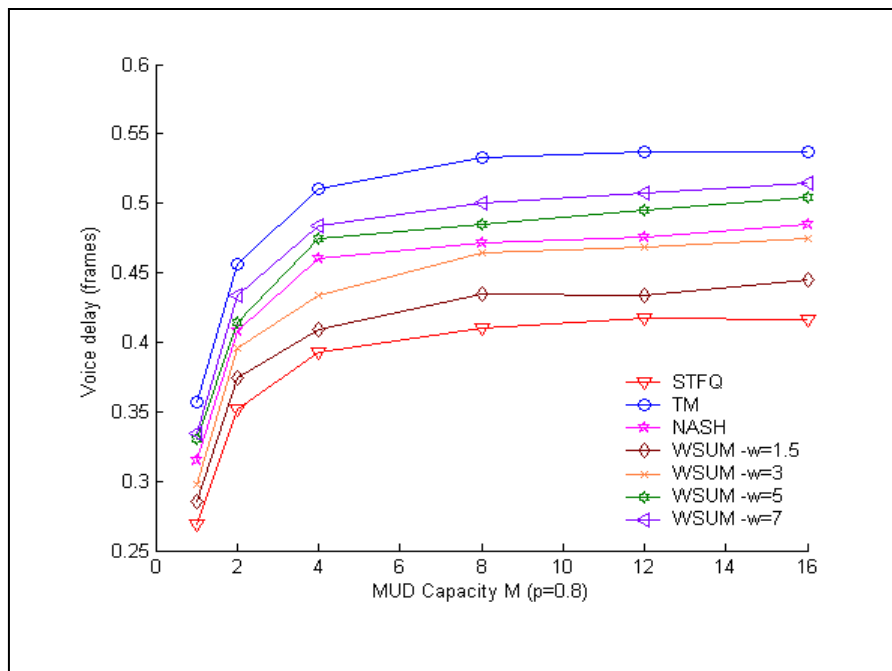


Figure 4.6 Délai des paquets voix du modèle d'ordonnancement WSUM 2.

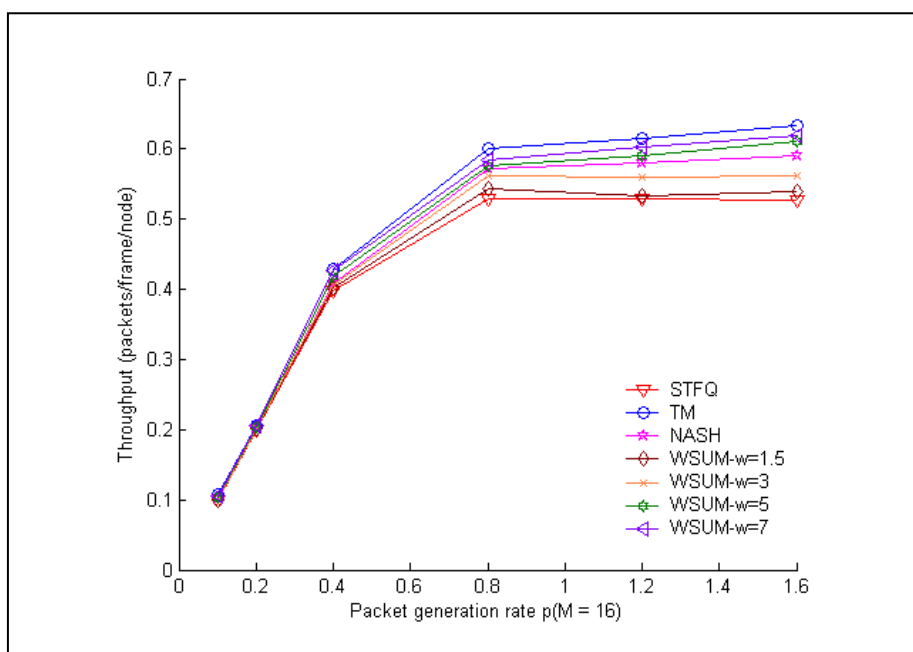


Figure 4.7 Débit du modèle d'ordonnancement WSUM 1.

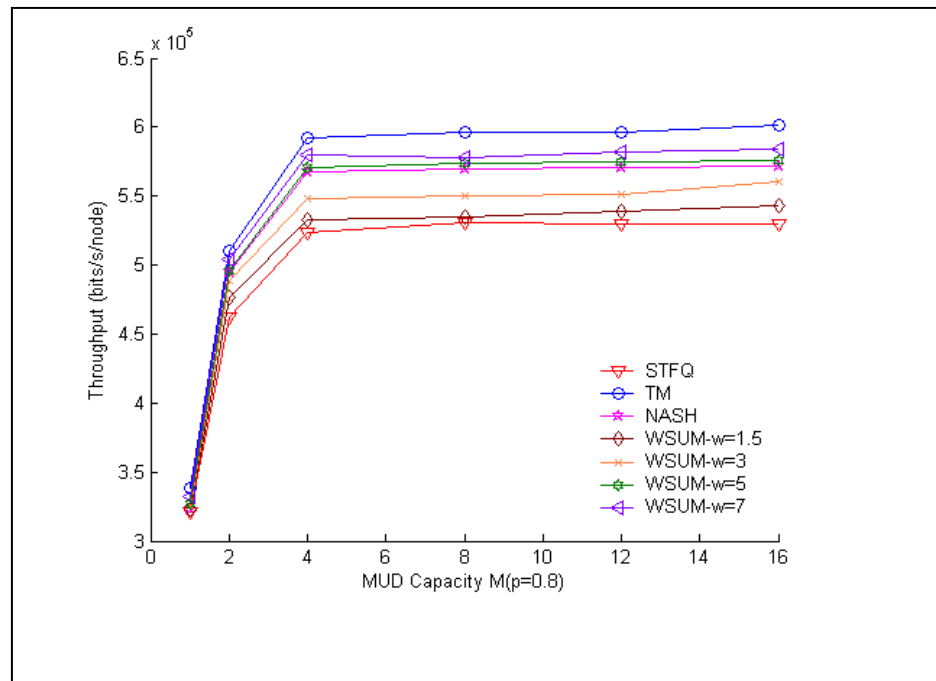


Figure 4.8 Débit du modèle d'ordonnancement WSUM 2.

CONCLUSION

Dans ce projet, nous avons abordé un sujet très intéressant, l'ordonnancement équitable dans la couche MAC avec détection multi-usagers dans les réseaux Ad hoc. Comme indiqué dans [2], la détection multi-usagers peut donner des gains importants dans le débit et la performance de la QoS. Néanmoins, la réalisation de ces gains nécessite une optimisation d'un ordonnancement distribué dans le voisinage afin d'obtenir les performances souhaitées. Pour atteindre cet objectif, nous avons proposé une plate-forme qui permet d'analyser et de comparer de façon optimale ou semi-optimale le modèle d'ordonnancement distribués dans le voisinage avec des objectifs différents.

L'approche est basée sur le flot et des matrices de dépendance des flots qui sont utilisées pour créer un ensemble de configurations possibles d'ordonnancement (également dénommé cliques). Ensuite, une configuration est sélectionnée pour la transmission est basée sur un modèle d'ordonnancement choisi. Pour démontrer la viabilité de cette approche nous avons proposé plusieurs algorithmes d'ordonnancement avec différents objectifs reposant sur STFQ, la maximisation de débit et les priorités de délai. En outre nous avons également proposé et mis en œuvre deux fonctions multi-objectives dans le but de fournir un compromis entre la minimisation de délai et la maximisation de débit du trafic. L'une est basée sur l'arbitrage de Nash et l'autre sur maximisation de la somme de fonctions de préférence.

Les résultats numériques ont montré que la mise en œuvre de l'algorithme STFQ peut améliorer considérablement la moyenne du délai des paquets voix, en particulier dans des conditions d'une grande charge, au détriment de certaine réduction du débit total du réseau. Les résultats pour modèles de multi-objectif confirment leur capacité à fournir un compromis efficace entre des objectifs contradictoires.

Les futurs travaux sont en cours dans plusieurs directions. L'approche proposée suppose que chaque nœud reçoit les informations de d'ordonnancement de tous les autres nœuds, ce qui est difficile à atteindre dans chaque trame. Par conséquent, nous prévoyons d'augmenter le saut 'the hop' pour le protocole de signalisation développée dans [2] à un protocole en deux-hop qui devrait être suffisant pour l'approche proposée.

LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES

- [1] G. Holland and N. Vaidya, "Analysis of TCP performance over mobile ad hoc networks," in Proc. ACM Mobicom '99, Seattle, WA, 1999.
- [2] J. Zhang, Z. Dziong, F. Gagnon, M.Kadoch, "Multiuser Detection Based MAC Design for Ad Hoc Networks," accepted for publication in IEEE Transactions on Wireless Communications (May 5, 2008).
- [3] D. Tse and S. Hanly, "Linear multiuser receivers: Effective interference, effective bandwidth and user capacity," IEEE Transactions on Information Theory, vol. 45, pp. 641–657, March 1999.
- [4] R. Lupas and S. Verdu, "Linear multiuser detectors for synchronous code-division multiple-access channels," IEEE Transactions on Information Theory, vol. 35, pp. 123–136, January 1989.
- [5] Z. Xie, R. T. Short, and C. K. Rushforth, "A family of sub optimum detectors for coherent multi-user communications," IEEE Journal On Selected Areas in Communications, vol. 8, pp. 683–690, May 1990.
- [6] Z. Dziong, "ATM Network Resource Management, Appendix C, Cooperative Game Theory", 1997 McGraw-Hill.
- [7] K.Sarkar, C. Puttamadappa, T.G. Basavaraju,"Ad hoc Mobile Wireless Networks, Principles, protocols, and application," Auerbach Publications edition,2008.
- [8] C. N. Thurwachter JR " Wireless Networking, " Perentice Hall edition, 2002.
- [9] S. Mangold, G. Hiertz, and B. Walke, "IEEE 802.11e Wireless LAN - Resource Sharing with Contention Based Medium Access," The 14th IEEE 2003 International Symposium on Personal, Indoor and Mobile Radio Communication Proceedings, 2003.
- [10] B.S. Manoj and C.Siva Ram Murthy, 'Real-Time Traffic support for Ad Hoc Wireless networks', Proceedings of IEEE ICON 2002,pp. 335-340, august 2002.
- [11] B.John Nagle. On packet switches with infinite storage. IEEE Transactions on Communications, COM-35(4):435--438, avril 1987.

- [12] K. Abhay Parekh and Robert G. Gallager, "A generalized processor sharing approach to flow control in integrated services networks: the single-node case" IEEE/ACM Transactions on Networking (TON), vol. 1, no. 3, pp. 344-357, 1993.
- [13] Hui Zhang, Service disciplines for guaranteed performance service in packet-switching networks," Proceedings of the IEEE, vol. 83, pp. 1374-1396, 1995.
- [14] S.J. Golestani, A self-clocked fair queueing scheme for broadband applications," in proceedings INFOCOM'94, 1994, pp. 12-16.
- [15] P. Goyal, Harrick M. Vin, and Haichen Chen, Start-time fair queueing : a scheduling algorithm for integrated services packet switching networks," in Conference proceedings on Applications, technologies, architectures, and protocols for computer communications, 1996, pp. 157-168.
- [16] Songwu Lu, Vaduvur Bharghavan, and R. Srikant, Fair scheduling in wireless packet networks," IEEE/ACM Transactions on Networking (TON), vol. 7, no. 4, pp. 473-489, 1999.
- [17] T.S.E. Ng, I. Stoica, and H. Zhang, Packet fair queueing algorithms for wireless networks with location-dependent errors," in proceedings INFOCOM'98, 1998, vol. 3, pp. 1103-1111.
- [18] P. Ramanathan and P. Agrawal, Adapting packet fair queuing algorithms to wireless networks," in Proceedings of the fourth annual ACM/IEEE international conference on Mobile computing and networking, 1998, pp. 1-9.
- [19] N. Vaidya, Anurag Dugar, Seema Gupta, and Paramvir Bahl, Senior Member, IEEE , Distributed Fair Scheduling in a Wireless LAN, IEEE TRANSACTIONS ON MOBILE COMPUTING, VOL. 4, NO. 6, NOVEMBER/DECEMBER 2005.
- [20] H. Luo, P. Medvedev, J. Cheng, Lu. C, A Self-Coordinating Approach to Distributed Fair Queueing in Ad Hoc Wireless Networks, IEEE, 2005.
- [21] H. Luo, P. LU. S, A topology –Independent Wireless Fair Queueing in Ad Hoc Wireless Networks, IEEE Journal on Selected AREAS In communications, Vol. 23. No 3, March 2005.
- [22] H. Luo, Medvedev. P, Cheng. J, Lu. C, A Self-Coordinating Localized Fair Queueing in Ad Hoc Wireless Networks, IEEE TRANSACTIONS ON MOBILE COMPUTING, VOL. 3, NO. 1, January/March 2005.
- [23] N. Vaidya, Anurag Dugar, Seema Gupta, and Paramvir Bahl, Senior Member, IEEE , Distributed Fair Scheduling in a Wireless LAN, IEEE TRANSACTIONS ON MOBILE COMPUTING, VOL. 4, NO. 6, NOVEMBER/DECEMBER 2005.

- [24] H.Luo, Medvedev. P, Cheng. J, Lu. C, A Self-Coordinating Approach to Distributed Fair Queueing in Ad Hoc Wireless Networks, IEEE,2005.
- [25] G.Owen. Game theory. San Diego, CA : Academic, 1995.
- [26] J. Nash, “Two-person cooperative games,” *Econometrica*, vol. 21, pp. 128–140, January 1953.
- [27] J. Zhang and J. M. Mark and X. Shen,“An adaptive resource reservation strategy for handoff in wireless cellular CDMA networks”, *Can. J. Electr.Comput. Eng.*, vol. 29, no. 1/2, pp. 77–83, Jan./Apr. 2004.